

MetaNO: How to transfer your knowledge on learning hidden physics[☆]

Lu Zhang^a, Huaqian You^a, Tian Gao^b, Mo Yu^c, Chung-Hao Lee^{d,e}, Yue Yu^{a,*}

^a Department of Mathematics, Lehigh University, Bethlehem, PA, USA

^b IBM Research, Yorktown Heights, NY, USA

^c Pattern Recognition Center, WeChat AI, Tencent Inc, China

^d School of Aerospace and Mechanical Engineering, The University of Oklahoma, Norman, OK, USA

^e Department of Bioengineering, University of California, Riverside, Riverside, CA, USA

Available online 28 July 2023

Abstract

Gradient-based meta-learning methods have primarily been applied to classical machine learning tasks such as image classification. Recently, partial differential equation (PDE)-solving deep learning methods, such as neural operators, are starting to make an important impact on learning and predicting the response of a complex physical system directly from observational data. Taking the material modeling problems for example, the neural operator approach learns a surrogate mapping from the loading field to the corresponding material response field, which can be seen as learning the solution operator of a hidden PDE. The microstructure and mechanical parameters of each material specimen correspond to the (possibly heterogeneous) parameter field in this hidden PDE. Due to the limitation on experimental measurement techniques, the data acquisition for each material specimen is commonly challenging and may also be costly. This fact calls for the utilization and transfer of existing knowledge to new and unseen material specimens, which corresponds to sampling efficient learning of the solution operator of a hidden PDE with a different parameter field. Herein, we propose a novel meta-learning approach for neural operators that can be seen as transferring the knowledge of solution operators between governing (unknown) PDEs with varying parameter fields. Our approach is a provably universal solution operator for multiple PDE solving tasks, with a key theoretical observation that underlying parameter fields can be captured in the first layer of neural operator models, in contrast to typical final-layer transfer in existing meta-learning methods. As applications, we demonstrate the efficacy of our proposed approach on PDE-based datasets and a real-world material modeling problem, illustrating that our method can handle complex and nonlinear physical response learning tasks while greatly improving the sampling efficiency in unseen tasks.

© 2023 Elsevier B.V. All rights reserved.

Keywords: Operator-regression neural networks; Data-driven physics modeling; Neural operators; Transfer learning; Meta-learning; Scientific machine learning

[☆] We dedicate this paper to Prof. Thomas J.R. Hughes for his lifetime achievements, pioneering contributions to computational mechanics and applied mathematics, and distinguished editorship of Computer Methods in Applied Mechanics and Engineering for more than four decades.

* Corresponding author.

E-mail address: yuy214@lehigh.edu (Y. Yu).

1. Introduction

Few-shot learning is an important problem in machine learning, where new tasks are learned with a very limited number of labeled datapoints [1]. In recent years, significant progress has been made on few-shot learning using meta-learning approaches [2–12]. Broadly speaking, given a family of tasks, some of which are used for training and others for testing, meta-learning approaches aim to learn a shared multi-task representation that can generalize across different training tasks and result in fast adaptation to new and unseen testing tasks. Meta-learning learning algorithms have been successfully applied to the conventional machine learning problems such as image classification, function regression and reinforcement learning; however, studies on few-shot learning approaches for complex physical system modeling problems have been limited. The call of developing a few-shot learning approach for complex physical system modeling problems, such as to learn the material response model from experimental measurements, is just as acute, while the typical understanding of how multi-task learning should be applied on this scenario is still nascent.

As a motivating example, we consider the scenario of new material discovery under the lab setting, where the material model is built based on experimental measurements of the material's responses subject to various loads. Since the physical properties (e.g., the mechanical and structural parameters) in different material specimens vary, the model learnt from the experimental measurements on one specimen would have large generalization errors on other specimens. As a result, the data-driven model has to be trained repeatedly with a large number of material specimens, making the learning process inefficient. Furthermore, experimental measurement acquisition of these specimens is often challenging and expensive. In some problems, a large amount of measurements are not even feasible. For example, in the design and testing of biosynthetic tissues, performing repeated loading would potentially induce the cross-linking and permanent set phenomenon, which notoriously alter the tissue durability [13]. Therefore, it is critical to learn the physical response model of a new specimen with a sample size as small as possible. In addition, since many characterization methods to obtain underlying material mechanistic and structural properties would require the use of destructive methods [14,15], in practice many physical properties are not measured and can only be treated as hidden and unknown variables. In other words, we likely only have limited access to the measurements on the complex system responses caused by the change of these physical properties.

Supervised operator learning methods are typically used to address this class of problems. They take a number of observations on the loading field as input and try to predict the physical system response field as output, corresponding to one underlying PDE (as one task). Herein, we consider the meta-learning of multiple complex physical systems as tasks such that all these tasks are governed by a common PDE with different (hidden) physical property or parameter fields. Formally, assume that we have a distribution $p(\mathcal{T})$ over tasks, each task $\mathcal{T}^\eta \sim p(\mathcal{T})$ corresponds to a hidden physical property field $\mathbf{b}^\eta(\mathbf{x}) \in \mathcal{B}(\mathbb{R}^{d_b})$ that contains the task-specific mechanistic and structural information in our material modeling example. On task $\mathcal{T}^\eta$, we have a number of observations on the loading field $\mathbf{g}_i^\eta(\mathbf{x}) \in \mathcal{A}(\mathbb{R}^{d_g})$ and the corresponding physical system response field $\mathbf{u}_i^\eta(\mathbf{x}) \in \mathcal{U}(\mathbb{R}^{d_u})$ according to a hidden parameter field $\mathbf{b}^\eta(\mathbf{x})$. Here, i is the sample index, $\mathcal{B}$, $\mathcal{A}$ and $\mathcal{U}$ are Banach spaces of function taking values in $\mathbb{R}^{d_b}$, $\mathbb{R}^{d_g}$ and $\mathbb{R}^{d_u}$, respectively. For task $\mathcal{T}^\eta$, our modeling goal is to learn the solution operator $\mathcal{G}_{\mathbf{b}^\eta}^\eta : \mathcal{A} \rightarrow \mathcal{U}$, such that the learnt model can predict the corresponding physical response field $\mathbf{u}(\mathbf{x})$ for any loading field $\mathbf{g}(\mathbf{x})$.¹ Without transfer learning, one needs to learn a surrogate solution operator for each task only based on the data pairs on this task, and repeat the training for every task. The learning procedure would require a relatively large amount of observation pairs and training time for each task. Therefore, this physical-based modeling scenario raises a key question: *Given data from a number of parametric PDE solving (training) tasks with different unknown parameters, how can one efficiently learn an accurate surrogate solution operator for a test task with new and unknown parameters, with few data on this task?*²

To address this question, we introduce the meta-learned neural operator (MetaNO), a novel meta-learning approach for transferring knowledge between neural operators as illustrated in Fig. 1, which can be seen as transferring the knowledge of solution operators between governing (potentially unknown) PDEs with varying hidden parameter fields. Our **main contributions** are:

¹ Without ambiguity, in the following context we will neglect the subscript and denote $\mathcal{G}_{\mathbf{b}^\eta}^\eta$ as $\mathcal{G}^\eta$ for notation simplicity.

² In some meta-learning literature, e.g., [16], these small sets of labeled data pairs on a new task (or any task) is called the context, and the learnt model will be evaluated on an additional set of unlabeled data pairs, i.e., the target.

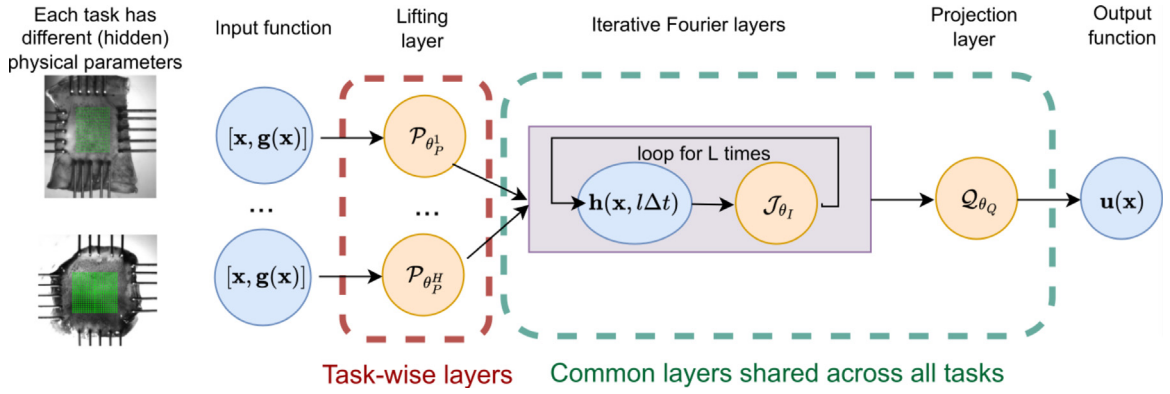


Fig. 1. The architecture of MetaNO based on an integral neural operator model.

- MetaNO is the first neural-operator-based meta-learning approach for multiple tasks, which not only preserves the generalizability to different resolutions and input functions from the integral neural operator architecture, but also improves sampling efficiency on new tasks—for comparable accuracy, MetaNO saves the number of measurements required by $\sim 90\%$.
- With rigorous operator approximation analysis, we made the key observation that the hidden parameter field can be captured by adapting the first layer of the neural operator model. Therefore, our MetaNO is *substantially different* from existed popular meta-learning approaches [5,10], since the later typically rely on the adaptation of their last layers [12]. By construction, MetaNO serves as a provably universal solution operator for multiple PDE solving tasks.
- On synthetic, benchmark, and real-world biological tissue datasets, the proposed method consistently outperforms existing non-meta transfer-learning baselines and other gradient-based meta-learning methods.

2. Background and related work

2.1. Hidden physics learning with neural networks

For many decades, physics-based PDEs have been commonly employed for predicting and monitoring complex system responses. Then traditional numerical methods were developed to solve these PDEs and provide predictions for desired system responses. However, three fundamental challenges usually present. First, the choice of governing PDE laws is often determined *a priori* and free parameters are often tuned to obtain agreement with experimental data, which makes the rigorous calibration and validation process challenging. Second, traditional numerical methods are solved for specific boundary and initial conditions, as well as loading or source terms. Therefore, they are not generalizable for other operating conditions and hence not effective for real-time prediction. Third, complex PDE systems such as turbulence flows and heterogeneous materials modeling problems usually require a very fine discretization, and are therefore very time-consuming for traditional solvers.

To provide an efficient surrogate model for physical responses, machine learning methods may hold the key. Recently, there has been significant progress in the development of deep neural networks (NNs) for learning the hidden physics of a complex system [17–25]. Among these methods, the neural operators show particular promises in resolving the above challenges, which aim to learn mappings between inputs of a dynamical system and its state, so that the network can serve as a surrogate for a solution operator [26–34].

Comparing with classical NNs, most notable advantages of neural operators are resolution independence and generalizability to different input instances. Moreover, comparing with the classical PDE modeling approaches, neural operators require only data with no knowledge of the underlying PDE. All these advantages make neural operators promising tools to PDE learning tasks. Examples include modeling the unknown physics law of real-world problems [35,36] and providing efficient solution operator for PDEs [26–28,37,38]. On the other hand, data in scientific applications are often scarce and incomplete. Utilization of other relevant data sources could alleviate such a problem, yet no existing work have addressed the transferability of neural operators. Through the

meta-learning techniques, our work fulfills the demand of such a transfer setting, with the same type of PDE system but different (hidden) physical properties.

2.2. Base model: Integral neural operators

We briefly introduce the integral neural operator model, which will be utilized as the base model of this work. The integral neural operators, first proposed in [26] and further developed in [27–29,39] comprises of three building blocks. First, the input function, $\mathbf{g}(\mathbf{x}) \in \mathcal{A}$, is lifted to a higher dimensional representation via $\mathbf{h}(\mathbf{x}, 0) = \mathcal{P}[\mathbf{g}](\mathbf{x}) := P(\mathbf{x})[\mathbf{x}, \mathbf{g}(\mathbf{x})]^T + \mathbf{p}(\mathbf{x})$. $P(\mathbf{x}) \in \mathbb{R}^{(s+d_g) \times d_h}$ and $\mathbf{p}(\mathbf{x}) \in \mathbb{R}^{d_h}$ define an affine pointwise mapping, which are often taken as constant parameters, i.e., $P(\mathbf{x}) \equiv P$ and $\mathbf{p}(\mathbf{x}) \equiv \mathbf{p}$. Then, the feature vector function $\mathbf{h}(\mathbf{x}, 0)$ goes through an iterative layer block where the layer update is defined via the action of the sum of a local linear operator, a nonlocal integral kernel operator, and a bias function: $\mathbf{h}(\cdot, l+1) = \mathcal{J}_{l+1}[\mathbf{h}(\cdot, l)]$. Here, $\mathbf{h}(\cdot, l)$, $l \in \{0, \dots, L\}$, is a sequence of functions representing values of the network at each hidden layer, taking values in $\mathbb{R}^{d_h}$. $\mathcal{J}_1, \dots, \mathcal{J}_L$ are nonlinear operator layers. In this work, we employ the implicit Fourier neural operator (IFNO) as the base model³ and take the iterative layers as $\mathcal{J}_1 = \dots = \mathcal{J}_L = \mathcal{J}$, where

$$\mathbf{h}(\mathbf{x}, l+1) = \mathcal{J}[\mathbf{h}(\mathbf{x}, l)] := \mathbf{h}(\mathbf{x}, l) + \frac{1}{L} \sigma(W\mathbf{h}(\mathbf{x}, l) + \mathcal{F}^{-1}[\mathcal{F}[\kappa(\cdot; \mathbf{v})] \cdot \mathcal{F}[\mathbf{h}(\cdot, l)]](\mathbf{x}) + \mathbf{c}(\mathbf{x})). \quad (2.1)$$

$\mathcal{F}$ and $\mathcal{F}^{-1}$ denote the Fourier transform and its inverse, respectively. $\mathbf{c} \in \mathbb{R}^{d_h}$ defines a constant bias, $W \in \mathbb{R}^{d_h \times d_h}$ is the weight matrix, and $\mathcal{F}[\kappa(\cdot; \mathbf{v})] := R$ is a circulant matrix that depends on the convolution kernel κ . σ is an activation function, which is often taken to be the popular rectified linear unit (ReLU) function. Finally, the output $\mathbf{u}(\cdot) \in \mathcal{U}$ is obtained through a projection layer, by mapping the last hidden layer representation $\mathbf{h}(\cdot, L)$ onto $\mathcal{U}$ as: $\mathbf{u}(\mathbf{x}) = \mathcal{Q}[\mathbf{h}(\cdot, L)](\mathbf{x}) := Q_2(\mathbf{x})\sigma(Q_1\mathbf{h}(\mathbf{x}, L) + \mathbf{q}_1(\mathbf{x})) + \mathbf{q}_2(\mathbf{x})$. $Q_1(\mathbf{x}) \in \mathbb{R}^{d_Q \times d_h}$, $Q_2(\mathbf{x}) \in \mathbb{R}^{d_u \times d_Q}$, $\mathbf{q}_1(\mathbf{x}) \in \mathbb{R}^{d_Q}$ and $\mathbf{q}_2(\mathbf{x}) \in \mathbb{R}^{d_u}$ are appropriately sized matrices and vectors that are part of the parameter set that we aim to learn, which are often taken as constant parameters and will be denoted as Q_1 , Q_2 , $\mathbf{q}_1$ and $\mathbf{q}_2$, respectively. In the following, we denote the set of trainable parameters in the lifting layer as θ_P , the set from the iterative layer block as θ_I , and the set in the projection layer as θ_Q .

