

Introducing a microstructure-embedded autoencoder approach for reconstructing high-resolution solution field data from a reduced parametric space

Rasoul Najafi Koopas^{1*}, Shahed Rezaei², Natalie Rauter¹, Richard Ostwald¹,
Rolf Lammering¹

^{1*}Institute of Solid Mechanics, Helmut-Schmidt University/University of the Federal Armed Forces, Holstenhofweg 85, Hamburg, 22043, Germany.

²Access e.V., Aachen, Germany.

*Corresponding author(s). E-mail(s): najafikr@hsu-hh.de;

Contributing authors: s.rezaei@access-technology.de; natalie.rauter@hsu-hh.de;
ostwald@hsu-hh.de; rolf.lammering@hsu-hh.de;

Abstract

In this study, we develop a novel multi-fidelity deep learning approach that transforms low-fidelity solution maps into high-fidelity ones by incorporating parametric space information into a standard autoencoder architecture. This method's integration of parametric space information significantly reduces the need for training data to effectively predict high-fidelity solutions from low-fidelity ones. In this study, we examine a two-dimensional steady-state heat transfer analysis within a highly heterogeneous materials microstructure. The heat conductivity coefficients for two different materials are condensed from a 101×101 grid to smaller grids. We then solve the boundary value problem on the coarsest grid using a pre-trained physics-informed neural operator network known as Finite Operator Learning (FOL). The resulting low-fidelity solution is subsequently upscaled back to a 101×101 grid using a newly designed enhanced autoencoder. The novelty of the developed enhanced autoencoder lies in the concatenation of heat conductivity maps of different resolutions to the decoder segment in distinct steps. Hence the developed algorithm is named microstructure-embedded autoencoder (MEA). We compare the MEA outcomes with those from finite element methods, the standard U-Net, and various other upscaling techniques, including interpolation functions and feedforward neural networks (FFNN). Our analysis shows that MEA outperforms these methods in terms of computational efficiency and error on test cases. As a result, the MEA serves as a potential supplement to neural operator networks, effectively upscaling low-fidelity solutions to high fidelity while preserving critical details often lost in traditional upscaling methods, particularly at sharp interfaces like those seen with interpolation.

Keywords: Multi-fidelity deep neural network, Multiscaling, U-Net, Neural operator learning

1 Introduction

High-fidelity simulations of physical phenomena are critically dependent on the balance between computational complexity and the accuracy of high-resolution models. In computational mechanics, where precision is essential, the trade-off between computational efficiency and the level of detail in the solution fields poses a significant challenge. High-fidelity models, while offering detailed insights, are computationally intensive and demand substantial resources in terms of both time and computational power. Conversely, low-resolution models, which are less demanding computationally, often lack important details, compromising the accuracy necessary in fields where precision is essential. This challenge highlights the necessity of optimizing both the performance and the cost of computational simulations in scientific research.

To address this challenge in both, scientific research and application, the multi-fidelity modeling method [1, 2] is considered a potential approach to effectively reduce the acquisition cost of high-fidelity data. Accurate but costly, high-fidelity data is typically obtained from detailed experiments or simulations. Conversely, more affordable low-fidelity data, though less accurate, offers a basic representation of physical phenomena. The multi-fidelity approach aims to combine low and high-fidelity data to improve both the cost-effectiveness and accuracy of the physical phenomena under consideration. Traditional multi-fidelity modeling approaches, such as Co-Kriging [3, 4] and multi-level Monte Carlo [5, 6], have proven effective in reducing the acquisition costs of high-fidelity data.

The use of deep neural networks (DNNs) for multi-fidelity modeling takes advantage of their exceptional ability to generalize. The integration of deep neural networks to construct multi-fidelity architectures has become a prominent research topic in recent years. Interested readers are referred to works such as [7–12], focusing on the development of deep neural network architectures. These architectures aim to reduce computational costs for generating high-fidelity solutions by incorporating both low- and high-fidelity data. A common issue with the aforementioned methods is their dependence on multiple neural networks for multi-fidelity modeling, necessitating the integration of various networks for practical inference. This approach often results in redundant structures within the serial multi-fidelity model and can lead to error accumulation and propagation across the series of models. Furthermore, they frequently overlook the physical principles fundamental to engineering within their frameworks. To address this deficiency and reduce the number of labeled data for training, efforts like those cited in [13–16] employ Physics-Informed Neural Networks (PINN), which utilize physical insights to broaden the multi-fidelity modeling scope. Although promising, integrating PINNs into multi-fidelity models is challenging. Up until now, these models require complex tuning due to the increased number of hyperparameters within their loss function, making their training costly and time-consuming. Therefore, utilizing PINNs for solving partial differential equations (PDEs) in a parametric way and for a class of problems with realistic engineering applications is hindered and shall be improved.

Another promising architecture for mapping two different fields is the so-called U-Net method. The U-Net architecture, originally developed for medical image segmentation [17], has attracted great interest in numerous disciplines, such as computational mechanics [18–20], due to its superior performance and adaptability. Known for its efficient management of complex spatial hierarchies, this deep learning framework has demonstrated remarkable versatility and extended its utility beyond its original healthcare context.

It is also worth mentioning that physics-informed operator learning techniques, which can be utilized in handling reduced parameter spaces, are increasingly recognized for their benefits (as investigated in [21–23] and reviewed in [24]). These methods aim to train networks to understand underlying physical equations **parametrically**, applicable across various boundary value problems rather than a specific boundary value problem [21, 25]. This approach extends the principles of PINNs, initially introduced by Raissi et al. [26]. To improve efficiency and accuracy of physics-informed operator learning techniques, Kontolati et al. [27] suggested mapping high-dimensional datasets to a low-dimensional latent space using specialized autoencoders and pre-trained decoders.

To address the challenges associated with multi-fidelity approaches, we introduced a novel architecture called the Microstructure-Embedded Autoencoder (MEA). In this architecture, the low-fidelity solution map is first obtained by a solver on the reduced parametric space. Afterwards, the resulting low-fidelity solution is then upscaled to a high-fidelity map via an enhanced autoencoder.

The innovation of the MEA model is twofold. First, unlike many architectures that rely on picture data for computations, MEA uses the spatial distribution of material properties as input, thereby indirectly incorporating physical influences into the deep neural network architecture. Second, instead

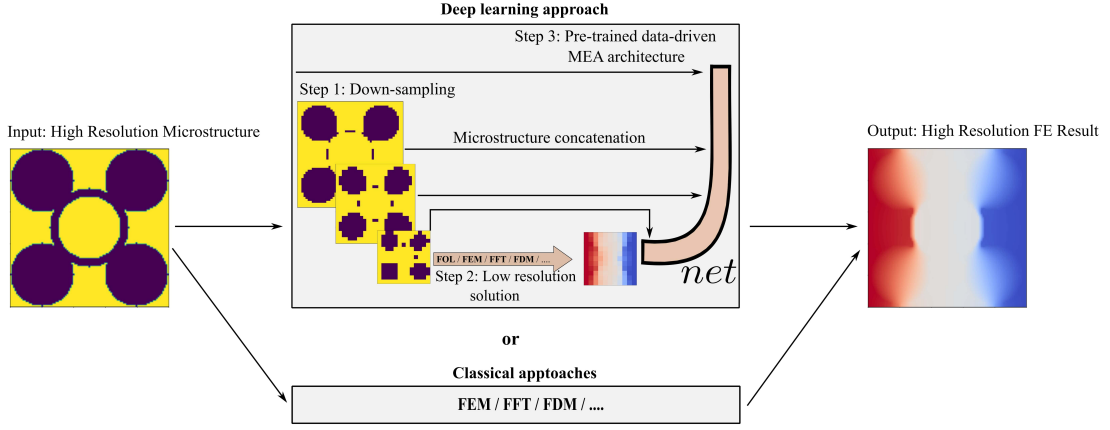


Fig. 1: In the developed MEA architecture, the initial step involves condensing a high-resolution heat conductivity map into various lower-resolution maps. In the next step, the coarsest grid—defined in this study as an 11×11 grid—is utilized to solve the boundary value problem. For this task, one of several techniques, including the Finite Element Method (FEM), Finite Difference Method (FDM), Fast Fourier Transform (FFT), Finite Operator Learning (FOL), or Physics-Informed Neural Networks (PINNs), among others can be implemented. Following the resolution of the boundary value problem at this lower scale, the resultant low-resolution output undergoes an upscaling process through the use of an enhanced autoencoder.

of employing multiple networks and inference processes, MEA uses a single autoencoder to enhance the resolution of the solution map.