The neural operator can be employed to learn an approximation for the solution operator, $\mathcal{G}$. Given $\mathcal{D} := \{(\mathbf{g}_i, \mathbf{u}_i)\}_{i=1}^N$, a labeled (context) set of observations, where the input $\{\mathbf{g}_i\} \subset \mathcal{A}$ is a set of independent and identically distributed (i.i.d.) random fields from a known probability distribution μ on $\mathcal{A}$, and $\mathbf{u}_i(\mathbf{x}) \in \mathcal{U}$ is the observed but possibly noisy corresponding solution. Let $\Omega \subset \mathbb{R}^s$ be the domain of interest, we assume that all observations can be modeled with a parametric PDE form:

$$\mathcal{K}_{b(\mathbf{x})}[\mathbf{u}_i](\mathbf{x}) = \mathbf{g}_i(\mathbf{x}), \quad \mathbf{x} \in \Omega. \quad (2.2)$$

Herein, $\mathcal{K}_b$ is the operator representing the possibly unknown governing law, e.g., balance laws. Then, the system response can be learnt by constructing a surrogate solution operator of Eq. (2.2): $\tilde{\mathcal{G}}[\mathbf{g}; \theta](\mathbf{x}) := Q_{\theta_Q} \circ (\mathcal{J}_{\theta_I})^L \circ P_{\theta_P}[\mathbf{g}](\mathbf{x}) \approx \mathbf{u}(\mathbf{x})$, where parameter set $\theta = [\theta_P, \theta_I, \theta_Q]$ is obtained by solving the optimization problem:

$$\min_{\theta \in \Theta} \mathcal{L}_{\mathcal{D}}(\theta) := \min_{\theta \in \Theta} \sum_{i=1}^N [C(\tilde{\mathcal{G}}[\mathbf{g}_i; \theta], \mathbf{u}_i)]. \quad (2.3)$$

Here C denotes a properly defined cost functional which is often taken as the relative mean square error.

2.3. Gradient-based meta-learning methods

Traditionally, machine learning models are designed to address one single task by learning from the dataset corresponding to such a task. Hence, its efficacy relies on the quality and abundance of available data. To improve the learning efficacy in the small data regime, a significant body of work has tackled the challenge of few-shot learning using transfer-learning or meta-learning approaches, with the basic idea of utilizing prior knowledge gained from related (source) tasks to improve the learning performance on the target task. Among these two types of approaches, the basic idea of transfer-learning is to train a model from the source task and then refine the pre-trained model in the target task [40]. In comparison, meta-learning approaches aim to learn a generalizable representations that shared across the different training tasks, which can accelerate learning of test tasks [5,41]. One of highly successful

³ We also point out that the proposed multi-task strategy is generic, and is hence also applicable to other neural operators [26–29,32].

meta-learning algorithms is Model Agnostic Meta-Learning (MAML) [5], which led to the development of a series of related gradient-based meta-learning (GBML) methods [7,9,10,42]. Almost-No-Inner-Loop algorithm (ANIL) [10] modifies MAML by freezing the final layer representation during local adaptation. Recently, theoretical analysis [12] found that the driving force causing MAML and ANIL to recover the general representation is the adaptation of the final layer of their models, which harnesses the underlying task diversity to improve the representation in all directions of interest.

Beyond applications such as image classification and reinforcement learning, a few approaches have studied hidden-physics learning under meta [43–46] or even transfer setting [47,48]. Among these meta-learning works, [43,44] are designed for specific physical applications, while [45,46] focus on dynamics forecasting by learning the temporal evolution information directly [45] or learning time-invariant features [46]. Hence, none of these works have provided a generic approach nor theoretical understanding on how to transfer the multi-task knowledge between a series of complex physical systems, such that all these tasks are governed by a common parametric PDE with different physical parameters.

3. Meta-learnt neural operator

To transfer the multi-task knowledge between a series of complex systems governed by different hidden physical parameters, we proposed to leverage the integral neural operator with a meta-learning setting. Before elaborating our novel meta-learnt neural operator architecture, MetaNO, we formally state the transfer-learning problem setting for PDE with different parameters.

Assume that we have a set of training tasks $\{\mathcal{T}^\eta\}$ such that $\mathcal{T}^\eta \sim p(\mathcal{T})$, and for each training task we have a set of observations of loading field/respond field data pairs $\mathcal{D}^\eta := \{(\mathbf{g}_i^\eta(\mathbf{x}), \mathbf{u}_i^\eta(\mathbf{x}))\}_{i=1}^{N^\eta}$. Each task can be modeled with a parametric PDE form

$$\mathcal{K}_{b^\eta(\mathbf{x})}[\mathbf{u}_i^\eta](\mathbf{x}) = \mathbf{g}_i^\eta(\mathbf{x}), \quad \mathbf{x} \in \Omega, \quad (3.1)$$

where $b^\eta(\mathbf{x})$ is the hidden task-specific physical parameter field for the common governing law. Given a new and unseen test task, $\mathcal{T}^{\text{test}}$, and a (usually small) context set of labeled samples $\mathcal{D}^{\text{test}} := \{(\mathbf{g}_i^{\text{test}}(\mathbf{x}), \mathbf{u}_i^{\text{test}}(\mathbf{x}))\}_{i=1}^{N^{\text{test}}}$ on it, our goal is to obtain the approximated solution operator model on the test task as $\tilde{\mathcal{G}}[\mathbf{g}; \theta^{\text{test}}]$. To provide a quantitative metric of the performance for each method, we reserve a separate set of labeled samples on the test task as the target set, and measure averaged relative errors of $\mathbf{u}$ on this set. In the few-shot learning context, we are particularly interested in the small-sample scenario where $N^{\text{test}} \ll N^\eta$.

3.1. A novel meta-learnt neural operator architecture

We now propose MetaNO, which applies *task-wise adaptation only to the first layer, i.e., the lifting layer*, with the full algorithm outlined in Algorithm 1. We point out that MetaNO is substantially different from existed popular meta-learning approaches such as MAML and ANIL, since the later rely on the adaptation of their last layer, as shown in [12]. This property makes MetaNO more suitable for PDE solving tasks as will be discussed in theoretical analysis below and confirmed in empirical evaluations of Section 4.

Similar as in other meta-learning approaches [49–52], the MetaNO algorithm consists of two phases: (1) a meta-train phase which learns shared iterative layers parameters θ_I and projection layer parameters θ_P from training tasks; (2) a meta-test phase which transfers the learned knowledge and rapidly learning surrogate solution operators for unseen test tasks with unknown physical parameter field, where only a few labeled samples are provided. In the meta-train phase, a batch $\{\mathcal{T}^\eta\}_{\eta=1}^H$ of H tasks is drawn from the training tasks set, with a context set of N^η numbers of labeled loading field/response field data pairs, $\mathcal{D}^\eta := \{(\mathbf{g}_i^\eta(\mathbf{x}), \mathbf{u}_i^\eta(\mathbf{x}))\}_{i=1}^{N^\eta}$, provided on each task. Then, we seek the common iterative (θ_I) and projection (θ_Q) parameters, and the task-wise lifting parameters θ_P^η by solving the optimization problem:

$$[\{\theta_P^{\eta,*}\}_{\eta=1}^H, \theta_I^*, \theta_Q^*] = \underset{\{\{\theta_P^\eta\}_{\eta=1}^H, \theta_I, \theta_Q\}}{\text{argmin}} \sum_{\eta=1}^H \mathcal{L}_{\mathcal{D}^\eta}([\theta_P^\eta, \theta_I, \theta_Q]). \quad (3.2)$$

Then, in the meta-test phase, we adapt the knowledge to a new and unseen test task $\mathcal{T}^{\text{test}}$, with limited data on the context set $\mathcal{D}^{\text{test}} := \{(\mathbf{g}_i^{\text{test}}(\mathbf{x}), \mathbf{u}_i^{\text{test}}(\mathbf{x}))\}_{i=1}^{N^{\text{test}}}$ on this task. In particular, we fix the common parameters θ_I^* and θ_Q^* ,

Algorithm 1 MetaNO**Meta-Train Phase:**

Input: a batch $\{\mathcal{T}^\eta\}_{\eta=1}^H$ of training tasks and labeled data pairs $\mathcal{D}^\eta := \{(\mathbf{g}_i^\eta(\mathbf{x}), \mathbf{u}_i^\eta(\mathbf{x}))\}_{i=1}^{N^\eta}$ on each task.

Output: common parameters θ_I^* and θ_Q^* across all tasks.

1. Initialize θ_I , θ_Q , and $\{\theta_P^\eta\}_{\eta=1}^H$.
2. Solve for $\{[\theta_P^{\eta,*}]_{\eta=1}^H, \theta_I^*, \theta_Q^*\}$ from the optimization problem in Eq. (3.2).

Meta-Test Phase:

Input: a test task $\mathcal{T}^{\text{test}}$ and few labeled data pairs $\mathcal{D}^{\text{test}} := \{(\mathbf{g}_i^{\text{test}}(\mathbf{x}), \mathbf{u}_i^{\text{test}}(\mathbf{x}))\}_{i=1}^{N^{\text{test}}}$ on it.

Output: the task-wise parameter $\theta_P^{\text{test},*}$ and the corresponding surrogate PDE solution operator $\tilde{\mathcal{G}}[\mathbf{g}; [\theta_P^{\text{test},*}, \theta_I^*, \theta_Q^*]](\mathbf{x})$ for the test task.

3. Solve for the lift layer parameter $\theta_P^{\text{test},*}$ from the optimization problem in Eq. (3.3).
4. (For cases with large N^{test} and/or small N^η), fine tune all parameters on the test task.

then solve for the task-wise parameter θ_P^{test} via:

$$\theta_P^{\text{test},*} = \underset{\theta_P^{\text{test}}}{\operatorname{argmin}} \mathcal{L}_{\mathcal{D}^{\text{test}}}([\theta_P^{\text{test}}, \theta_I^*, \theta_Q^*]). \quad (3.3)$$

One can then fine tune all test task parameters $[\theta_P^{\text{test}}, \theta_I, \theta_Q]$ for further improvements. Finally, the surrogate PDE solution operator on the test task is obtained as:

$$\tilde{\mathcal{G}}[\mathbf{g}; [\theta_P^{\text{test},*}, \theta_I^*, \theta_Q^*]](\mathbf{x}) := \mathcal{Q}_{\theta_Q^*} \circ (\mathcal{J}_{\theta_I^*})^L \circ \mathcal{P}_{\theta_P^{\text{test},*}}[\mathbf{g}](\mathbf{x}).$$

and will be evaluated on a reserved target data set on the test task.

3.2. Universal solution operator

To see the inspiration of the proposed architecture, without loss of generality, we assume that the underlying task parameter field $\mathbf{b}^\eta(\mathbf{x})$, modeling the physical property field, is normalized and satisfying $\|\mathbf{b}^\eta(\mathbf{x}) - \bar{\mathbf{b}}(\mathbf{x})\|_{L^2(\Omega)} \leq 1$ for all $\eta \in \{1, \dots, H\}$, where $\bar{\mathbf{b}} := \mathbb{E}_{\mathcal{T}^\eta \sim p(\mathcal{T})}[\mathbf{b}^\eta]$. Denoting $\mathcal{F}_u[\mathbf{b}] := \mathcal{K}_b[\mathbf{u}]$ as a function from physical parameter fields $\mathcal{B}$ to loading fields $\mathcal{A}$, we take the Fréchet derivative of $\mathcal{F}$ with respect to $\mathbf{b} - \bar{\mathbf{b}}$ and obtain:

$$\mathcal{K}_b[\mathbf{u}] = \mathcal{F}_u[\bar{\mathbf{b}}] + D\mathcal{F}_u[\bar{\mathbf{b}}](\mathbf{b}^\eta - \bar{\mathbf{b}}) + o(\|\mathbf{b}^\eta - \bar{\mathbf{b}}\|_{L^2(\Omega)}).$$

Substituting the above formulation into Eq. (3.1) yields:

$$\mathcal{F}_{u_i^\eta}[\bar{\mathbf{b}}] + D\mathcal{F}_{u_i^\eta}[\bar{\mathbf{b}}](\mathbf{b}^\eta - \bar{\mathbf{b}}) \approx \mathbf{g}_i^\eta.$$

Denoting $\mathbf{F}_1[\mathbf{b}^\eta] := [\mathbf{1}, \mathbf{b}^\eta - \bar{\mathbf{b}}]$ and $\mathbf{F}_2[\mathbf{u}_i^\eta] := [\mathcal{F}_{u_i^\eta}[\bar{\mathbf{b}}], D\mathcal{F}_{u_i^\eta}[\bar{\mathbf{b}}]]$, we can reformulate Eq. (3.1) into a more generic form:

$$\mathbf{F}_1[\mathbf{b}^\eta](\mathbf{x}) \cdot \mathbf{F}_2[\mathbf{u}_i^\eta](\mathbf{x}) = \mathbf{g}_i^\eta(\mathbf{x}), \quad \mathbf{x} \in \Omega. \quad (3.4)$$

Note that this parametric PDE form is very general and applicable to many science and engineering applications — besides our motivating example on material modeling, other examples include the monitoring of tissue degeneration problems [13], the detection of subsurface flows [53], the nondestructive inspection in aviation [54], and the prediction of concrete structures deterioration [55], etc.

In the following, we show that MetaNOs are universal solution operators for the multi-task PDE solving problem in Eq. (3.4), in the sense that they can approximate a fixed point method to a desired accuracy. For simplicity, we consider a 1D domain $\Omega \subset \mathbb{R}$, and scalar-valued functions $\mathbf{F}_1[\mathbf{b}^\eta]$, $\mathbf{F}_2[\mathbf{u}_i^\eta]$. These functions are assumed to be sufficiently smooth and measured at uniformly distributed nodes $\chi := \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_M\}$, with $\mathbf{F}_1[\mathbf{b}^\eta](\mathbf{x}_j) \neq 0$ for

all η and j . Then, Eq. (3.4) can be formulated as an implicit system of equations:

$$\mathcal{H}(U_i^{\eta,*}; \tilde{\mathbf{G}}_i^\eta) := \begin{bmatrix} \mathbf{F}_2[\mathbf{u}_i^\eta](\mathbf{x}_1) - \mathbf{g}_i^\eta(\mathbf{x}_1)/\mathbf{F}_1[\mathbf{b}^\eta](\mathbf{x}_1) \\ \vdots \\ \mathbf{F}_2[\mathbf{u}_i^\eta](\mathbf{x}_M) - \mathbf{g}_i^\eta(\mathbf{x}_M)/\mathbf{F}_1[\mathbf{b}^\eta](\mathbf{x}_M) \end{bmatrix} = \mathbf{0}, \quad (3.5)$$

where $U_i^{\eta,*} := [\mathbf{u}_i^\eta(\mathbf{x}_1), \dots, \mathbf{u}_i^\eta(\mathbf{x}_M)]$ is the solution we seek, $\tilde{\mathbf{G}}_i^\eta := [\mathbf{g}_i^\eta(\mathbf{x}_1)/\mathbf{F}_1[\mathbf{b}^\eta](\mathbf{x}_1), \dots, \mathbf{g}_i^\eta(\mathbf{x}_M)/\mathbf{F}_1[\mathbf{b}^\eta](\mathbf{x}_M)]$ is the reparameterized loading vector, and $\mathbf{G}_i^\eta := [\mathbf{g}_i^\eta(\mathbf{x}_1), \mathbf{g}_i^\eta(\mathbf{x}_2), \dots, \mathbf{g}_i^\eta(\mathbf{x}_M)]$ is the original loading vector. Here, we notice that all task-specific information is encoded in $\tilde{\mathbf{G}}_i^\eta$ and can be captured in the lifting layer parameter. Therefore, when seeing Eq. (3.5) as an implicit problem of $U_i^{\eta,*}$ and $\tilde{\mathbf{G}}_i^\eta$, it is actually independent of the task parameter field $\mathbf{b}^\eta$, i.e., this problem is task-independent. In the following, we refer to Eq. (3.5) without the task index, as $\mathcal{H}(U^*; \tilde{\mathbf{G}})$, for notation simplicity.

To solve for U^* from the nonlinear system $\mathcal{H}(U^*; \tilde{\mathbf{G}}) = \mathbf{0}$, a popular approach would be to use fixed-point iteration methods such as the Newton–Raphson method. With an initial guess of the solution (denoted as U^0), the process is repeated to produce successively better approximations to the roots of Eq. (3.5), from the solution of iteration l (denoted as U^l) to that of $l+1$ (denoted as U^{l+1}) as:

$$U^{l+1} = U^l - (\nabla \mathcal{H}(U^l; \tilde{\mathbf{G}}))^{-1} \mathcal{H}(U^l; \tilde{\mathbf{G}}) := U^l + \mathcal{R}(U^l, \tilde{\mathbf{G}}), \quad (3.6)$$

until a sufficiently precise value is reached. In the following, we show that as long as [Assumptions 1](#) and [2](#) hold, i.e., there exists a converging fixed point method, then MetaNO can be seen as an resemblance of the fixed point method in Eq. (3.6) and hence acts as an universal approximator of the solution operator for Eq. (3.4).

Assumption 1. There exists a fixed point equation, $U = U + \mathcal{R}(U, \tilde{\mathbf{G}})$ for the implicit problem in Eq. (3.5), such that $\mathcal{R} : \mathbb{R}^{2M} \mapsto \mathbb{R}^M$ is a continuous function satisfying $\mathcal{R}(U, \tilde{\mathbf{G}}) = \mathbf{0}$ and $\|\mathcal{R}(\hat{U}, \tilde{\mathbf{G}}) - \mathcal{R}(\tilde{U}, \tilde{\mathbf{G}})\|_{l^2(\mathbb{R}^M)} \leq m \|\hat{U} - \tilde{U}\|_{l^2(\mathbb{R}^M)}$ for any two vectors $\hat{U}, \tilde{U} \in \mathbb{R}^M$. Here, $m \geq 0$ is a constant independent of $\tilde{\mathbf{G}}$.

Assumption 2. With the initial guess $U^0 := [\mathbf{x}_1, \dots, \mathbf{x}_M]$, the fixed-point iteration $U^{l+1} = U^l + \mathcal{R}(U^l, \tilde{\mathbf{G}})$ ($l = 0, 1, \dots$) converges, i.e., for any given $\varepsilon > 0$, there exists an integer L such that

$$\|U^l - U^*\|_{l^2(\mathbb{R}^M)} \leq \varepsilon, \quad \forall l > L,$$

for all possible input instances $\tilde{\mathbf{G}} \in \mathbb{R}^M$ and their corresponding solutions U^* .

Intuitively, [Assumptions 1](#) and [2](#) ensure the hidden PDEs to be numerically solvable with a converging iterative solver, which is a typical required condition of numerical PDE solving problems. Then, we have our universal approximation theorem as below. The main result of this theorem is to show that for any desired accuracy $\varepsilon > 0$, one can find a sufficiently large $L > 0$ and sets of parameters $\theta^\eta = \{\theta_p^\eta, \theta_I, \theta_Q\}$, such that the resultant MetaNO model acts as a fixed point method with the desired prediction for all tasks and samples.

Theorem 1 (Universal Approximation). Given [Assumptions 1–2](#), let the activation function σ for all iterative kernel integration layers be the ReLU function, and the activation function in the projection layer be the identity function. Then for any $\varepsilon > 0$, there exist sufficiently large layer number $L > 0$ and feature dimension number $d_h > 0$, such that one can find a parameter set for the multi-task problem, $\theta^\eta = [\theta_p^\eta, \theta_I, \theta_Q]$, such that the corresponding MetaNO model satisfies

$$\left\| \mathcal{Q}_{\theta_Q} \circ (\mathcal{J}_{\theta_I})^L \circ \mathcal{P}_{\theta_p^\eta}([U^0, \mathbf{G}^\eta]^T) - U^{\eta,*} \right\| \leq \varepsilon,$$

for all loading instance $\mathbf{G}^\eta \in \mathbb{R}^M$ and tasks.

We now provide the detailed proof for [Theorem 1](#), based on [Assumptions 1](#) and [2](#). Intuitively, these assumptions mean the underlying implicit problem is solvable with a converging fixed point method. This condition is a basic requirement by numerical PDEs, and it generally holds true in many applications governed by nonlinear and complex PDEs, such as in our three experiments.

Here, we prove that the MetaNO is universal, i.e., given a fixed point method satisfying [Assumptions 1](#) and [2](#), one can find parameter sets θ^η whose output approximates $U^{\eta,*}$ to a desired accuracy, $\varepsilon > 0$, for all $\eta = 1, \dots, H$

tasks. For the task-wise parameters, with a slight abuse of notation, we denote $P^\eta \in \mathbb{R}^{d_h M \times (d_g + s)M}$ as the collection of the pointwise weight matrices at each discretization point in χ for the η th task, and $\mathbf{p}^\eta \in \mathbb{R}^{d_h M}$ for the bias in the lifting layer. Then, for the parameters shared among all tasks, in the iterative layer we denote $\mathbf{C} = [\mathbf{c}(\mathbf{x}_1), \dots, \mathbf{c}(\mathbf{x}_M)] \in \mathbb{R}^{d_h M}$ as the collection of pointwise bias vectors $\mathbf{c}(\mathbf{x}_i)$, $W \in \mathbb{R}^{d_h \times d_h}$ for the local linear transformation, and $R = \mathcal{F}[\kappa(\cdot; \mathbf{v})] \in \mathbb{C}^{d_h \times d_h \times M} \in \mathbb{C}^{d_h \times d_h \times M}$ for the Fourier coefficients of the kernel κ . For simplicity, here we have assumed that the Fourier coefficient is not truncated, and all available frequencies are used. Then, for the projection layer we seek $Q_1 \in \mathbb{R}^{d_Q M \times d_h M}$, $Q_2 \in \mathbb{R}^{d_u M \times d_Q M}$, $\mathbf{q}_1 \in \mathbb{R}^{d_Q M}$ and $\mathbf{q}_2 \in \mathbb{R}^{d_u M}$. For the simplicity of notation, in this section we organize the feature vector $\mathbf{H} \in \mathbb{R}^{d_h M}$ in a way such that the components corresponding to each discretization point are adjacent, i.e., $\mathbf{H} = [\mathbf{H}(\mathbf{x}_1), \dots, \mathbf{H}(\mathbf{x}_M)]$ and $\mathbf{H}(\mathbf{x}_i) \in \mathbb{R}^{d_h}$.

We point out that under this circumstance, the (discretized) iterative layer can be written as

$$\begin{aligned} \mathcal{J}[\mathbf{H}(l)] &= \mathbf{H}(l) + \frac{1}{L} \sigma \left(\tilde{W} \mathbf{H}(l) + \text{Re}(\mathcal{F}_{\Delta x}^{-1}(R \cdot \mathcal{F}_{\Delta x}(\mathbf{H}(l)))) + \mathbf{C} \right) \\ &= \mathbf{H}(l) + \frac{1}{L} \sigma (V \mathbf{H}(l) + \mathbf{C}), \end{aligned}$$

with

$$V := \text{Re} \begin{bmatrix} \sum_{n=0}^{M-1} R_{n+1} + W & \sum_{n=0}^{M-1} R_{n+1} \exp(\frac{2i\pi \Delta x n}{M}) & \dots & \sum_{n=0}^{M-1} R_{n+1} \exp(\frac{2i\pi(M-1)\Delta x n}{M}) \\ \sum_{n=0}^{M-1} R_{n+1} \exp(\frac{2i\pi \Delta x n}{M}) & \sum_{n=0}^{M-1} R_{n+1} + W & \dots & \sum_{n=0}^{M-1} R_{n+1} \exp(\frac{2i\pi(M-2)\Delta x n}{M}) \\ \vdots & \vdots & \ddots & \vdots \\ \sum_{n=0}^{M-1} R_{n+1} \exp(\frac{2i\pi(M-1)\Delta x n}{M}) & \sum_{n=0}^{M-1} R_{n+1} \exp(\frac{2i\pi(M-2)\Delta x n}{M}) & \dots & \sum_{n=0}^{M-1} R_{n+1} + W \end{bmatrix}.$$

Here, $R \in \mathbb{C}^{M \times d_h \times d_h}$ with $R_i \in \mathbb{C}^{d_h \times d_h}$ being the component associated with each discretization point $\mathbf{x}_i \in \chi$, $V \in \mathbb{R}^{d_h M \times d_h M}$, $\mathbf{C} \in \mathbb{R}^{d_h M}$, $\tilde{W} := W \oplus W \oplus \dots \oplus W$ is a $d_h M \times d_h M$ block diagonal matrix formed by $W \in \mathbb{R}^{d_h \times d_h}$, $\mathcal{F}_{\Delta x}$ and $\mathcal{F}_{\Delta x}^{-1}$ denote the discrete Fourier transform and its inverse, respectively. By further taking $R_2 = \dots = R_M = W = 0$, a $d_h \times d_h$ matrix with all its elements being zero, it suffices to show the universal approximation property for an iterative layer as follows:

$$\mathcal{J}(\mathbf{H}(l)) := \mathbf{H}(l) + \frac{1}{L} \sigma \left(\tilde{V} \mathbf{H}(l) + \mathbf{C} \right)$$

where $\tilde{V} := \mathbf{1}_{[M, M]} \otimes V$ with $V \in \mathbb{R}^{d_h \times d_h}$ and $\mathbf{1}_{[m, n]}$ being an m by n all-ones matrix.

For the proof of this main theorem, we need the following approximation property of a shallow neural network, with its detailed proof provided in [39]:

Lemma 1. Given a continuous function $\mathcal{T} : \mathbb{R}^{2M} \mapsto \mathbb{R}^M$, and a non-polynomial and continuous activation function σ , for any constant $\hat{\varepsilon} > 0$ there exists a shallow neural network model $\hat{\mathcal{T}} := S\sigma(B\mathbf{X} + A)$ such that

$$\|\mathcal{T}(\mathbf{X}) - \hat{\mathcal{T}}(\mathbf{X})\|_{l^2(\mathbb{R}^M)} \leq \hat{\varepsilon}, \quad \forall \mathbf{X} \in \mathbb{R}^{2M},$$

for sufficiently large feature dimension $\hat{d} > 0$. Here, $S \in \mathbb{R}^{M \times \hat{d}M}$, $B \in \mathbb{R}^{\hat{d}M \times 2M}$, and $A \in \mathbb{R}^{\hat{d}M}$ are matrices/vectors which are independent of $\mathbf{X}$.

We now proceed to the proof of [Theorem 1](#):

Proof. Since all $U^{\eta,*}$ satisfies [Assumptions 1–2](#), for any $\varepsilon > 0$, we first pick a sufficiently large integer L such that the L th layer iteration result of this fixed point formulation satisfies $\|U^L - U^{\eta,*}\|_{l^2(\mathbb{R}^M)} \leq \frac{\varepsilon}{2}$ for all tasks. By taking $\hat{\varepsilon} := \frac{m\varepsilon}{2(1+m)^L}$ in [Lemma 1](#), there exists a sufficiently large feature dimension $\hat{d}$ and one can find $S \in \mathbb{R}^{M \times \hat{d}M}$, $B \in \mathbb{R}^{\hat{d}M \times 2M}$, and $A \in \mathbb{R}^{\hat{d}M}$, such that $\hat{\mathcal{R}}(U^\eta, \tilde{\mathbf{G}}^\eta) := S\sigma(B[U^\eta, \tilde{\mathbf{G}}^\eta]^T + A)$ satisfies

$$\|\mathcal{R}(U^\eta, \tilde{\mathbf{G}}^\eta) - \hat{\mathcal{R}}(U^\eta, \tilde{\mathbf{G}}^\eta)\|_{l^2(\mathbb{R}^M)} = \|\mathcal{R}(U^\eta, \tilde{\mathbf{G}}^\eta) - S\sigma(B[U^\eta, \tilde{\mathbf{G}}^\eta]^T + A)\|_{l^2(\mathbb{R}^M)} \leq \hat{\varepsilon} = \frac{m\varepsilon}{2(1+m)^L},$$

where m is the contraction parameter of $\mathcal{R}$, as defined in [Assumption 1](#). By this construction, we know that S has independent rows. Denoting $\tilde{d} := \hat{d} + 1 > 0$, there exists the right inverse of S , which we denote as $S^+ \in \mathbb{R}^{(\tilde{d}-1)M \times M}$, such that

$$SS^+ = I_M, \quad S^+S := \tilde{I}_{(\tilde{d}-1)M},$$

where I_M is the M by M identity matrix, $\tilde{I}_{(\tilde{d}-1)M}$ is a $(\tilde{d}-1)M$ by $(\tilde{d}-1)M$ block matrix with each of its element being either 1 or 0. Hence, for any vector $Z \in \mathbb{R}^{(\tilde{d}-1)M}$, we have $\sigma(\tilde{I}_{(\tilde{d}-1)M}Z) = \tilde{I}_{(\tilde{d}-1)M}\sigma(Z)$. Moreover, we note that S has a very special structure: from the $((i-1)(\tilde{d}-1)+1)$ th to the $(i(\tilde{d}-1))$ th column of S , all nonzero elements are on its i th row. Correspondingly, we can also choose S^+ to have a special structure: from the $((i-1)(\tilde{d}-1)+1)$ th to the $(i(\tilde{d}-1))$ th row of S^+ , all nonzero elements are on its i th column. Hence, when multiplying S^+ with U , there will be no entanglement between different components of U . That means, S^+ can be seen as a pointwise weight function.