In this work, we have used a finite operator learning (FOL) architecture to solve the boundary value problem, which is steady-state heat transfer analysis, in reduced parameter space. This approach eliminates problems commonly associated with other solvers, such as those specific to a boundary value problem, and increases the versatility of MEA model.

As can be seen in Fig. 1, within the MEA architecture, detailed high-resolution fields are reconstructed by integrating the spatial distribution of material properties of different resolutions in different stages of the decoder segment. The general MEA architecture involves a three-step process: first, it starts with a condensation process to generate material property distribution maps of different grid sizes, second, solving the boundary value problem on the coarsest grid with an available numerical solver (FEM, FDM, PINNs, FOL, ...), and finally, scaling to a high-resolution solution map using an enhanced autoencoder architecture. As shown in this work, this method significantly reduces the calculation time compared to traditional methods such as the finite element method, making it both efficient and accurate for computational mechanics. Moreover, to demonstrate the power and robustness of the MEA approach, we conduct a comparative analysis against several established methods. These include the standard U-Net architecture, as well as other techniques that utilize interpolation functions to upscale the solution. The comparative results indicate the superior performance of MEA in terms of accuracy and computational cost.

The structure of this paper is organized as follows: Section 2 describes the boundary value problem employed in this study for data generation and outlines the methodologies used to solve the heat transfer problem, specifically through the finite element method and finite operator learning as numerical solvers. Section 3 details the architecture of the algorithms used to obtain high-fidelity solution maps. These algorithms include the standard U-Net architecture, the microstructure-embedded autoencoder, interpolation techniques, and feed-forward neural networks. In Section 4, following a comprehensive hyperparameter study, the performance of the MEA approach is evaluated against other established methods in the literature. Finally, the paper concludes by summarizing the main results.

2 Problem description

2.1 Necessity of reducing parameter space

Regardless of the numerical method used to solve a given boundary value problem parametrically, it is always beneficial from a computational cost standpoint to reduce the parameter space as much as possible. This reduction facilitates a faster training phase for operator learning algorithms. Moreover, working with a fewer number of inputs is not only more convenient but also easier to interpret. Therefore, we propose a simple approach to reduce the parametric space (microstructure resolution) for our problem. This approach is inspired by algorithms used in convolutional neural networks and is illustrated in Fig. 2.

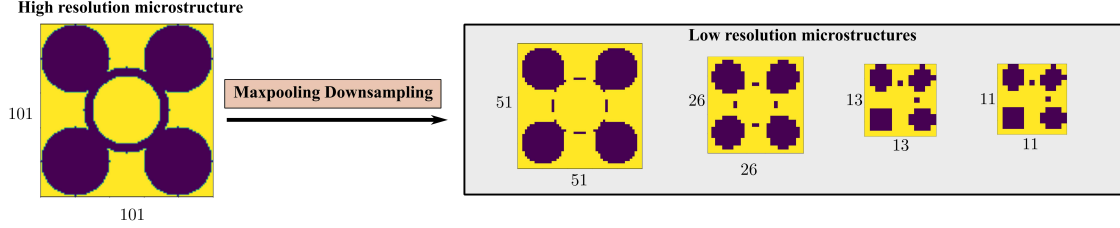


Fig. 2: Condensation of high-resolution parametric space into lower-resolution parametric spaces using the `MaxPool` function from the `SciPy` library.

The reduction of the parametric space involves transforming the high-resolution material property map into lower-resolution representations through **Max Pooling** operations. **Max Pooling** is a process commonly used in image processing and deep learning for reducing the spatial size of an image or matrix while preserving the most significant features, such as edges or textures. This technique involves sliding a window across the input matrix and replacing the output's value with the maximum value within that window. The size of the window determines the degree of downsampling where larger windows result in more significant reduction and coarser representations of the original matrix. In this study, the high-resolution parametric space is condensed into four different resolutions, namely: $N_{51R} = 51$, $N_{26R} = 26$, $N_{13R} = 13$, and $N_{11R} = 11$.

2.2 Problem statement

As shown in Fig. 3, this study addresses a two-dimensional steady-state heat transfer problem in a heterogeneous material. Same derivation principles are applied to analogous problems governed by Poisson's equation, such as electrostatics, electric conduction, magnetostatic problems [28], chemical diffusion [29], and Darcy flow [30] under some assumptions.

In the current work, the parametric input space consists of the two-dimensional spatial distribution of heat conductivity coefficients. Moreover, the solution space represents the temperature field throughout the microstructure. The analysis domain, as depicted on the left side of Fig. 3, is discretized into equally shaped quadrilateral finite element meshes. Additionally, boundary conditions, including Neumann (flux-free) and Dirichlet (with fixed temperature), are illustrated in Fig. 3. The two material phases, differentiated by distinct colors, are homogeneous within their respective phase.

The heat flux vector $\mathbf{q}$ is modelled according to Fourier's assumption,

$$\mathbf{q} = -k(x, y) \nabla T. \quad (1)$$

Here $k(x, y)$ is the phase-dependent heat conductivity coefficient and ∇T represents the spatial gradient of the temperature field T . The governing equation for this problem reads

$$\text{div}(\mathbf{q}) + Q = 0 \quad \text{in } \Omega, \quad (2)$$

$$T = \bar{T} \quad \text{on } \Gamma_D, \quad (3)$$

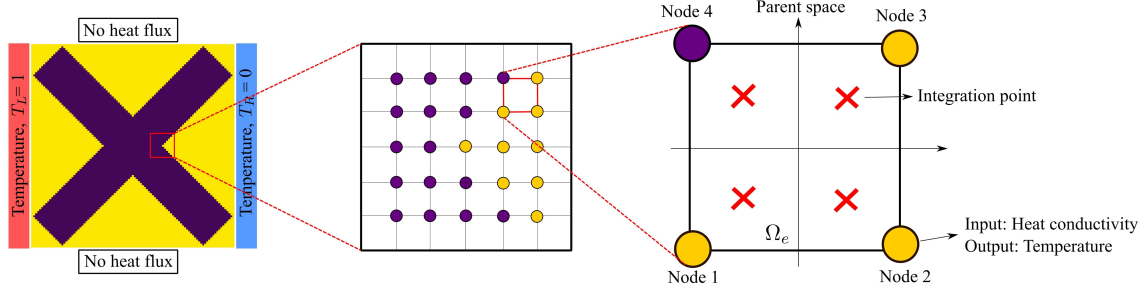


Fig. 3: Description of problem setup, boundary conditions, and finite element meshes. The information of the domain is conveyed through a finite number of sensor points or nodes.

$$\mathbf{q} \cdot \mathbf{n} = q_n = \bar{q} \quad \text{on } \Gamma_N. \quad (4)$$

In the above relations, Q is the heat source term. Moreover, the Dirichlet and Neumann boundary conditions are introduced in Eq. 3 and Eq. 4, respectively.