We now construct the parameters of MetaNO as follows. In this construction, we choose the feature dimension as $d_h := \tilde{d}M$. With the input $[U^0, G^\eta] \in \mathbb{R}^{2M}$, for the lift layer we set

$$P^\eta := \mathbf{1}_{[M,1]} \otimes \begin{bmatrix} S^+ & \mathbf{0} \\ \mathbf{0} & D^\eta \end{bmatrix} = \underbrace{\begin{bmatrix} S^+ & \mathbf{0} & S^+ & \mathbf{0} & \cdots & S^+ & \mathbf{0} \\ \mathbf{0} & D^\eta & \mathbf{0} & D^\eta & \cdots & \mathbf{0} & D^\eta \end{bmatrix}}_{\text{repeated for } M \text{ times}} \in \mathbb{R}^{d_h M \times 2M},$$

and $p^\eta := \mathbf{0} \in \mathbb{R}^{d_h M}$. Here, $D^\eta := \text{diag}[1/F_1[\mathbf{b}^\eta](\mathbf{x}_1), \dots, 1/F_1[\mathbf{b}^\eta](\mathbf{x}_M)]$. As such, the initial layer of feature is then given by

$$\mathbf{H}(0) = P^\eta([U^0, G^\eta]^T) = \mathbf{1}_{[M,1]} \otimes [S^+U^0, D^\eta G^\eta]^T = \mathbf{1}_{[M,1]} \otimes [S^+U^0, \tilde{G}^\eta]^T \in \mathbb{R}^{d_h M}.$$

Here, we point out that P^η and p^η can be seen as pointwise weight and bias functions, respectively.

Next we construct the shared iterative layer $\mathcal{J}$, by setting

$$V := \begin{bmatrix} \tilde{I}_{(\tilde{d}-1)M} B/M \\ \mathbf{0} \end{bmatrix} \begin{bmatrix} LS & \mathbf{0} \\ \mathbf{0} & LI_M \end{bmatrix}, \quad \tilde{V} := \mathbf{1}_{[M,M]} \otimes V, \quad \text{and } C := \mathbf{1}_{[M,1]} \otimes \begin{bmatrix} L\tilde{I}_{(\tilde{d}-1)M} A \\ \mathbf{0} \end{bmatrix}.$$

Note that $\tilde{V}$ is independent of η , and falls into the formulation of V , by letting $R_1 = V$ and $R_2 = R_3 = \dots = R_M = W = 0$. For the $l+1$ th layer of feature vector, we then arrive at

$$\begin{aligned} \mathbf{H}(l+1) &= \mathbf{H}(l) + \frac{1}{L} \sigma \left(\tilde{V} \mathbf{H}(l) + C \right) \\ &= \mathbf{H}(l) + \left(I_M \otimes \begin{bmatrix} S^+S & \mathbf{0} \\ \mathbf{0} & I_M \end{bmatrix} \right) \sigma \left(\left(\mathbf{1}_{[M,1]} \otimes \begin{bmatrix} B/M \\ \mathbf{0} \end{bmatrix} \right) \left(\mathbf{1}_{[1,M]} \otimes \begin{bmatrix} S & \mathbf{0} \\ \mathbf{0} & I_M \end{bmatrix} \right) \mathbf{H}(l) + \mathbf{1}_{[M,1]} \otimes \begin{bmatrix} A \\ \mathbf{0} \end{bmatrix} \right), \end{aligned}$$

where $\mathbf{H}(l) = [\hat{\mathbf{h}}_1^l, \hat{\mathbf{h}}_2^l, \dots, \hat{\mathbf{h}}_{2M-1}^l, \hat{\mathbf{h}}_{2M}^l]^T$ denotes the (spatially discretized) hidden layer feature at the l th iterative layer of the MetaNO. Subsequently, we note that the second part of the feature vector, $\hat{\mathbf{h}}_{2j}^l \in \mathbb{R}^M$, satisfies

$$\hat{\mathbf{h}}_{2j}^{l+1} = \hat{\mathbf{h}}_{2j}^l = \dots = \hat{\mathbf{h}}_{2j}^0 = \tilde{G}^\eta, \quad \forall l = 0, \dots, L-1, \forall j = 1, \dots, M$$

Hence, the first part of the feature vector, $\hat{\mathbf{h}}_{2j-1}^l \in \mathbb{R}^{(\tilde{d}-1)M}$, satisfies the following iterative rule:

$$\hat{\mathbf{h}}_{2j-1}^{l+1} = \hat{\mathbf{h}}_{2j-1}^l + S^+S\sigma(B[S\hat{\mathbf{h}}_{2j-1}^l, \tilde{G}^\eta]^T + A), \quad \forall l = 0, \dots, L-1, \forall j = 1, \dots, M,$$

and

$$\hat{\mathbf{h}}_1^{l+1} = \hat{\mathbf{h}}_3^{l+1} = \dots = \hat{\mathbf{h}}_{2M-1}^{l+1}.$$

Finally, for the projection layer $\mathcal{Q}$, we set the activation function in the projection layer as the identity function, $\mathcal{Q}_1 := I_{d_h M}$ (the identity matrix of size $d_h M$), $\mathcal{Q}_2 := [S, \mathbf{0}] \in \mathbb{R}^{M \times d_h M}$, $\mathbf{q}_1 := \mathbf{0} \in \mathbb{R}^{d_h M}$, and $\mathbf{q}_2 := \mathbf{0} \in \mathbb{R}^M$. Denoting the output $U^\eta := \mathcal{Q}_{\theta_Q} \circ (\mathcal{J}_{\theta_l})^L \circ \mathcal{P}_{\theta_p}([U^0, G^\eta]^T)$, we now show that U^η can approximate $U^{\eta,*}$ with a

desired accuracy ε :

$$\begin{aligned}
\|\mathbf{U}^\eta - \mathbf{U}^{\eta,*}\| &\leq \|\mathbf{U}^\eta - \mathbf{U}^L\|_{l^2(\mathbb{R}^M)} + \|\mathbf{U}^L - \mathbf{U}^{\eta,*}\|_{l^2(\mathbb{R}^M)} \\
&\leq \|\hat{\mathbf{S}}\mathbf{h}_1^L - \mathbf{U}^L\|_{l^2(\mathbb{R}^M)} + \frac{\varepsilon}{2} \quad (\text{by Assumption 2}) \\
&\leq \|\hat{\mathbf{S}}\mathbf{h}_1^{L-1} - \mathbf{U}^{L-1}\|_{l^2(\mathbb{R}^M)} + \|\hat{\mathcal{R}}(\hat{\mathbf{S}}\mathbf{h}_1^{L-1}, \tilde{\mathbf{G}}) - \mathcal{R}(\mathbf{U}^{L-1}, \tilde{\mathbf{G}})\|_{l^2(\mathbb{R}^M)} + \frac{\varepsilon}{2} \\
&\leq \|\hat{\mathbf{S}}\mathbf{h}_1^{L-1} - \mathbf{U}^{L-1}\|_{l^2(\mathbb{R}^M)} + \|\hat{\mathcal{R}}(\hat{\mathbf{S}}\mathbf{h}_1^{L-1}, \tilde{\mathbf{G}}b) - \mathcal{R}(\hat{\mathbf{S}}\mathbf{h}_1^{L-1}, \tilde{\mathbf{G}}b)\|_{l^2(\mathbb{R}^M)} \\
&\quad + \|\mathcal{R}(\hat{\mathbf{S}}\mathbf{h}_1^{L-1}, \tilde{\mathbf{G}}b) - \mathcal{R}(\mathbf{U}^{L-1}, \tilde{\mathbf{G}}b)\|_{l^2(\mathbb{R}^M)} + \frac{\varepsilon}{2} \\
&\leq (1+m)\|\hat{\mathbf{S}}\mathbf{h}_1^{L-1} - \mathbf{U}^{L-1}\|_{l^2(\mathbb{R}^M)} + \frac{m\varepsilon}{2(1+m)^L} + \frac{\varepsilon}{2} \quad (\text{by Lemma 1 and Assumption 1}) \\
&\leq \frac{m\varepsilon}{2(1+m)^L} (1 + (1+m) + (1+m)^2 + \dots + (1+m)^{L-1}) + \frac{\varepsilon}{2} \\
&\leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon. \quad \square
\end{aligned}$$

4. Empirical evaluation

In this section, we demonstrate the empirical effectiveness of the proposed MetaNO approach. Specifically, we conduct experiments on a synthetic dataset from a nonlinear PDE solving problem, a benchmark dataset of heterogeneous materials subject to large deformation, and a real-world dataset from biological tissue mechanical testing. We compare the proposed method against competitive GBML methods as well as two non-meta transfer-learning baselines. All of the experiments are implemented using PyTorch with Adam optimizer, and performed on a workstation with two 2.2 GHz 24-core CPUs and eight NVIDIA Tesla T4 GPUs. For a fair comparison, for each algorithm, we tune the hyperparameters, including the learning rate from $\{0.1, 0.01, 0.001, 0.0001, 0.00001, 0.000001\}$, the decay rate from $\{0.5, 0.7, 0.9\}$, the weight decay parameter from $\{0.01, 0.001, 0.0001, 0.00001, 0.000001\}$, and the inner loop learning rate for the baseline GBML methods (MAML and ANIL) from $\{0.01, 0.001, 0.0001, 0.00001, 0.000001\}$, to minimize the error on a separate validation dataset. In all experiments we decrease the learning rate with a ratio of learning rate decay rate every 100 epochs. A detailed description of each baseline method is provided in [Appendix A](#).

In all experiments, we considered the averaged relative error, $\|\mathbf{u}_{i,pred} - \mathbf{u}_i\|_{L^2(\Omega)} / \|\mathbf{u}_i\|_{L^2(\Omega)}$, as the error metric. We repeat each experiment for 5 times, and report the averaged relative errors and their standard errors.

4.1. Example 1: Synthetic data sets and ablation study

We first consider the PDE-solution-finding problem of the Holzapfel–Gasser–Ogden (HGO) model [56], which describes the deformation of hyperelastic, anisotropic, and fiber-reinforced materials. In this example, the goal is to find its displacement field $\mathbf{u} : [0, 1]^2 \rightarrow \mathbb{R}^2$ under different boundary loadings. The specimen is assumed to be subject to a uniaxial tension $T_y(\mathbf{x})$ on the top edge (see [Fig. 2\(a\)](#)). Therefore, we take the input function $\mathbf{g}(\mathbf{x})$ as the padded traction loading field, and the output function as the corresponding displacement field.

Data Generation. To generate training and test samples, the Holzapfel–Gasser–Ogden (HGO) model [56] was employed to describe the constitutive behavior of the material in this example, with its strain energy density function given as:

$$\begin{aligned}
\eta &= \frac{E}{4(1+\nu)}(\bar{I}_1 - 2) - \frac{E}{2(1+\nu)} \ln(J) \\
&\quad + \frac{k_1}{2k_2} (\exp(k_2 \langle S(\alpha) \rangle^2) + \exp(k_2 \langle S(-\alpha) \rangle^2) - 2) + \frac{E}{6(1-2\nu)} \left(\frac{J^2 - 1}{2} - \ln J \right).
\end{aligned}$$

Here, $\langle \cdot \rangle$ denotes the Macaulay bracket, and the fiber strain of the two fiber groups is defined as:

$$S(\alpha) = \frac{\bar{I}_4(\alpha) - 1 + |\bar{I}_4(\alpha) - 1|}{2}.$$

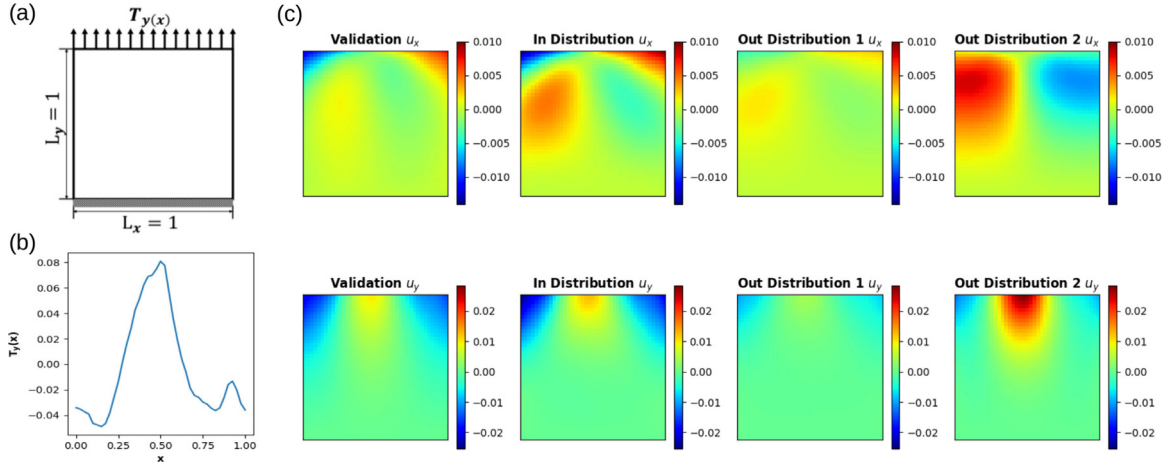


Fig. 2. Problem setup of example 1: the synthetic data sets. (a) A unit square specimen subject to uniaxial tension with Neumann-type boundary condition. (b) & (c) Visualization of an instances of the loading field $T_y(x)$, and the corresponding ground-truth solutions $\mathbf{u}^\eta(\mathbf{x})$ from the in-distribution and out-of-distribution tasks, showing the solution diversity across different tasks, due to the change of underlying hidden material parameter set.

where k_1 and k_2 are fiber modulus and the exponential coefficient, respectively, E is the Young's modulus for the non-fibrous ground matrix, and ν is the Poisson ratio. Moreover, $\bar{I}_1 = \text{tr}(\mathbf{C})$ is the first invariant of the right Cauchy–Green deformation tensor $\mathbf{C} = \mathbf{F}^T \mathbf{F}$, $\mathbf{F}$ is the deformation gradient, and J is related with $\mathbf{F}$ such that $J = \det \mathbf{F}$. For the fiber group with angle direction α from the reference direction, $\bar{I}_4(\alpha) = \mathbf{n}^T(\alpha) \mathbf{C} \mathbf{n}(\alpha)$ is the fourth invariant of the right Cauchy–Green tensor $\mathbf{C}$, where $\mathbf{n}(\alpha) = [\cos(\alpha), \sin(\alpha)]^T$. To generate samples for different specimens, different specimens (tasks) correspond to different material parameter sets, $\{k_1, k_2, E, \nu, \alpha\}$.

To investigate the performance of MetaNO in few-shot learning, we generate 59 training and one validation tasks, as well as five in-distribution (ID) test tasks by sampling different physical parameters k_1, k_2, E, ν, α from the same uniform distribution. To further evaluate the generalizability when the physical parameters of test tasks are outside the training regime, we also generate two out-of-distribution (OOD) test tasks with physical parameters from different distributions. For the training tasks, the validation task, and the in-distribution (ID) test task, their physical parameters are sampled from: $k_1, k_2 \sim \mathcal{U}[0.1, 1]$, $E \sim \mathcal{U}[0.55, 1.5]$, $\nu \sim \mathcal{U}[0.01, 0.49]$, and $\alpha \sim \mathcal{U}[\pi/10, \pi/2]$. For the two out-of-distribution (OOD) test tasks, we sample their parameters following $k_1, k_2 \sim \mathcal{U}[1, 1.9]$, $E \sim \mathcal{U}[1.5, 2] \cup \mathcal{U}[0.5, 0.55]$, $\nu \sim \mathcal{U}[0.01, 0.49]$,⁴ and $\alpha \sim \mathcal{U}[\pi/2, 3\pi/4] \cup [0, \pi/10]$. To generate the high-fidelity (ground-truth) dataset, we sampled 500 different vertical traction conditions $T_y(\mathbf{x})$ on the top edge from a random field, following the algorithm in [36,58]. $T_y(\mathbf{x})$ is taken as the restriction of a 2D random field, $\phi(\mathbf{x}) = \mathcal{F}^{-1}(\gamma^{1/2} \mathcal{F}(I))(\mathbf{x})$, on the top edge. Here, $I(\mathbf{x})$ is a Gaussian white noise random field on $\mathbb{R}^2$, $\gamma = (w_1^2 + w_2^2)^{-\frac{5}{4}}$ represents a correlation function, and w_1, w_2 are the wave numbers on x and y directions, respectively. Then, for each sampled traction loading, we solved the displacement field on the entire domain by minimizing potential energy using the finite element method implemented in FEniCS [59]. In particular, the displacement field was approximated by continuous piecewise linear finite elements with triangular mesh, and the grid size was taken as 0.025. The finite element solution was interpolated onto χ , a structured 41×41 grid which will be employed as the discretization in our neural operators. The distribution of training and ID/OOD tasks are demonstrated in Fig. 3, where one can see that the first OOD task (denoted as “OOD Task1”) corresponds to a stiffer material sample and smaller deformation for each given loading, while the second OOD task (denoted as “OOD Task2”) generates a softer material sample and larger deformation. For each training task, we generate 500 data pairs $\mathcal{D}^\eta := \{(\mathbf{g}_i^\eta, \mathbf{u}_i^\eta)\}_{i=1}^{500}$, by sampling the vertical traction loading from a Gaussian random field. Then, the corresponding ground-truth displacement field is obtained using the finite element method implemented in FEniCS [59]. For test tasks, we train with $N^{\text{test}} = \{2, 4, 8, 12, 20, 100, 300\}$ numbers of labeled data pairs (the

⁴ Here we sample both ID and OOD tasks from the same range of ν , due to the fact that $[0.01, 0.49]$ is the range of Poisson ratio for common materials [57].

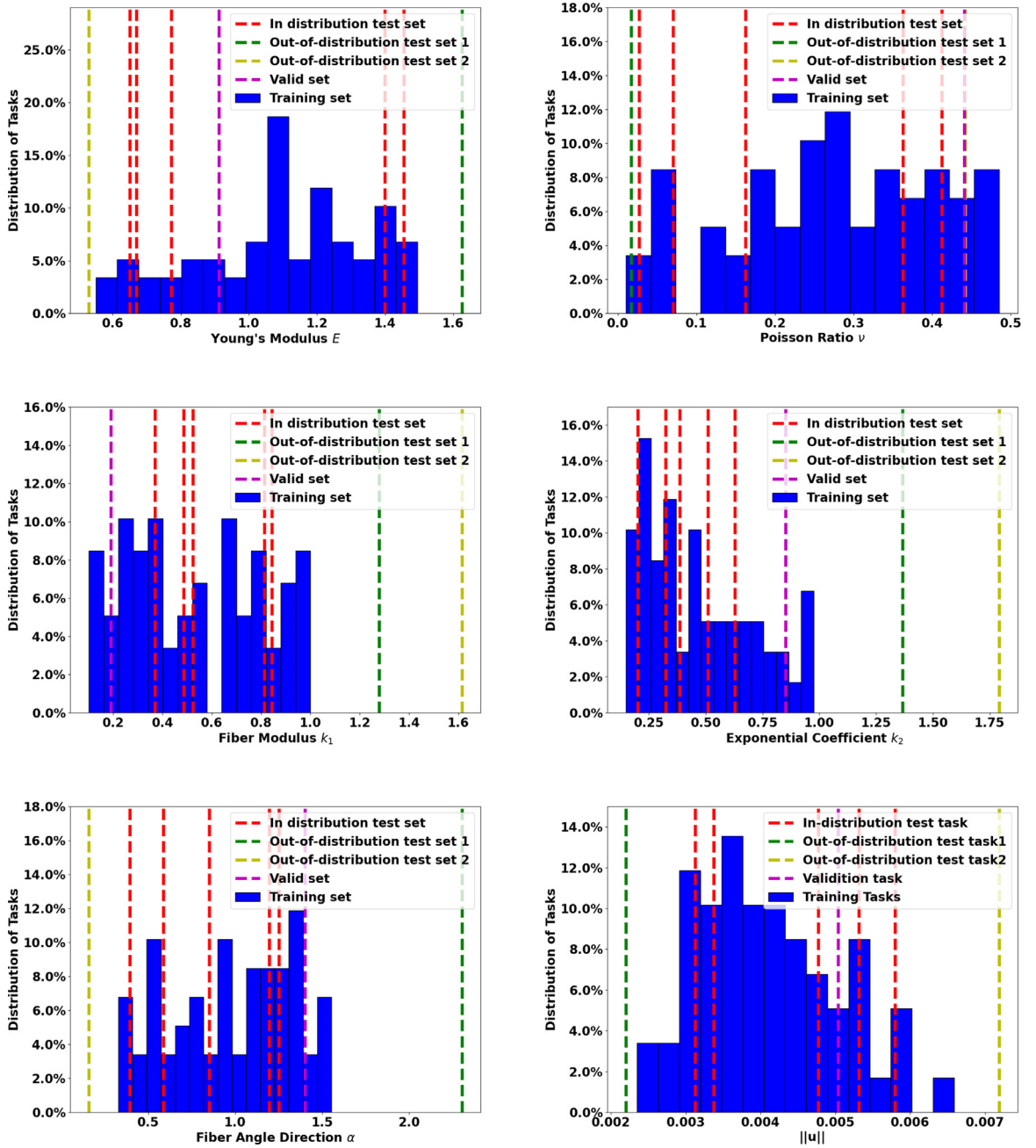


Fig. 3. Distribution of physical parameters of different tasks in example 1, and the resultant magnitude of material response, $\|u^\eta(x)\|_{L^2(\Omega)}$, on an exemplar loading instance shown in Fig. 2(b).

context set), and evaluate the model on a reserved dataset with 200 data pairs (the target set) on each test task. An 8-layer IFNO is employed as the base model.

Ablation Study. We first conduct an ablation study on 4 variants of the proposed algorithm: (1) to use the full meta-train and meta-test phases as in Algorithm 1 (denotes as “MetaNO”); (2) to perform steps 1–3 of Algorithm 1, such that only the lifting layer is adapted in the meta-test phase (denotes as “MetaNO–”); (3) to apply task-wise adaptation only to the iterative layers instead of the lift layer in both meta-train and meta-test phases (denotes as “MetaMid”); (4) to apply task-wise adaptation only to the projection layer instead of the lift layer in both meta-train

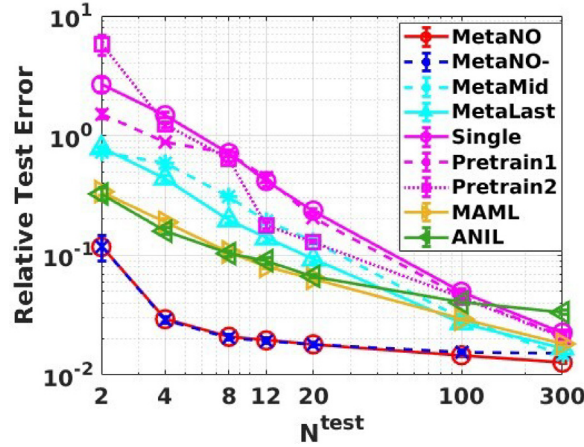


Fig. 4. Results of the ablation study on example 1: the synthetic data set. Comparison on test errors in the in-distribution test, when using the full context set ($N^\eta = 500$) on training tasks and different sizes of context set (N^{test}) on test tasks.

Table 1

A comparison of the computational costs (in seconds) and test errors (in relative L^2 norm) of MetaNO and MetaNO–, in the ablation study on example 1. Here, the total time of MetaNO– includes the Meta-Train Phase (Steps 1 and 2) and the step of solving for θ_p^{test} in the Meta-Test Phase (Step 3), and the total time of MetaNO includes all steps of both the Meta-Train and Meta-Test Phases (Steps 1–4).

N^{test}			MetaNO–		MetaNO	
	Step 3	Step 4	Total time	Test error	Total time	Test error
2	24	25	45636	$11.70\% \pm 14.01\%$	45661	$11.72\% \pm 14.02\%$
4	48	49	45660	$2.91\% \pm 0.90\%$	45709	$2.93\% \pm 0.90\%$
8	93	98	45705	$2.07\% \pm 0.53\%$	45803	$2.08\% \pm 0.53\%$
12	141	147	45753	$1.94\% \pm 0.56\%$	45900	$1.92\% \pm 0.55\%$
20	223	247	45835	$1.82\% \pm 0.41\%$	46082	$1.80\% \pm 0.39\%$
100	976	1025	46588	$1.56\% \pm 0.35\%$	47613	$1.46\% \pm 0.26\%$
300	2116	2282	47728	$1.53\% \pm 0.34\%$	50010	$1.27\% \pm 0.15\%$

and meta-test phases (denoted as “MetaLast”). We study if the successful “adapting last layers” strategy of MAML and ANIL in image classification problems would apply for our PDE solving problem. Besides these four settings, we also report the few-shot learning results with five baseline methods: (1) Learn a neural operator model only based on the context data set of the test task (denoted as “Single”); (2) Pretrain a neural operator model based on all training task data sets, then fine-tune it based on the context test task data set (denoted as “Pretrain1”); (3) Pretrain a single neural operator model based on the context data set of one training task, then fine-tune it based on the context test task data set (denoted as “Pretrain2”); To remove the possible dependency on the pre-training task, in this baseline we randomly select five training tasks for the purpose of pretraining and report the averaged results. (4) MAML, and (5) ANIL. For all experiments we use the full context data set on each training task ($N^\eta = 500$).

As shown in Fig. 4, MetaNO– and MetaNO are both able to quickly adapt with few data pairs – to achieve a test error below 5%, “Single” and the two transfer-learning baselines (“Pretrain1”, “Pretrain2”) require 100+ data pairs, while MetaNO– and MetaNO requires only 4 data pairs. On the other hand, MetaMid, MetaLast, MAML and ANIL have similar performance. They all require 100 data pairs to achieve a $< 5\%$ test error. In particular, when using only 4 data pairs on the target task, “MetaNO–” has $2.91\% \pm 0.90\%$ relative test error, while “MetaMid” and “MetaLast” have $58.41\% \pm 9.92\%$ and $43.99\% \pm 11.36\%$ relative test errors, respectively. This observation verifies our finding on the multi-task parametric PDE solution operator learning problem, where one should adapt the first layer, not the middle layers or the last ones. A detailed comparison on the computational time and test errors of these “MetaNO–” and “MetaNO” is provided in Table 1. Herein, the total time of MetaNO– includes the Meta-Train Phase (Steps 1 and 2) and the step of solving for θ_p^{test} in the Meta-Test Phase (Step 3) of Algorithm 1, and the total time of MetaNO includes all steps of both the Meta-Train and Meta-Test Phases (Steps 1–4). Since the Meta-Train