2.3 Finite Element Modeling

By introducing δT as a test function and with integration by parts, the weak form of the steady-state diffusion problem reads

$$\int_{\Omega} k(x, y) \nabla^T T \delta(\nabla T) \, dV + \int_{\Gamma_N} \bar{q} \delta T \, dA - \int_{\Omega} Q \delta T \, dV = 0. \quad (5)$$

Following the standard finite element method, the temperature field T , conductivity field k as well as their first spatial derivatives, are approximated as

$$T = \sum N_{T_i} T_i = \mathbf{N}_T \mathbf{T}_e, \quad \nabla T = \sum B_{T_i} T_i = \mathbf{B}_T \mathbf{T}_e, \quad k = \sum N_{T_i} k_i = \mathbf{N}_T \mathbf{k}_e. \quad (6)$$

Here, T_i and k_i are the nodal values of the temperature and conductivity field of node i of element e , respectively. Matrices $\mathbf{N}_T$ and $\mathbf{B}_T$ store the corresponding shape functions and their spatial derivatives [24, 31]. To compute these derivatives, we utilize the Jacobian matrix $\mathbf{J} = \partial \mathbf{X} / \partial \boldsymbol{\xi}$, where $\mathbf{X} = [x, y]$ and $\boldsymbol{\xi} = [\xi, \eta]$ represent the physical and parent coordinate systems, respectively [32, 33]. Finally, one can write the discretized version of the weak form for one general finite element as

$$\mathbf{r}_{eT} = \int_{\Omega_e} [\mathbf{B}_T]^T k [\mathbf{B}_T \mathbf{T}_e] \, dV - \int_{\Omega_t} [\mathbf{N}_T]^T k [\mathbf{B}_T \mathbf{T}_e]^T \mathbf{n} \, dS - \int_{\Omega} [\mathbf{N}_T]^T Q \, dV. \quad (7)$$

2.4 Finite operator learning

Motivated by the original formulation of the FEM and the concept of PINNs [26], in a recent study [24], the authors proposed to construct the neural network's loss function based on the governing physical equations derived from the discretized weak form of the problem. Consequently, one can train the deep learning model in a parametric manner and utilize ideas from operator learning to map the physical input field to the solution field by training the deep learning model on a feasible amount of collocation fields (denoted as $k(x, y)$ in this work). It is demonstrated that by adopting this approach, one can achieve accuracy comparable to that of the finite element method while being computationally faster, as it only requires evaluating the network even for unseen cases [24].

Based on the above introduction to the method, the output layer $\mathbf{Y} = T_i$ consists of the components of the temperature field at each node. One can employ separate or fully connected feed-forward neural networks for each output variable. The neural network model is summarized as

$$T_i = \mathcal{N}_i(\mathbf{X}; \boldsymbol{\theta}_i), \quad \mathbf{X} = \{k_j\}, \quad \boldsymbol{\theta}_i = \{\mathbf{W}_i, \mathbf{b}_i\}, \quad i, j = 1 \cdots N. \quad (8)$$

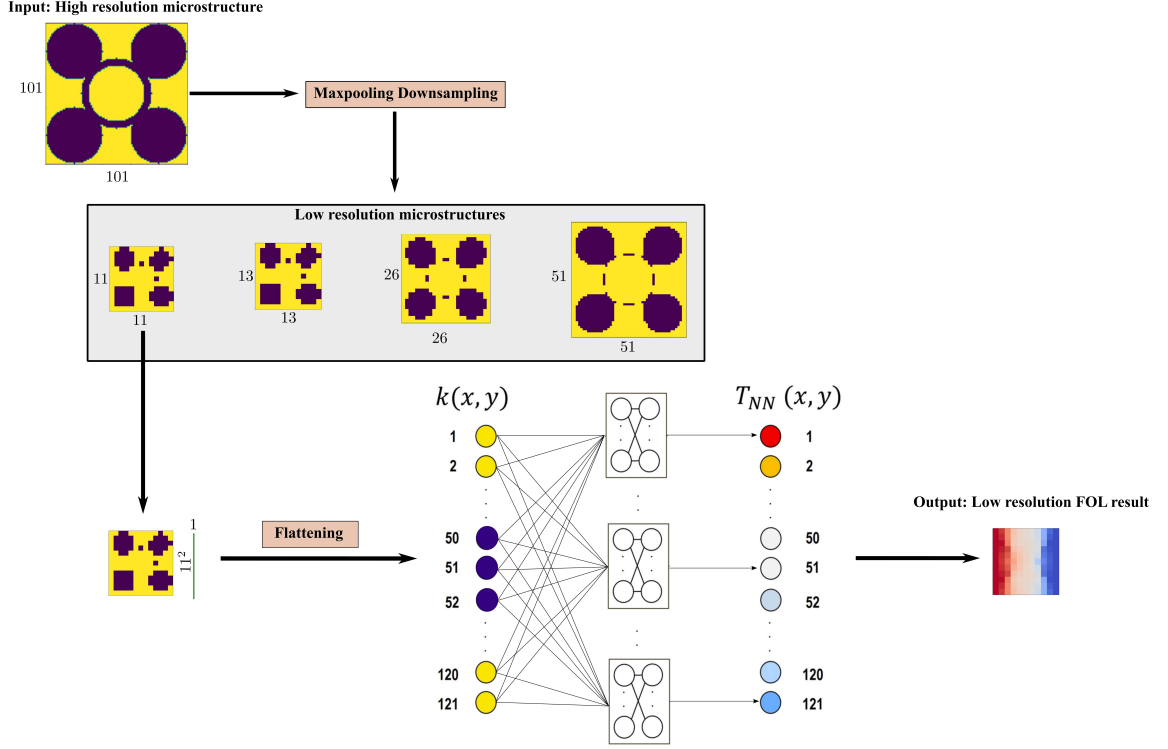


Fig. 4: Network architecture for finite operator learning, where information about the input parameter at each grid point goes in and the The output layer is evaluated by training the network based on the discretized weak form. By satisfying the residual or energy form of the problem through the optimization process in the deep learning model, physical outcomes or solutions are obtained for a parametric input space.

Here, the trainable parameters of the i -th network are denoted by θ_i . Moreover, $\mathbf{W}_i$ and $\mathbf{b}_i$ encompass all the weights and biases within the i -th neural network. Additionally, k_j represents the conductivity value at node j . The number of nodes is denoted by N . Next, we introduce the loss function where we assume no heat source term, i.e., $Q = 0$. Second, we consider the upper and right edges as isolated, resulting in vanishing normal heat flux at the Neumann boundaries (see Fig. 3). The temperature values at the left and right edges are fixed at 1.0 and 0.0 which are applied in a hard way in the network by omitting these nodes from the output layer. Finally, the total loss term $\mathcal{L}_{\text{tot}}$ which is implemented in Python using the JAX [34] and Sciann [35] reads

$$\mathcal{L}_{\text{tot}} = \lambda_e \sum_{e=1}^{n_{el}} \mathbf{T}_e^T \left[\sum_{n=1}^{n_{int}} (N_T \mathbf{k}_e) \mathbf{B}_T^T \mathbf{B}_T \right] \mathbf{T}_e. \quad (9)$$

As depicted in the upper right section of Fig. 3, we utilize the positions of the standard four Gaussian points within the parent space to compute the energy integral. Moreover, we use quadrilateral elements with four integration points (i.e. $w_n = 1$), where the determinant of the Jacobian matrix remains constant. Therefore, we define $\lambda_e = \frac{w_n}{2} \det(\mathbf{J})$. Fig. 4 illustrates the architecture of the FOL used in this study to solve the boundary value problem.