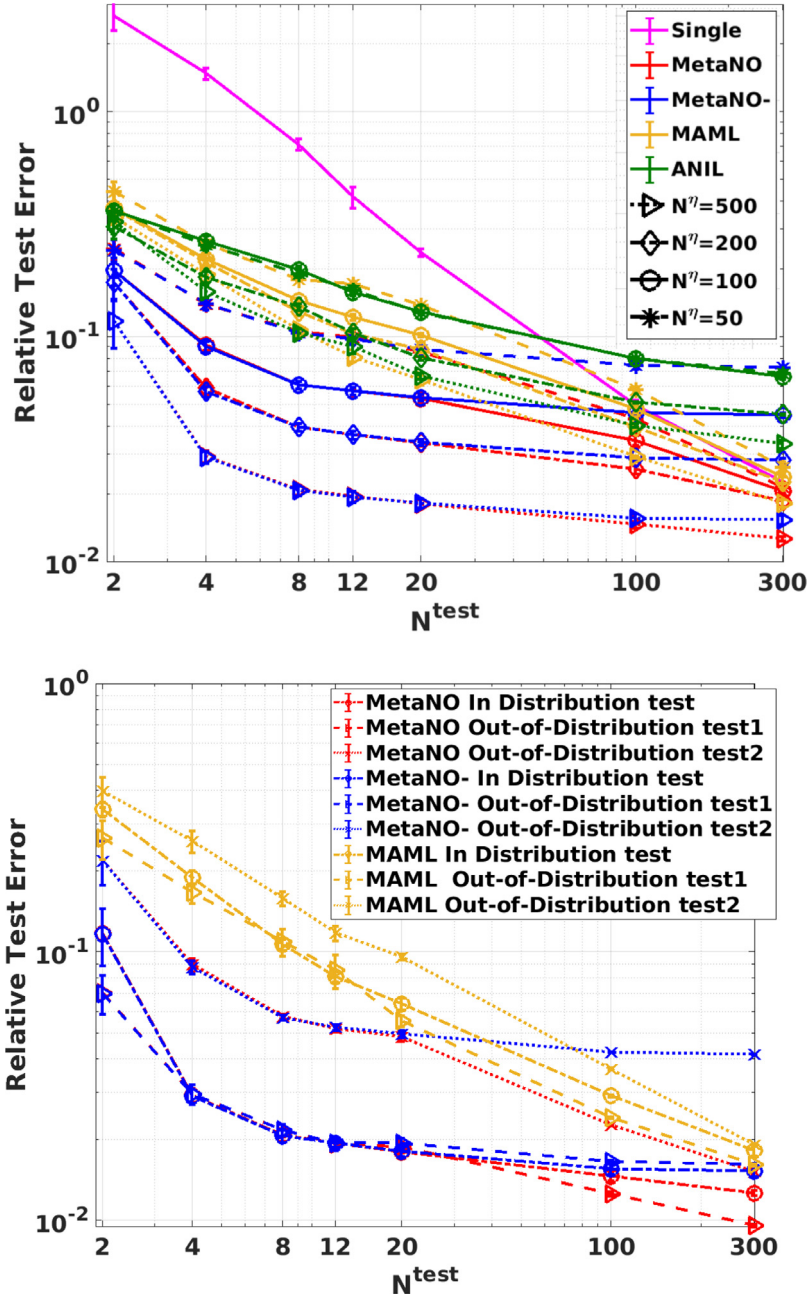


Fig. 5. Results on example 1: the synthetic data set. Top: The relative test error results showing the effect of varying training task context set sizes $N^\eta \in \{50, 100, 200, 500\}$. Bottom: The relative error of MetaNO and MAML in in-distribution and out-of distribution tests.

Phase took the majority of training time (45612 seconds), although “MetaNO” requires an additional fine-tuning step, its computational cost is of a similar scale as “MetaNO-”. When comparing the test errors, we can see that the additional fine-tune step only improves the performance in the larger-sample regime (when $N^{\text{test}} \geq 100$). This fact shows that when given sufficient training context sets, adapting the first layer can capture the underlying task diversity so further fine-tuning may not be needed, especially in the small data regime.

Effect of Varying Training Context Set Sizes. In this study, we investigate the effect of different training task context sizes $N^\eta = \{50, 100, 200, 500\}$ on four meta-learned models: MetaNO, MetaNO-, MAML, and ANIL. The comparison on test errors are demonstrated in the top plot of Fig. 5. Here, MetaNO- and MetaNO did not have

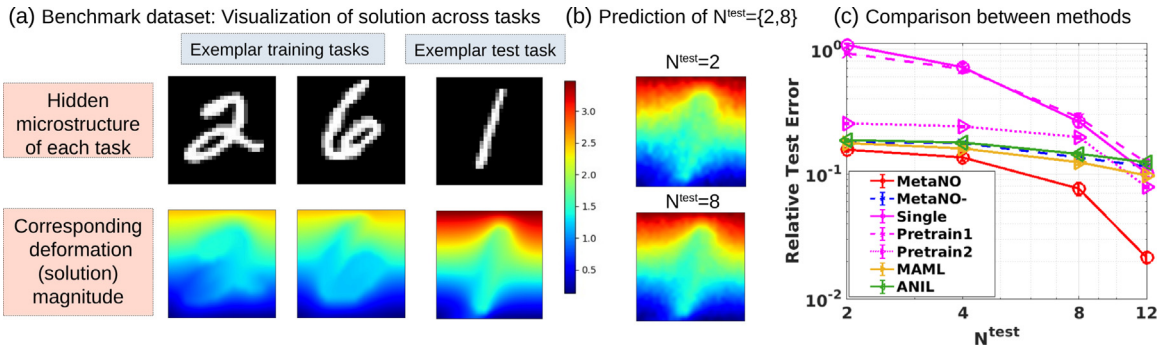


Fig. 6. Results on example 2: the benchmark (Mechanical MNIST [60]) dataset. (a) The visualization of different tasks, their underlying microstructure field $\mathbf{b}^\eta$, and the corresponding ground-truth solution. (b) Prediction results based on few samples ($N^{\text{test}} = 2$ and $N^{\text{test}} = 8$) on a test task. (c) Comparison of MetaNO and five baseline methods.

any inner loop updates. All parameters from all training tasks are optimized together. In MAML and ANIL we use half of the context set for inner loop updates (support set) and the other half for outer loop updates (target set). One can see that when $N^{\text{test}} \leq 20$, MetaNO– and MetaNO have similar performance and consistently beat MAML and ANIL for both context set sizes. With the increase of N^{test} , the fine-tuning strategy on the test context set becomes more helpful where we see MetaNO becomes more accurate than MetaNO– and MAML beats ANIL. Such effect is more evident on small training context set cases. Perhaps unsurprisingly, with the training task context size varying from 50 to 500, all methods have improved performance with decreasing relative test errors (with the same colors for the same methods across different context dataset). In addition, as the context set size in the test task grows, fine-tuning will gradually have better performance as MetaNO and MAML beats MetaNO– and ANIL, respectively. Overall, in all combinations of N^η and N^{test} , MetaNO achieves the best performance among all models.

In-Distribution and Out-Of-Distribution Tests. On bottom Fig. 5, we demonstrate the relative test error of MetaNO against MAML in both ID and OOD tasks. We can see that test errors of these three tasks are at a similar scale as the error on training tasks. The error from OOD task1 is comparable to the averaged ID test task error, while the error from OOD task2 is much larger, as expected by the data property: due to the fact that the solutions in OOD task1 generally have smaller magnitude and hence its solution operator lies more in a linear regime, which makes the solution operator learning task easier. Probably due to the fact that the solutions in OOD task1 generally have smaller magnitude and hence its solution operator lies more in a linear regime, which makes the solution operator learning task easier. In all three cases, MetaNO outperforms MAML, hence validating the good generalization performance of MetaNO.

4.2. Example 2: Benchmark mechanical MNIST datasets

We further test MetaNO and five baseline methods on benchmark Mechanical MNIST [60]. Mechanical MNIST is a benchmark dataset of heterogeneous material undergoing large deformation, modeled by the Neo-Hookean material with a varying modulus converted from the MNIST bitmap images [60]. It contains 70,000 heterogeneous material specimens, and each specimen is governed by the Neo-Hookean material with a varying modulus converted from the MNIST bitmap images. On each specimen, we have 32 loading/response data pairs on a 27×27 structured grid, under the uniaxial extension, shear, equibiaxial extension, and confined compression load scenarios, respectively.⁵

Data Generation and Settings. Here in, we randomly select one specimen corresponding to hand-written number 0 and 2 – 9 respectively as training tasks. Then, among the specimens corresponding to 1, we randomly select six specimens: one for validation and the rest five as the test tasks. Visualization of the ground-truth solutions corresponding to one common loading from different tasks is provided in Fig. 6(a), together with the underlying (hidden) microstructure pattern which determines the parameter set $\mathbf{b}^\eta$. On the meta-train phase, we use the full context data set of all 32 samples for each training task. On the meta-test phase, we reserve 20 data pairs on

⁵ We have excluded small deformation samples with the maximum displacement magnitude ≤ 0.1 .

the test task as the target set for evaluation, then train each model under the few-shot learning setting with $N^{\text{test}} = \{2, 4, 8, 12\}$ labeled data pairs as the context set. All approaches are developed based on an 32-layer IFNO model.

Results. Besides the diversity of tasks as seen in Fig. 6(a), notice that we also have a small number of training tasks ($H = 9$), and a relatively small training context set size ($N^\eta = 32$). All these facts make the transfer learning on this benchmark dataset challenging. We present the results in Fig. 6(b) and (c). The neural operator model learned by MetaNO again outperforms the baseline single/transfer learning models and the state-of-the-art GBML models. Our MetaNO model achieves 15% error when using only 2 labeled data pair on the test task, while the Single model has high errors due to overfitting. This fact highlights the importance of learning across multi-tasks: when the total number of measurements on each specimen is limited, it is necessary to transfer the knowledge across specimens. Moreover, while MetaNO—, MAML, and ANIL all have a similar performance in this example, the fine-tuning step in MetaNO seems to substantially improve the accuracy, especially when N^{test} gets larger. This observation is consistent with previous finding on varying training task context sizes.

4.3. Example 3: Application on real-world data sets

We now take a step further to demonstrate the performance of our method on a real-world physical response dataset, which is not generated by solving PDEs. We consider the problem of learning the mechanical response of multiple biological tissue specimens from DIC displacement tracking measurements. We measure the biaxial loading of tricuspid valve anterior leaflet (TVAL) specimens from a porcine heart, such that each specimen (as a task) corresponds to a different region of the leaflet. Due to material heterogeneity of biological tissues, these specimens contain different mechanical and structural properties. However, these properties are generally not observable without damaging the tissue. Therefore, this example provides a proof-of-principle demonstration that how MetaNO applies to learning tasks where the governing equations and physical properties are both unknown, and the dataset has unavoidable measurement noise.

Data Generation. We now briefly provide the data generation procedure for the tricuspid valve anterior leaflet (TVAL) response modeling example. To generate the data, we firstly followed the established biaxial testing procedure, including acquisition of a healthy porcine heart and retrieval of the TVAL [61,62]. Then, we sectioned the leaflet tissue and applied a speckling pattern to the tissue surface using an airbrush and black paint [63–65]. The painted specimen was then mounted to a biaxial testing device (BioTester, CellScale, Waterloo, ON, Canada). To generate samples for each specimen, we performed 7 protocols of displacement-controlled testing to target various biaxial stresses: $P_{11} : P_{22} = \{1 : 1, 1 : 0.66, 1 : 0.33, 0.66 : 1, 0.33 : 1, 0.05 : 1, 1 : 0.1\}$. Here, P_{11} and P_{22} denote the first Piola–Kirchhoff stresses in the x - and y -directions, respectively. Each stress ratio was performed for three loading/unloading cycles. Throughout the test, images of the specimen were captured by a CCD camera, and the load cell readings and actuator displacements were recorded at 5 Hz. After testing, the acquired images were analyzed using the digital image correlation (DIC) module of the BioTester’s software. The pixel coordinate locations of the DIC-tracked grid were then exported and extrapolated to a 21 by 21 uniform grid. In this example, we have the DIC measurements on 16 specimens, with 500 data pairs of loadings and material responses from the 7 protocols on each specimen. These specimens are divided into three groups: 12 for the purpose of meta-train, 1 for validation, and 3 for test. This reflects a common challenge in scientific applications, we not only have limited samples per task, the number of available training tasks is also limited. To demonstrate the diversity of these specimens due to the material heterogeneity in biological tissues, in Fig. 7 we plot the processed displacement field of two exemplar training specimens and the validation and test specimens. For each model, the results are reported as the average of all 3 test tasks.

Results. Herein, we aim to model the tissue response by learning a neural operator mapping the boundary displacement loading to the interior displacement field on each tissue specimen. We use 13 specimens for training and validation with context size $N^\eta = 500$, and provide the test results as the average on the rest 3 specimens. With a 4-layer IFNO as the base model, we train each model based on $N^{\text{test}} \in [2, 300]$ samples, and then evaluate the performance on another 200 samples. The results are provided in Fig. 8. MetaNO performs the best among all the methods across all N^{test} , beating MAML and ANIL by a significant margin. Interestingly, MAML and ANIL did not even beat the “Pretrain1” method, possibly due to the low efficacy of the adapting last layers strategy and the small number of training tasks.

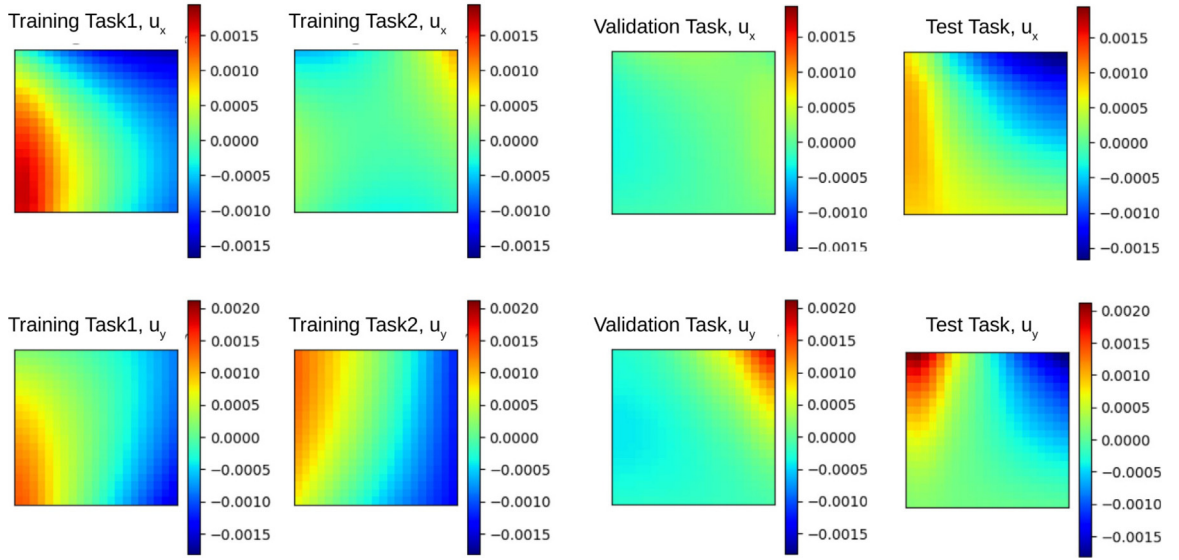


Fig. 7. Visualization of the processed dataset in example 3: learning the biological tissue responses. Subject to the same loading instance, different columns show the corresponding ground-truth solutions $\mathbf{u}^n(x)$ from different tasks, showing the solution diversity across different tasks due to the change of underlying hidden material parameter field.