3 Sample generation

In the development of our deep learning models for predicting the heat conductivity of two-phase materials, generating a diverse and representative training dataset is crucial. This dataset comprises spatial distributions that simulate 2D heat conductivity problems. The process begins by defining a square domain discretized into an $N \times N$ grid, with each grid point representing a spatial coordinate in our simulated material sample. For each sample, elliptical inclusions with varying lengths of inner and

outer axes and orientations are introduced, representing the two distinct phases of the material. The conductivities within these phases are determined by predefined values (k_{in} and k_{out}), which represent the thermal conductivities of the inner and outer phases, respectively.

To systematically investigate the space of possible microstructures, a parametric approach is employed where the number of ellipses (n_c), the axes lengths of the ellipses ($a_{\text{outer}}, b_{\text{outer}}, a_{\text{inner}}, b_{\text{inner}}$), and the orientation (θ) are varied according to predefined ranges and steps. This method allows for the generation of a comprehensive set of microstructural configurations. Each configuration's thermal conductivity distribution is computed using a spatial function, which calculates the conductivity at each point in the domain based on the defined microstructure. In total, about 5,670 samples are generated in high resolution (i.e. $N_{\text{HR}} = 101$), with some of them depicted in Fig. 5 for the sake of illustration only.

These datasets are then fed into a finite element solver to compute the corresponding temperature fields. The set of high-resolution input morphology (i.e., the spatial distribution of the heat conductivity coefficients) as well as high-resolution output temperatures are then used for the following three purposes:

- The set of high-resolution input-output pairs is used for training the U-Net model, as will be described in Section 4.3.
- The high-resolution inputs, with $N_{\text{HR}} = 101$, are condensed to lower resolutions of $N_{51\text{R}} = 51$, $N_{26\text{R}} = 26$, $N_{13\text{R}} = 13$, and $N_{11\text{R}} = 11$. Subsequently, the lowest resolution, $N_{11\text{R}} = 11$, is used for unsupervised training of the finite operator learning model described in Section 2.4.
- Finally, the high-resolution heat conductivity maps along with their corresponding condensed maps will be used for training the new microstructure-embedded autoencoder developed in this work.

Referring to Fig. 5, which showcases randomly selected samples from the training dataset, and considering Fig. 6 showing the test samples, it is evident that some of the test cases significantly diverge from the training distribution, categorizing them as out-of-distribution (OOD). These OOD test cases represent scenarios substantially different from those observed during training (e.g. there are no ring, triangular, or rectangular shapes in the training set). This deliberate selection aims to rigorously assess the robustness and generalization capabilities of the neural networks.

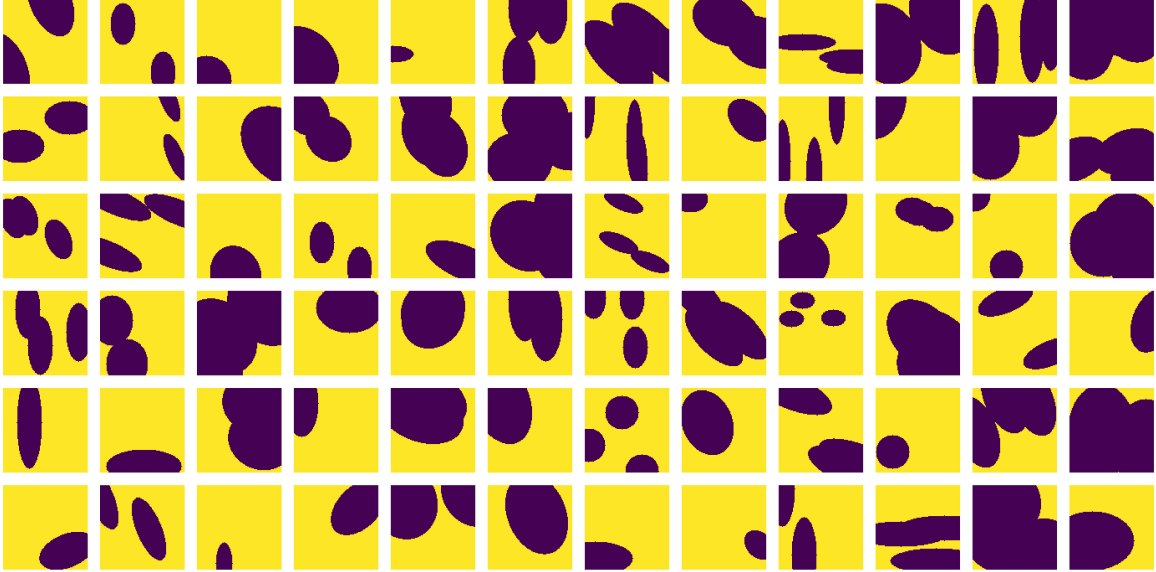


Fig. 5: Randomly selected training samples from a dataset of 5,670 illustrate the spatial conductivity maps of two-phase materials within the training set.

Additionally, for visualization and further analysis, the phase fraction of the generated microstructures, along with the test cases, are illustrated in Fig. 7 to showcase the diversity of the dataset. It's worth mentioning that the chosen database leans more towards higher phase fractions, which means the selected test cases are not among those with the highest frequency.



Fig. 6: Out-of-Distribution test samples. These samples, not seen during training, test the model’s ability to adapt to new, unseen cases.

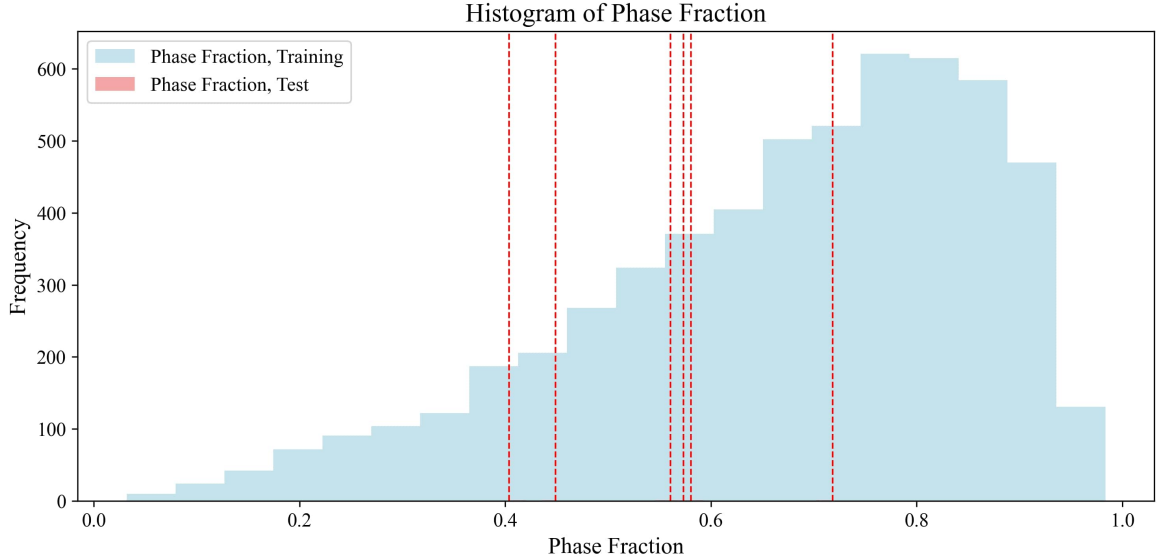


Fig. 7: Phase fraction of the generated microstructures. Phase fraction refers to the proportion of each phase within a material.

For the results in Fig. 7, phase fraction is calculated by determining the ratio of the volume or area occupied by the phase of interest to the total volume or area of the material.

4 Enhancing resolution of solution: A methodological Study

In this section, we comprehensively describe the architecture of the implemented approaches to increase the resolution of the solution map. The methods analyzed include a simple interpolation approach, a fully connected feed-forward neural network, two different versions of the MEA architecture, as the novelty of this paper, and the standard U-Net architecture. The motivation behind this study is to compare the performance of various approaches against our developed algorithm. It should be noted that the codes for all architectures have been developed in PyTorch.

4.1 Interpolation approach

The simplest and most direct method for rescaling a solution field to a higher resolution involves interpolating the solutions using a selected order, such as linear, quadratic, or cubic. This process mirrors the method used in the finite element method, which utilizes the capabilities of shape functions. More specifically, we utilized the `ndimage.zoom` command from the `SciPy` library to resize an n -dimensional array. The `order` parameter within this command determines the interpolation method applied during the resizing process, effectively interpolating the values between existing data points to fill gaps when the array is enlarged. This interpolation can lead to smoothing, as the chosen method (specified by the `order` parameter) introduces smoother transitions between pixel values. For instance,

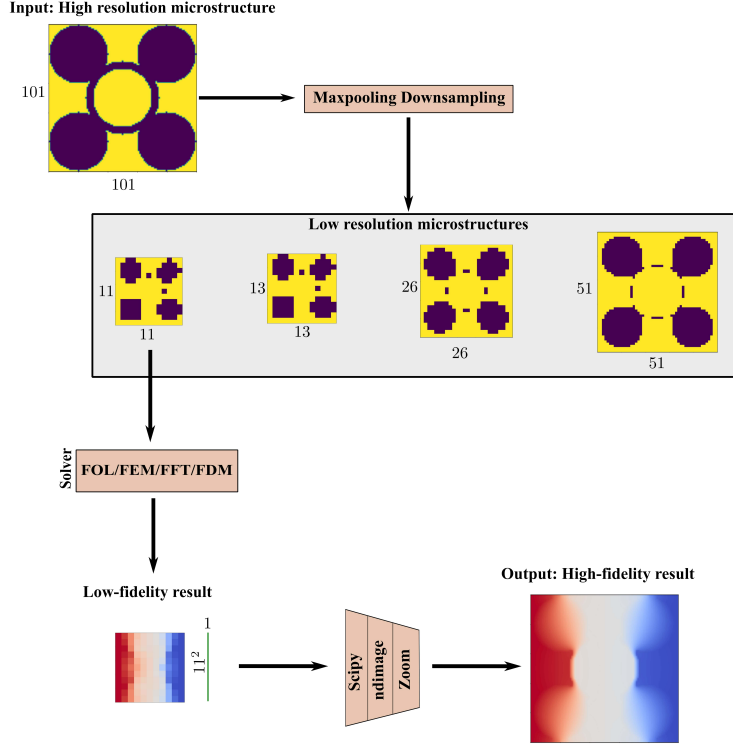


Fig. 8: Implemented interpolation architecture to obtain a high-fidelity solution map from a low-fidelity one. After condensing the high-resolution parametric space into lower resolutions, the low-fidelity solution is obtained by solving the boundary value problem on the coarsest grid using FEM/FOL solver. This low-fidelity solution is then upscaled using an interpolation function to achieve higher fidelity.

higher-order interpolation methods, such as cubic, typically produce smoother results compared to lower-order methods like linear. This aspect of interpolation is comparable to the post-processing phase in the finite element method, where different shape functions of varying orders are employed to evaluate and interpolate solutions across the element domain based on the nodal solution values.

Fig. 8 illustrates the architecture of the interpolation approach used in this study. The input consists of the high-resolution spatial distribution of each coefficient, which is then condensed to lower resolutions $N_{51R} = 51$, $N_{26R} = 26$, $N_{13R} = 13$, and $N_{11R} = 11$ using the **Max Pooling** function. In this work, By adjusting the **pool-window-size** parameter, the original 101×101 heat conductivity maps are condensed into four different lower-resolution versions. In the next step, the lowest resolution, $N_{11R} = 11$, is then fed into a numerical solver which potentially can be methodologies such as Finite Element Method (FEM), Finite Difference Method (FDM), or Fast Fourier Transform (FFT), to generate corresponding low-resolution FE results. In this work physics operator learning (FOL) is used to capture the low-fidelity solution. Afterward, this solution is upscaled using the `SciPy.ndimage.zoom` function.

4.2 Feedforward neural network approach

This method is based on a fully connected Feed Forward Neural Network (FFNN). Each neural network follows the standard structure, featuring a single input layer, potentially several hidden layers, and an output layer. Every layer is interconnected, transferring information to the next layer [36]. Within each layer, the neurons are not interconnected. Thus, we denote the information transfer from the $l-1$ layer to l using the vector $\mathbf{z}^l$. Each component of vector $\mathbf{z}^l$ is computed by

$$z_m^l = a\left(\sum_{n=1}^{N_l} w_{mn}^l z_n^{l-1} + b_m^l\right), \quad l = 1, \dots, L. \quad (10)$$