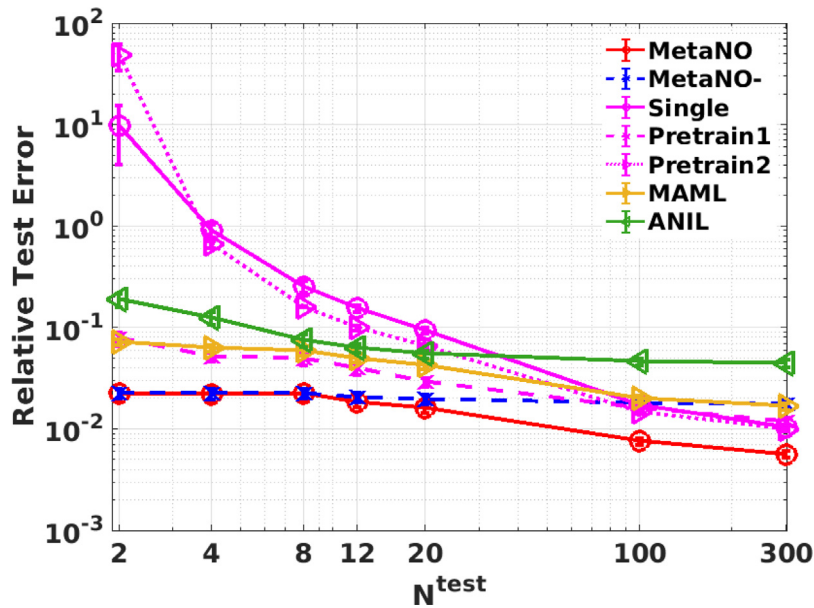


Fig. 8. Comparison of MetaNO and five baseline methods in example 3: learning the biological tissue responses. This example features measurement noise and a small number of available tasks.

5. Summary and future directions

In this paper, we have proposed MetaNO, the first neural-operator-based meta-learning approach that is designed to achieve good transferability in learning complex physical system responses. In particular, we focused on the scenario that the physical responses of different systems can be modeled by a common parametric PDE, and the behavior of each system (specimen) is governed by a different underlying parameter field. Then, based on the observed data from several known systems, the goal was to efficiently learn a surrogate model, which acts as the

surrogate solution operator for the hidden PDE, for a new and unseen system. Comparing with existed transfer-learning techniques, we have shown that our MetaNO features a novel *first layer adaption architecture*, which is theoretically motivated and shown to be a universal solution operator for multiple parametric PDE solving tasks. We have further demonstrated the effectiveness of our proposed MetaNO algorithm on various synthetic, benchmark, and real-world datasets, showing promises with significant improvement in sample efficiency over baseline methods.

Although in this work we focus on material modeling problems in experiments, the proposed approach is a machine learning model that can handle different PDEs with heterogeneous coefficients, which is common in physics problems and not limited in material modeling. For example, another suitable example is climate modeling, where cities responses, temperature exposures, and impacts are all heterogeneous, resulting in different parameters in the PDE model. Here, we chose material as the application mainly due to the availability of data. For future work, we will investigate the applicability of the proposed approach to other scientific domains. Moreover, we point out that the universal approximator property of the MetaNO relies on the assumption that all systems are governed by a common parametric PDE. As such, a set of common characteristic solution operator features can be obtained from the training tasks in the form of the common layer parameters, and then can be utilized to accelerate the learning efficacy in the test task. Therefore, we anticipate the method to suffer from performance deterioration when handling the transferability between different types of PDEs, e.g., when all training tasks are to solve diffusion problems but the test task is on an advection problem. It would be an interesting future direction of extending the proposed MetaNO to handle such a challenging scenario. In this paper, we employed the implicit Fourier neural operator [39] as the base model. As another natural extension, one might consider the proposed meta-learning approach based on other types of neural operators and neural networks, such as the graph neural operators [26,29] and the general graph neural networks [66,67]. As such, the model will be capable to handle measurements from unstructured grid points and domains, which be valuable for applications with complex and possibly evolving domain geometries.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data availability

Data will be made available on request.

Acknowledgments

L. Zhang, H. You and Y. Yu would like to acknowledge support by the National Science Foundation, USA under award DMS-1753031 and the AFOSR, USA grant FA9550-22-1-0197. C.-H. Lee would like to thank the support from the National Institutes of Health (NIH), USA Grant R01 HL159475 and the Presbyterian Health Foundation Team Science Grants. Portions of this research were conducted on Lehigh University's Research Computing infrastructure partially supported by National Science Foundation, USA Award 2019035.

Appendix A. Formulation of baseline methods

In this section, we discuss each baseline methods in details and how they are used in our experiments. A meta-learning baseline in our problem setting would be to apply MAML and ANIL to a neural operator architecture. Here we formally state the implementation of ANIL and MAML for the problem described above, and they will serve as the baseline meta-based methods in our empirical experiments.

MAML. The MAML algorithm proposed in [5] aims to find an initialization, $\tilde{\theta}$, across all tasks, so that new tasks can be learnt with very few gradient updates and examples. First, a batch $\{\mathcal{T}^\eta\}_{\eta=1}^H$ of H tasks are drawn from the training task set. For each task $\mathcal{T}^\eta$, the context set of loading field/response field data pairs $\mathcal{D}^\eta$ is split to a support set of samples, $\mathcal{S}^\eta$, which will be used for inner loop updates, and a target set of samples, $\mathcal{Z}^\eta$, for outer loop updates. Then, for the inner loop, let $\theta^{\eta,0} := \tilde{\theta}$ and $\theta^{\eta,i}$ be the task-wise parameter after i th gradient update. During each inner loop update, the task-wise parameter is updated via

$$\theta^{\eta,i} = \theta^{\eta,i-1} - \alpha \nabla_{\theta^{\eta,i-1}} \mathcal{L}_{\mathcal{S}^\eta}(\theta^{\eta,i-1}), \text{ for } \eta = 1, \dots, H, \quad (\text{A.1})$$

where $\mathcal{L}_{S^\eta}(\theta^{\eta,i-1})$ is the loss on the support set of the η th task, and α is the step size. After m inner loop updates, the initial parameter $\tilde{\theta}$ is updated with a fixed step size β :

$$\tilde{\theta} \leftarrow \tilde{\theta} - \beta \nabla_{\tilde{\theta}} \mathcal{L}_{\text{meta}}(\tilde{\theta}), \text{ where the meta-loss } \mathcal{L}_{\text{meta}}(\tilde{\theta}) := \sum_{\eta=1}^H \mathcal{L}_{Z^\eta}(\theta^{\eta,m}). \quad (\text{A.2})$$

Then, on the test task, $\mathcal{T}^{\text{test}}$, an inner loop adaptation is performed based on few labeled samples $\mathcal{D}^{\text{test}}$ until convergence, and the approximated solution operator model is obtained on the test task as $\tilde{\mathcal{G}}[\mathbf{g}; \theta^{\text{test}}]$.

ANIL. In [10], ANIL was proposed as a modified version of MAML with inner loop updates only for the final layer. The inner loop update formulation of Eq. (A.1) is modified as

$$\theta_Q^{\eta,i} = \theta_Q^{\eta,i-1} - \alpha \nabla_{\theta_Q^{\eta,i-1}} \mathcal{L}_{S^\eta}(\theta_Q^{\eta,i-1}), \text{ for } \eta = 1, \dots, H, \quad (\text{A.3})$$

where $\theta_Q^{\eta,i}$ is the task-wise parameter on the final (projection) layer after i th gradient update. Then, the same outer loop updates are performed following Eq. (A.2).

Single/Pretrain1/Pretrain2. We also implemented 3 non-meta-learning baseline approaches.

- **Single:** Learn a neural operator model only based on the context data set of the test task.
- **Pretrain1:** Pretrain a neural operator model based on all training task data sets, then fine-tune it based on the context test task data set.
- **Pretrain2:** Pretrain a single neural operator model based on the context data set of one training task, then fine-tune it based on the context test task data set. To remove the possible dependency on the pre-training task, in this baseline we randomly select five training tasks for the purpose of pretraining and report the averaged results.

Appendix B. Training details

In the following we briefly describe the hyperparameter settings employed in running of each algorithm. The code and the processed datasets will be publicly released at Github for readers to reproduce the experimental results.

B.1. Example 1: Synthetic data sets

Base model: As the base model for all algorithms, we construct an architecture for IFNO [39] as follows. First, the input loading field instance $\mathbf{g}(\mathbf{x}) \in \mathcal{A}$ is lifted to a higher dimensional representation via lift layer $\mathcal{P}[\mathbf{g}](\mathbf{x})$, which is parameterized as a 1-layer feed forward linear layer with width (3, 32). Then for the iterative layer in Eq. (2.1), we implement $\mathcal{F}^{-1}[\mathcal{F}[\kappa(\cdot; \mathbf{v})] \cdot \mathcal{F}[\mathbf{h}(\cdot, l)]](\mathbf{x})$ with 2D fast Fourier transform (FFT) with input channel and output channel widths both set as 32 and the truncated Fourier modes set as 8. The local linear transformation parameter, W , is parameterized as a 1-layer feed forward network with width (32, 32). In the projection layer, a 2-layer feed forward network with width (32, 128, 2) is employed. To accelerate the training procedure, we apply the shallow-to-deep training technique to initialize the optimization problem. In particular, we start from the NN model with depth $L = 1$, train until the loss function reaches a plateau, then use the resultant parameters to initialize the parameters for the next depth, with $L = 2, L = 4$, and $L = 8$. In the synthetic experiments, we set the layer depth as $L = 8$.

MetaNO: During the meta-train phase, we train for the task-wise parameters θ_p^η and the common parameters θ_l and θ_Q on all 59 training tasks, with the context set of 500 samples on each task. After meta-train phase, we load θ_l and θ_Q and the averaged θ_p^η among all 59 tasks as initialization, then tune the hyperparameters based on the validation task. In particular, the 500 samples on the validation task is split into two parts: 300 samples are reserved for the purpose of training (as the context set) and the rest 200 samples are used for evaluation (as the target set). Then we train for the lift layer on the validation task, and tune the learning rate, the decay rate, and the weight decay parameter for different context set sizes (N^{test}), to minimize the loss on the target set. Based on the chosen hyperparameters, we perform the test on the test task by training for the lift layer on different numbers of samples on its context set, then evaluate and report the performance based on its target set. We repeat the procedure on the test task with selected hyperparameters with five different random seeds, and calculate means and standard errors for the resultant test errors on target set.

MAML & ANIL: For MAML and ANIL, we use the same architecture as the base model, and also split the training tasks for the purpose of training (59 tasks) and validation (1 task) as in MetaNO. During the meta-train phase, for each task we randomly split the available 500 samples to two sets: 250 samples in the support set used for inner loop updates, and the rest in the target set for outer loop updates. During the inner loop update, we train for the task-wise parameter with one epoch, following the standard settings of MAML and ANIL [5,10]. Then, the model hyperparameters, including the learning rate, weight decay, decay rate, and inner loop learning rate, are tuned. In the meta-test phase, we load the initial parameter and train for all parameters (in MAML) or the last-layer parameters (in ANIL) until the optimization algorithm converges. Similar as in MetaNO, we first tune the hyperparameters on the validation task, then evaluate the performance on the test task.

B.2. Example 2: Mechanical MNIST

Base model: As the base model for all algorithms, we construct two IFNO architectures, for the prediction of u_x and u_y , the displacement fields in the x - and y -directions, respectively. On each architecture, the input loading field instance $\mathbf{g}(\mathbf{x}) \in \mathcal{A}$ is mapped to a higher dimensional representation via a lifting layer $\mathcal{P}[\mathbf{g}(\mathbf{x})]$ parameterized as a 1-layer feed forward linear layer with width (4, 64). Then for the iterative layer in Eq. (2.1), we set the number of truncated Fourier mode as 13, and parameterize the local linear transformation parameter, W , as a 1-layer feed forward network with width (64, 64). In the projection layer, a 2-layer feed forward network with width (64, 128, 1) is employed. In this example we also apply the shallow-to-deep technique to accelerate the training, and set the layer depth as $L = 32$.

MetaNO: During the meta-train phase, we train for the task-wise parameters θ_p^η and the common parameters θ_I and θ_Q on all 9 training tasks, with the context set of 32 samples on each task. After the meta-train phase, we load θ_I and θ_Q and the averaged θ_p^η among all 9 tasks as initialization, then train θ_p on the validation task. In particular, the 32 samples on the validation task is split into two parts: 12 samples are reserved for the purpose of training (as the context set) and the rest 20 samples are used for the purpose of evaluation (as the target set). Then we train the lift layer on the validation task, and tune the learning rate, the decay rate, and the weight decay parameter for different context set sizes (N^{test}), to minimize the loss on the target set. Based on the chosen hyperparameters, we perform the meta-test phase on the test task by training for the lift layer on different numbers of samples on its context set, then evaluate and report the performance based on its target set. We repeat the procedure with different 5 random seeds on each of the 5 test tasks, and calculate means and standard errors for the resultant test errors on the target set.

MAML & ANIL: For MAML and ANIL, we use the same architecture as the base model. During the meta-train phase, for each task we randomly split the available 32 samples to two sets: 16 samples in the support set used for inner loop updates, and the rest in the target set for outer loop updates. During the inner loop update, we also follow the standard settings of MAML and ANIL [5,10], and tune the hyperparameters following the same procedure as elaborated above for Example 1.

B.3. Example 3: Experimental measurements on biological tissues

Base model: As the base model, we first construct the lifting layer as a 1-layer feed forward linear layer with width (4, 16). Then for the iterative layer in we keep 8 truncated Fourier modes and parameterize the local linear transformation parameter, W , a 1-layer feed forward network with width (16, 16). In the projection layer, a 2-layer feed forward network with width (16, 64, 1) is employed. We construct two 4-layer IFNO architectures, for the prediction of u_x and u_y , the displacement fields in the x - and y -directions, respectively.

MetaNO: During the meta-train phase, we train for the task-wise parameters θ_p^η and the common parameters θ_I and θ_Q on all 12 tasks, with the context set of 500 samples on each task. After meta-train phase, we load θ_I and θ_Q and the averaged θ_p^η among all 12 tasks as initialization, then tune the hyperparameters based on the validation task. In particular, the 500 samples on the validation task is divided into two parts: 300 samples are reserved for the purpose of training (as the context set) and the rest 200 samples are used for evaluation (as the target set). Based on the chosen hyperparameters, we perform the test on the test tasks by training for the lift layer on different numbers of samples on its context set, then evaluate the performance based on its target set.

MAML & ANIL: For MAML and ANIL, we use the same architecture as the base model, and also split the training tasks for the purpose of training and validation as in MetaNO. During the meta-train phase, for each task we randomly split the available 500 samples to two sets: 250 samples in the support set used for inner loop updates, and the rest in the target set for outer loop updates. During the inner loop update, we train for the task-wise parameter with one epoch, following the standard settings of MAML and ANIL [5,10].