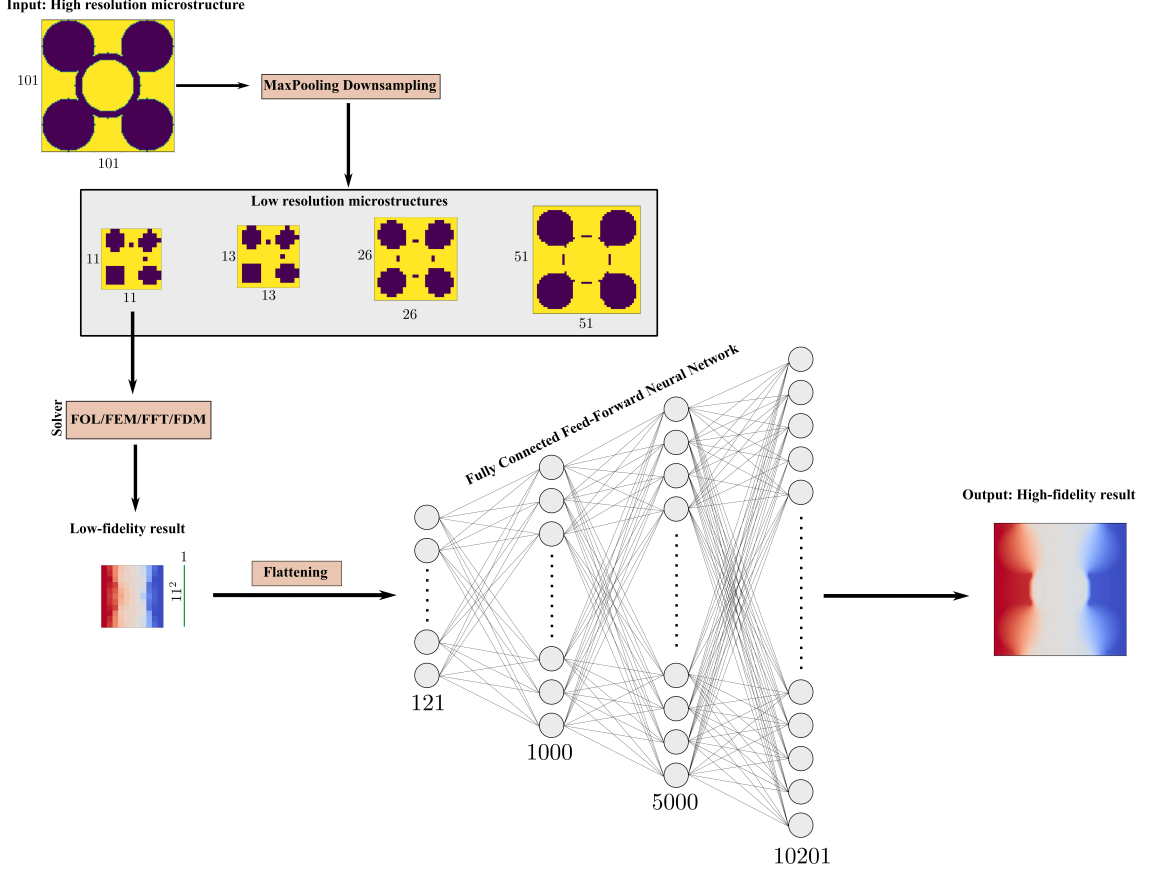


Fig. 9: Fully connected feed-forward neural network (FFNN) architecture to obtain a high-fidelity solution from a low-fidelity one. Similar to the interpolation approach, the low-fidelity solution is first obtained using an FEM/FOL solver. In the FFNN architecture, however, the low-fidelity solution is first flattened and then fed into a fully connected neural network. The output of the FFNN is designed to yield the high-fidelity solution.

In Eq. (10), z_n^{l-1} , is the n -th neuron within the $l-1$ -th layer. The component w_{mn} shows the connection weight between the n -th neuron of the layer $l-1$ and the m -th neuron of the layer l . Every neuron in the l -th hidden layer owns a bias variable b_m^l . The number N_l corresponds to the number of neurons in the l -th hidden layer. The total number of hidden layers is L . The letter a stands for the activation function in each neuron. The activation function $a(\cdot)$ is usually non-linear.

The principle idea behind the implemented FFNN is depicted in Fig. 9. Initially, similar to the interpolation approach prescribed in section. 4.1, a high-resolution microstructural representation, represented by a 101×101 grid, undergoes dimensionality condensation via the **Max Pooling** algorithm, yielding a series of lower-resolution counterparts. After resolution degradation, the microstructure with a resolution of 11×11 is subjected to the FOL solver.

In the FFNN approach, a critical preprocessing step involves the transformation of the two-dimensional low-fidelity solution into a one-dimensional array, commonly referred to as flattening. This is a prerequisite for input compatibility with the FFNN architecture. The FFNN in this study is structured with an input layer succeeded by hidden layers, comprising 121, 1000, and 5000 neurons, respectively, resulting in an output layer with 10201 neurons. Within this approach, the low-fidelity solution is converted to the high-fidelity one. We utilized the Swish activation function, with a learning rate set to $1e-4$. We utilized a batch size of 100 and employed the Adam optimizer for training.

4.3 Standard U-Net

In this study, we utilize a classic U-Net architecture to predict a high-fidelity solution map. The network architecture, illustrated in Fig. 10, diverges from the traditional U-Net application where images are

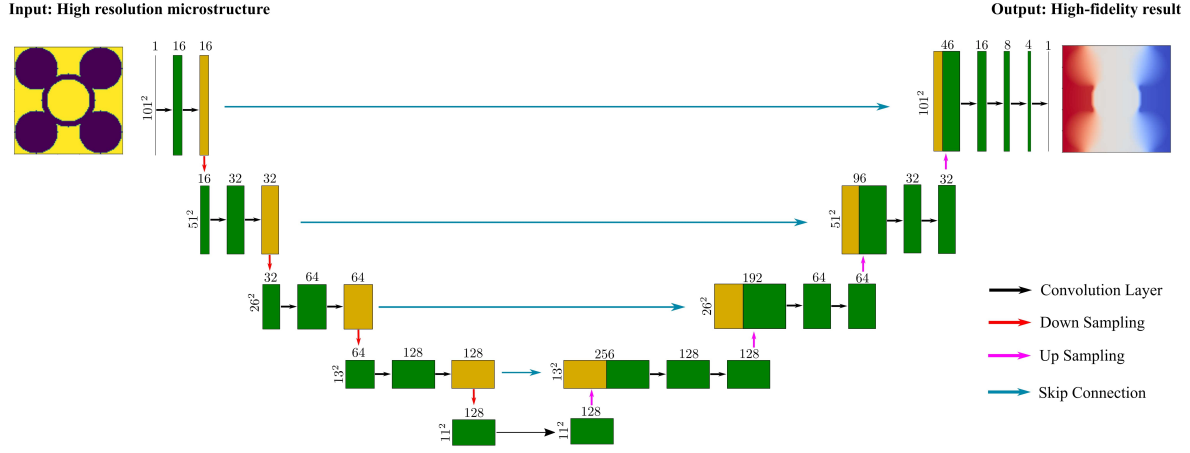


Fig. 10: The classical U-Net architecture employed in the current study. The input to the U-Net is a high-resolution heat conductivity coefficient map, and the output is a high-resolution solution. Within the U-Net’s decoder component, the input undergoes downsampling through four stages, resulting in a 10×10 grid featuring 11 nodes in both horizontal and vertical directions, alongside 128 channels. Following this, the input is upsampled in four stages, with skip connections from the corresponding layers of the decoder being applied at each step to enhance information flow and preserve spatial details.

the typical input; here, the spatial distribution of heat conductivity coefficients serves as the input. The U-Net structure in this work processes an input of a $1 \times 101 \times 101$ high-resolution heat conductivity map of the microstructure, generating a high-fidelity temperature distribution as the output.

In the architecture of the standard U-Net, the encoder is designed to sequentially compress the input data while extracting critical features. It starts with a $1 \times 101 \times 101$ input representing the high-resolution thermal conductivity coefficients of the microstructure. The encoder consists of twelve convolutional layers grouped in blocks, each designed to progressively reduce the spatial dimensions and increase the channel depth to encode the input data into a compact, feature-rich representation. The convolutional layers are characterized by a kernel size of 3, with strides of 1 or 2 and padding 1 to maintain or reduce dimensions as needed. Each convolution operation is followed by batch normalization and ReLU activation functions that ensure efficient feature extraction and introduction of non-linearity. The design of the encoder facilitates the acquisition of hierarchical features from the spatial heat conductivity map, reducing the final output dimension to $128 \times 11 \times 11$, ready for the decoder phase.

The decoder part of the U-Net architecture inversely mirrors the encoder and aims to reconstruct the high-resolution output from the encoded features. It utilizes a series of convolutional layers and up-convolutional layers along with skip connections that integrate corresponding encoder features to enrich the decoder input at various stages. Starting from a compact $128 \times 11 \times 11$ representation, the decoder gradually increases the spatial dimensions while refining the feature maps to reconstruct the final output. This is achieved by alternating convolutional layers that increase feature resolution and upsample layers that expand spatial dimensions, aided by skip connections that reintegrate spatial details from the encoder. The upsampling process ends in a final convolution layer that produces the high-fidelity solution with dimensions that match the original input, effectively capturing the detailed temperature distribution within the microstructure.

4.4 Microstructure-embedded autoencoder (MEA)

Similar to the previous approaches, a crucial step in MEA architecture involves creating lower-resolution representations of heat conductivity maps. To achieve this, once again we employed the **Max Pooling** operation by utilizing the **maximum-filter** function from the **scipy.ndimage** library in Python.

In the subsequent phase of the analysis, the lowest resolution heat conductivity map, now represented by an 11×11 matrix, is inputted into a pre-trained FOL network to solve the boundary value problem. It’s important to note, however, that solving this problem isn’t strictly limited to FOL; alternative solvers can also be utilized based on the specific requirements or constraints of the study.

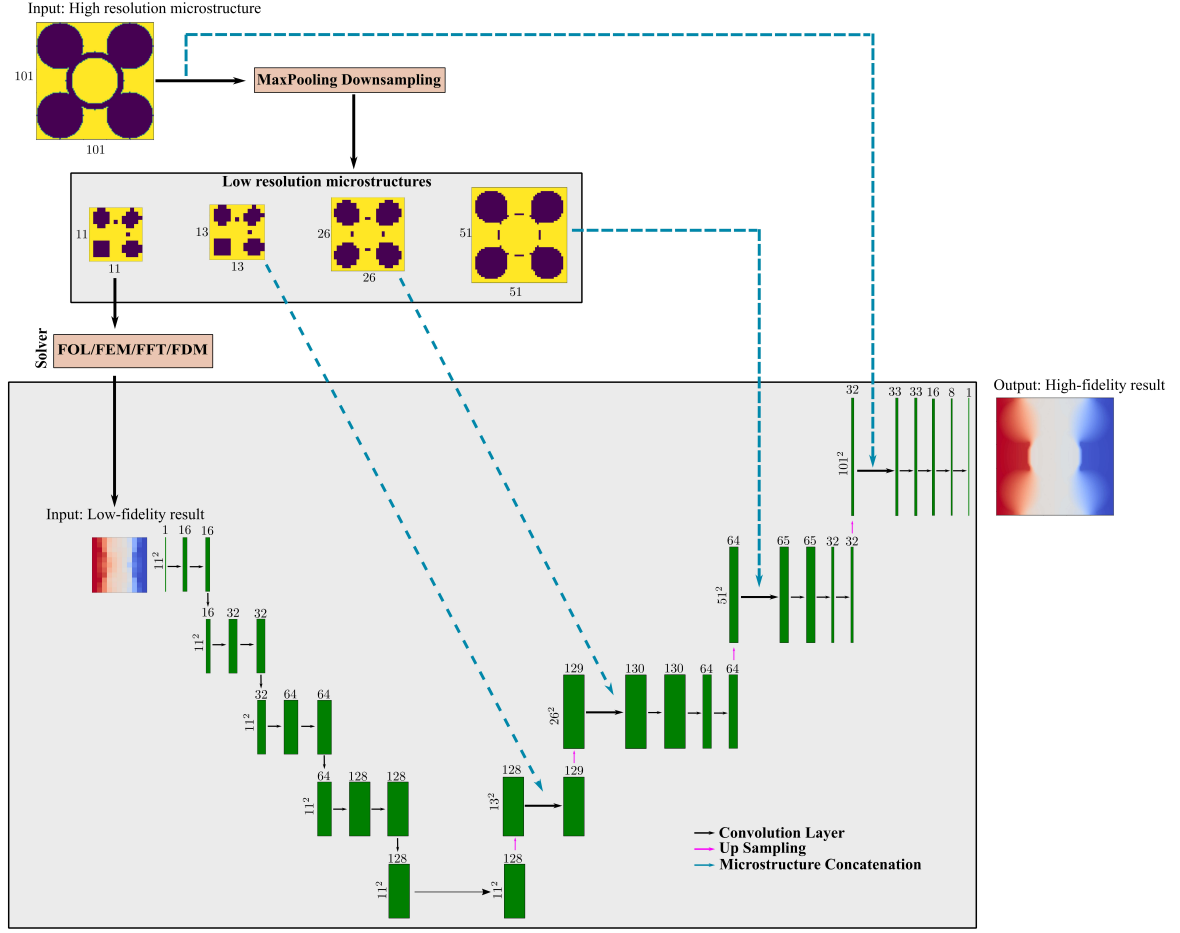


Fig. 11: microstructure-embedded autoencoder architecture for a high-resolution output generation: Input is a low-resolution solution that is obtained from PINNs on an 11×11 node grid in this study and output is a high-resolution 100×100 grid with 101 nodes in the vertical and horizontal direction. The decoder integrates heat conductivity maps of sizes 13×13 , 26×26 , 51×51 , and 101×101 , concatenating them at corresponding stages of the decoder for enhanced output resolution.

As can be seen in Fig. 11, the encoder segment of MEA architecture starts with an input of dimensions $1 \times 11 \times 11$, which represents a low-resolution solution of a steady-state heat transfer problem. It consists of twelve convolutional layers, each with a kernel size of 3, a stride of 1, and a padding of 1. This configuration ensures that the spatial dimensions are maintained while the depth of the feature maps increases, allowing detailed feature extraction without loss of spatial information. Starting from a single channel, the encoder increases the depth to 128 channels through successive convolutional stages, each followed by batch normalization and ReLU activation to improve feature normalization and introduce nonlinearity.

The decoder part of the MEA architecture is specifically designed to reconstruct the input's high-resolution spatial features from the feature representation (channels) produced by the encoder. This process involves expanding the spatial dimensions of the feature maps while integrating additional morphological information to enhance the fidelity of the reconstruction. The decoder segment begins its process with a feature map of dimensions $128 \times 11 \times 11$, which is the output of the encoder section. The decoder section encompasses four upsampling operations, systematically enlarging the spatial dimensions of the feature maps at each stage.

Beginning with the channels received from the encoder, the decoder initiates its first upsampling to expand the dimensions to 13×13 . This step is followed by subsequent enlargements to 26×26 , then to 51×51 , and ultimately achieving a spatial dimension of 101×101 . After each upsampling operation, the corresponding morphological map which is the spatial distribution of the heat conductivity is integrated with the upsampled feature map in a concatenation step, so that the decoder's input is enriched