References

- [1] Y. Wang, Q. Yao, J.T. Kwok, L.M. Ni, Generalizing from a few examples: A survey on few-shot learning, *ACM Comput. Surv.* 53 (3) (2020) 1–34.
- [2] G. Koch, R. Zemel, R. Salakhutdinov, et al., Siamese neural networks for one-shot image recognition, in: *ICML Deep Learning Workshop*, vol. 2, Lille, 2015.
- [3] O. Vinyals, C. Blundell, T. Lillicrap, D. Wierstra, et al., Matching networks for one shot learning, *Adv. Neural Inf. Process. Syst.* volume=29.
- [4] J. Snell, K. Swersky, R. Zemel, Prototypical networks for few-shot learning, *Adv. Neural Inf. Process. Syst.* 30.
- [5] C. Finn, P. Abbeel, S. Levine, Model-agnostic meta-learning for fast adaptation of deep networks, in: *International Conference on Machine Learning*, vol. 70, PMLR, 2017, pp. 1126–1135.
- [6] A. Santoro, S. Bartunov, M. Botvinick, D. Wierstra, T. Lillicrap, Meta-learning with memory-augmented neural networks, in: *International Conference on Machine Learning*, PMLR, 2016, pp. 1842–1850.
- [7] A. Antoniou, H. Edwards, A. Storkey, How to train your maml, *arXiv preprint arXiv:1810.09502*.
- [8] S. Ravi, H. Larochelle, Optimization as a Model for Few-Shot Learning.
- [9] A. Nichol, J. Schulman, Reptile: A scalable metalearning algorithm, 2, (3) 2018, p. 4, *arXiv preprint arXiv:1803.02999*.
- [10] A. Raghu, M. Raghu, S. Bengio, O. Vinyals, Rapid learning or feature reuse? Towards understanding the effectiveness of maml, *arXiv preprint arXiv:1909.09157*.
- [11] N. Tripuraneni, C. Jin, M. Jordan, Provable meta-learning of linear representations, in: *International Conference on Machine Learning*, PMLR, 2021, pp. 10434–10443.
- [12] L. Collins, A. Mokhtari, S. Oh, S. Shakkottai, Maml and anil provably learn representations, *arXiv preprint arXiv:2202.03483*.
- [13] W. Zhang, M.S. Sacks, Modeling the response of exogenously crosslinked tissue to cyclic loading: The effects of permanent set, *J. Mech. Behav. Biomed. Mater.* 75 (2017) 336–350.
- [14] M. Misfeld, H.-H. Sievers, Heart valve macro- and microstructure, *Philos. Trans. R. Soc. B* 362 (1484) (2007) 1421–1436.
- [15] J. Rieppo, J. Hallikainen, J.S. Jurvelin, I. Kiviranta, H.J. Helminen, M.M. Hyttinen, Practical considerations in the use of polarized light microscopy in the analysis of the collagen network in articular cartilage, *Microsc. Res. Tech.* 71 (4) (2008) 279–287.
- [16] J. Xu, J.-F. Ton, H. Kim, A. Kosiorek, Y.W. Teh, Metafun: Meta-learning with iterative functional updates, in: *International Conference on Machine Learning*, PMLR, 2020, pp. 10617–10627.
- [17] J. Ghaboussi, D.A. Pecknold, M. Zhang, R.M. Haj-Ali, Autoprogressive training of neural network constitutive models, *Internat. J. Numer. Methods Engng.* 42 (1) (1998) 105–126.
- [18] J. Ghaboussi, J. Garrett, X. Wu, Knowledge-based modeling of material behavior with neural networks, *J. Eng. Mech.* 117 (1) (1991) 132–153.
- [19] G. Carleo, I. Cirac, K. Cranmer, L. Daudet, M. Schuld, N. Tishby, L. Vogt-Maranto, L. Zdeborová, Machine learning and the physical sciences, *Rev. Modern Phys.* 91 (4) (2019) 045002.
- [20] G.E. Karniadakis, I.G. Kevrekidis, L. Lu, P. Perdikaris, S. Wang, L. Yang, Physics-informed machine learning, *Nat. Rev. Phys.* 3 (6) (2021) 422–440.
- [21] L. Zhang, J. Han, H. Wang, R. Car, E. Weinan, Deep potential molecular dynamics: A scalable model with the accuracy of quantum mechanics, *Phys. Rev. Lett.* 120 (14) (2018) 143001.
- [22] S. Cai, Z. Mao, Z. Wang, M. Yin, G.E. Karniadakis, Physics-informed neural networks (PINNs) for fluid mechanics: A review, *Acta Mech. Sin.* (2022) 1–12.
- [23] D. Pfau, J.S. Spencer, A.G. Matthews, W.M.C. Foulkes, Ab initio solution of the many-electron schrödinger equation with deep neural networks, *Phys. Rev. Res.* 2 (3) (2020) 033429.
- [24] Q. He, D.W. Laurence, C.-H. Lee, J.-S. Chen, Manifold learning based data-driven modeling for soft biological tissues, *J. Biomech.* 117 (2021) 110124.
- [25] G. Besnard, F. Hild, S. Roux, Finite-element displacement fields analysis from digital images: Application to portevin-le chatelier bands, *Exp. Mech.* 46 (6) (2006) 789–803.
- [26] Z. Li, N. Kovachki, K. Azizzadenesheli, B. Liu, K. Bhattacharya, A. Stuart, A. Anandkumar, Neural operator: Graph kernel network for partial differential equations, *arXiv preprint arXiv:2003.03485*.
- [27] Z. Li, N. Kovachki, K. Azizzadenesheli, B. Liu, A. Stuart, K. Bhattacharya, A. Anandkumar, Multipole graph neural operator for parametric partial differential equations, *Adv. Neural Inf. Process. Syst.* 33 (2020) NeurIPS 2020.
- [28] Z. Li, N.B. Kovachki, K. Azizzadenesheli, K. Bhattacharya, A. Stuart, A. Anandkumar, Fourier NeuralOperator for parametric partial differential equations, in: *International Conference on Learning Representations*, 2020.
- [29] H. You, Y. Yu, M. D’Elia, T. Gao, S. Silling, Nonlocal kernel network (NKN): A stable and resolution-independent deep neural network, *J. Comput. Phys.* (2022) *arXiv preprint arXiv:2201.02217*.
- [30] Y.Z. Ong, Z. Shen, H. Yang, IAE-NET: Integral autoencoders for discretization-invariant learning, doi:10.13140/RG.2.2.25120.87047/2.

- [31] G. Gupta, X. Xiao, P. Bogdan, Multiwavelet-based operator learning for differential equations, in: A. Beygelzimer, Y. Dauphin, P. Liang, J.W. Vaughan (Eds.), *Adv. Neural Inf. Process. Syst.* (2021) <https://openreview.net/forum?id=LZDiWaC9CGL>.
- [32] L. Lu, P. Jin, G.E. Karniadakis, Deeponet: Learning nonlinear operators for identifying differential equations based on the universal approximation theorem of operators, *arXiv preprint arXiv:1910.03193*.
- [33] L. Lu, P. Jin, G. Pang, Z. Zhang, G.E. Karniadakis, Learning nonlinear operators via DeepONet based on the universal approximation theorem of operators, *Nat. Mach. Intell.* 3 (3) (2021) 218–229.
- [34] S. Goswami, A. Bora, Y. Yu, G.E. Karniadakis, Physics-informed neural operators, 2022, *arXiv preprint arXiv:2207.05748*.
- [35] M. Yin, E. Ban, B.V. Rego, E. Zhang, C. Cavinato, J.D. Humphrey, G. Em Karniadakis, Simulating progressive intramural damage leading to aortic dissection using DeepONet: An operator–regression neural network, *J. R. Soc. Interface* 19 (187) (2022) 20210670.
- [36] M. Yin, E. Zhang, Y. Yu, G.E. Karniadakis, Interfacing finite elements with deep neural operators for fast multiscale modeling of mechanics problems, *Comput. Methods Appl. Mech. Engrg.* 398 (2022) 115027, <http://dx.doi.org/10.1016/j.cma.2022.115027>.
- [37] L. Lu, H. He, P. Kasimbeg, R. Ranade, J. Pathak, One-shot learning for solution operators of partial differential equations, *arXiv preprint arXiv:2104.05512*.
- [38] L. Lu, X. Meng, S. Cai, Z. Mao, S. Goswami, Z. Zhang, G.E. Karniadakis, A comprehensive and fair comparison of two neural operators (with practical extensions) based on fair data, *arXiv preprint arXiv:2111.05512*.
- [39] H. You, Q. Zhang, C.J. Ross, C.-H. Lee, Y. Yu, Learning deep implicit fourier neural operators (IFNOs) with applications to heterogeneous material modeling, *Comput. Methods Appl. Mech. Engrg.* 398 (2022) 115296, <http://dx.doi.org/10.1016/j.cma.2022.115296>.
- [40] S.J. Pan, Q. Yang, A survey on transfer learning, *IEEE Trans. Knowl. Data Eng.* 22 (10) (2009) 1345–1359.
- [41] S. Thrun, L. Pratt, Learning to learn: Introduction and overview, in: *Learning to Learn*, Springer, 1998, pp. 3–17.
- [42] T. Hospedales, A. Antoniou, P. Micaelli, A. Storkey, Meta-learning in neural networks: A survey, 2020, *arXiv preprint arXiv:2004.05439*.
- [43] H. Mai, T.C. Le, T. Hisatomi, D. Chen, K. Domen, D.A. Winkler, R.A. Caruso, Use of meta models for rapid discovery of narrow bandgap oxide photocatalysts, *iScience* (2021) 103068.
- [44] L. Zhang, H. You, Y. Yu, Metanor: A meta-learned nonlocal operator regression approach for metamaterial modeling, *arXiv preprint arXiv:2206.02040*.
- [45] Y. Yin, I. Ayed, E. de Bézenac, N. Baskiotis, P. Gallinari, Leads: Learning dynamical systems that generalize across environments, *Adv. Neural Inf. Process. Syst.* 34 (2021) 7561–7573.
- [46] R. Wang, R. Walters, R. Yu, Meta-learning dynamics forecasting using task inference, *arXiv preprint arXiv:2102.10271*.
- [47] B. Kailkhura, B. Gallagher, S. Kim, A. Hiszpanski, T.Y.-J. Han, Reliable and explainable machine-learning methods for accelerated material discovery, *NPJ Comput. Mater.* 5 (1) (2019) 1–9.
- [48] S. Goswami, K. Kontolati, M.D. Shields, G.E. Karniadakis, Deep transfer learning for partial differential equations under conditional shift with deeponet, *arXiv preprint arXiv:2204.09810*.
- [49] J. Yoon, T. Kim, O. Dia, S. Kim, Y. Bengio, S. Ahn, Bayesian model-agnostic meta-learning, in: *Proceedings of the 32nd International Conference on Neural Information Processing Systems, NeurIPS 2018*, 2018, pp. 7343–7353.
- [50] J. Vanschoren, Meta-learning: A survey, 2018, *arXiv preprint arXiv:1810.03548*.
- [51] H. Yang, J. Kwok, Efficient variance reduction for meta-learning, in: *International Conference on Machine Learning*, PMLR, 2022, pp. 25070–25095.
- [52] K. Kalais, S. Chatzis, Stochastic deep networks with linear competing units for model-agnostic meta-learning, in: *International Conference on Machine Learning*, PMLR, 2022, pp. 10586–10597.
- [53] M. Dejam, H. Hassanzadeh, Z. Chen, Pre-darcy flow in porous media, *Water Resour. Res.* 53 (10) (2017) 8187–8210.
- [54] F. Fallah, M. Ghajari, Y. Safa, Computational modelling of dynamic delamination in morphing composite blades and wings, *Int. J. Multiphys.* 13 (4) (2019) 393–430.
- [55] C. Wei, C.S. Wojnar, C. Wu, Hydro-chemo-mechanical phase field formulation for corrosion induced cracking in reinforced concrete, *Cem. Concr. Res.* 144 (2021) 106404.
- [56] G.A. Holzapfel, T.C. Gasser, R.W. Ogden, A new constitutive framework for arterial wall mechanics and a comparative study of material models, *J. Elast. Phys. Sci. Solids* 61 (1) (2000) 1–48.
- [57] I.B. Bischofs, U.S. Schwarz, Effect of Poisson ratio on cellular structure formation, *Phys. Rev. Lett.* 95 (6) (2005) 068102.
- [58] A. Lang, J. Potthoff, Fast simulation of gaussian random fields, *Monte Carlo Methods Appl.* 17 (3) (2011) 195–214, <http://dx.doi.org/10.1515/mcma.2011.009>.
- [59] M. Alnæs, J. Blechta, J. Hake, A. Johansson, B. Kehlet, A. Logg, C. Richardson, J. Ring, M.E. Rognes, G.N. Wells, The fenics project version 1.5, *Arch. Numer. Softw.* 3, (100).
- [60] E. Lejeune, Mechanical MNIST: A benchmark dataset for mechanical metamodels, *Extreme Mech. Lett.* 36 (2020) 100659.
- [61] C.J. Ross, D.W. Laurence, J. Richardson, A.R. Babu, L.E. Evans, E.G. Beyer, R.C. Childers, Y. Wu, R.A. Towner, K.-M. Fung, A. Mir, H.M. Burkhart, G.A. Holzapfel, C.-H. Lee, An investigation of the glycosaminoglycan contribution to biaxial mechanical behaviours of porcine atrioventricular heart valve leaflets, *J. R. Soc. Interface* 16 (156) (2019) 20190069.
- [62] D. Laurence, C. Ross, S. Jett, C. Johns, A. Echols, R. Baumwart, R. Towner, J. Liao, P. Bajona, Y. Wu, et al., An investigation of regional variations in the biaxial mechanical properties and stress relaxation behaviors of porcine atrioventricular heart valve leaflets, *J. Biomech.* 83 (2019) 16–27.
- [63] D.S. Zhang, D.D. Arola, Applications of digital image correlation to biological tissues, *J. Biomed. Opt.* 9 (4) (2004) 691–699.
- [64] G. Lionello, L. Cristofolini, A practical approach to optimizing the preparation of speckle patterns for digital-image correlation, *Meas. Sci. Technol.* 25 (10) (2014) 107001.
- [65] M. Palanca, G. Tozzi, L. Cristofolini, The use of digital image correlation in the biomechanical area: A review, *Int. Biomech.* 3 (1) (2016) 1–21.

- [66] A. Li, Y.J. Zhang, Isogeometric analysis-based physics-informed graph neural network for studying traffic jam in neurons, *Comput. Methods Appl. Mech. Engrg.* 403 (2023) 115757.
- [67] A. Li, A. Barati Farimani, Y.J. Zhang, Deep learning of material transport in complex neurite networks, *Sci. Rep.* 11 (1) (2021) 11280.