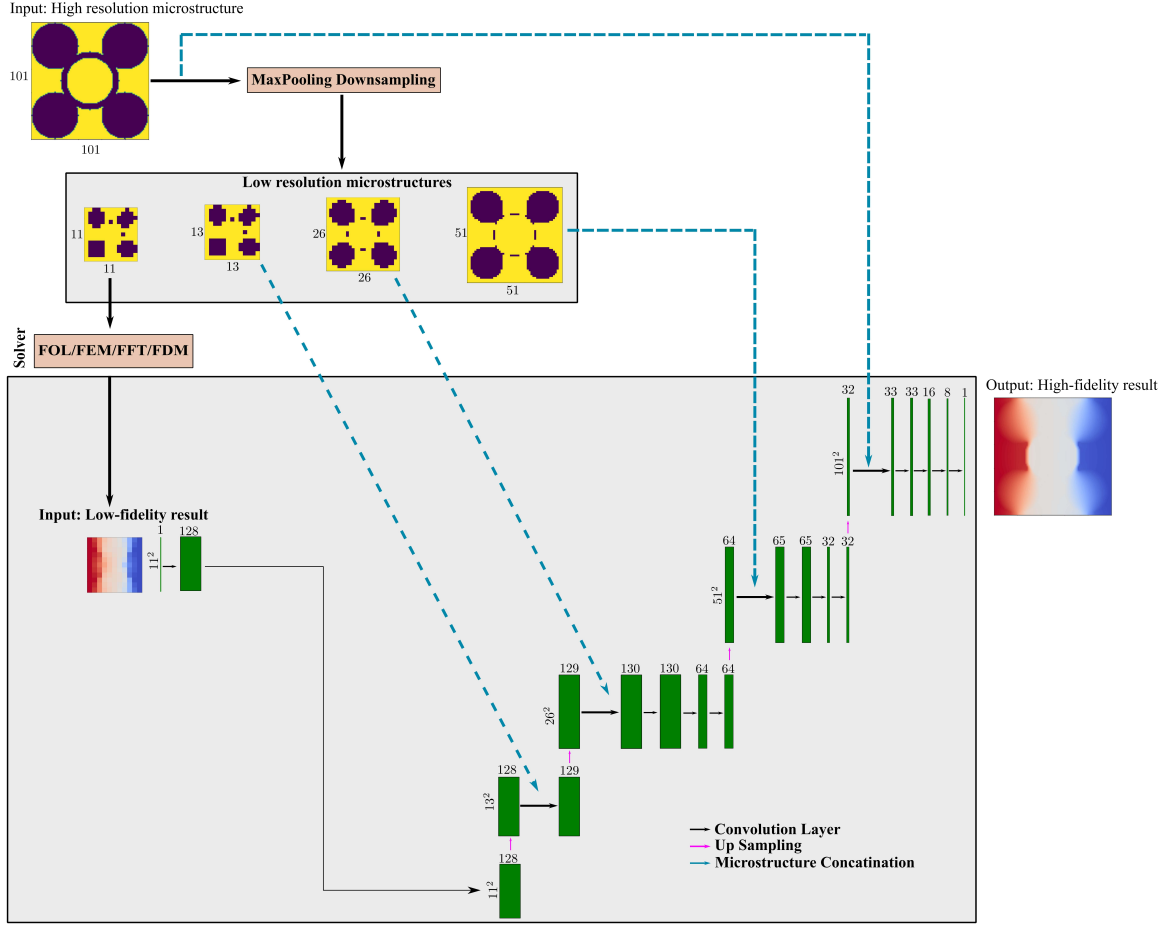


Fig. 12: Second version of the microstructure-embedded autoencoder with a simplified encoder featuring a single convolution layer, maintaining the original decoder structure with microstructure concatenation as shown in Fig.11.

with essential external spatial details that are crucial for accurate reconstruction. In general, the concatenation operation involves joining two or more matrices along a specified direction. For example, if matrix $A_{m \times n}$ and matrix $B_{m \times p}$, then concatenating these two matrices along the first dimension would result in matrix $C_{m \times (n+p)}$. Subsequent convolutional layers refine these merged feature maps by adjusting the channel depths and enhancing the spatial details to ensure that the morphological information is perfectly integrated.

This intricate decoder structure, characterized by the integration of morphological maps and a careful feature refinement process, is critical to the network's ability to reconstruct a high-fidelity solution map from the low-resolution one. Through a strategic combination of convolutional and transposed convolutional operations and the selective incorporation of external morphological data, the architecture ensures that the reconstructed output authentically reflects the high-resolution attributes of the original input.

Considering the architecture of the modified autoencoder, in the encoder section, the low-fidelity solution undergoes twelve convolutional transformations to generate an encoded output with 128 channels. This raises a question about the necessity of utilizing 9 convolutional layers, given that a similar number of channels can be produced by a single convolutional layer in the decoder section. To address this question, an alternative architecture is presented, as shown in Fig. 12, where only a single convolutional layer is employed to generate an output that meets the input requirements for the decoder section. In the results section, the impact of the number of convolutional layers in the decoder section on the prediction accuracy of the enhanced autoencoder is discussed.

5 Results

In this section, we analyze the prediction outcomes of the implemented architecture to obtain the high-fidelity solution. As observed in Table 1, the analyzed methods include the interpolation approach, FFNN, two different MEA algorithms, and a standard U-Net, which are applied across six test cases (see Fig. 6). Additionally, Table 1 also presents the methods implemented to obtain a low-fidelity solution after the condensation of the high-fidelity parametric space.

Table 1: Methods for obtaining a high-fidelity solution map from high-fidelity microstructure data analyzed in this study.

Method	Solver
Interpolation approach	only low-fidelity FOL/FEM
Feed Forward Neural Network	low+high-fidelity FOL/FEM
Microstructure-embedded autoencoder, Type 1	low+high-fidelity FOL/FEM
Microstructure-embedded autoencoder, Type 2	low+high-fidelity FOL/FEM
Classical U-Net	only offline high-fidelity FEM

5.1 Hyperparameter analysis

In this section, we investigate the effects of two crucial hyperparameters, batch size and the extent of microstructure concatenation, on the prediction accuracy of the models. We set the learning rate for all DNNs in this study to 1×10^{-4} and use the Adam optimizer to optimize weights and biases. The dataset is consistently split, with 80% used for training and the remaining 20% for validation. All DNN analyses are performed over a uniform duration of 500 epochs, with mean squared error (MSE) serving as the loss function for the analyses.

In the first study, we investigate the effects of microstructural concatenation, focusing on the integration of different resolutions. The MEA architecture shown in Fig. 11 is characterized by a unique approach in the decoder part: The spatial distribution of thermal conductivity coefficients from four different resolutions is concatenated once each. This approach prompts us to investigate the effects of concatenating microstructure information from multiple resolutions. To address this question, we conducted a comprehensive study to investigate the influence of different numbers of resolution-specific concatenations.

To investigate the effects of varying concatenations, modifications were made to the decoder segment of the MEA architecture. Specifically, the approach to concatenation was adjusted to evaluate different numbers of concatenations involving spatial distributions of heat conductivity coefficients at various resolutions. In *Concatenation 1*, only the distribution with a resolution of $N_{13R} = 13$ is concatenated. *Concatenation 2* involves combining the spatial heat conductivity coefficients at resolutions of $N_{13R} = 13$ and $N_{26R} = 26$. For *Concatenation 3*, resolutions of $N_{13R} = 13$, $N_{26R} = 26$, and $N_{51R} = 51$ are concatenated, and finally, in *Concatenation 4*, heat conductivity coefficients from the resolutions of $N_{13R} = 13$, $N_{26R} = 26$, $N_{51R} = 51$, and $N_{101R} = 101$ are concatenated together. A batch size of 50 is used in this study.

Figure 13 presents the results from analyzing the impact of the number of microstructural concatenations. It is evident from the error graph that an increase in the number of concatenations correlates with a reduction in error. Given that in the process of microstructural concatenation, only a single channel is added to the corresponding stage of the decoder section, this approach does not significantly affect computational cost. Therefore, for all subsequent studies, the *Concatenation 4* method is employed.

The other hyperparameter analyzed is the batch size, a crucial factor in deep learning, which refers to the number of samples processed in a single forward and backward pass of the network. It establishes a balance between accuracy, efficiency, and speed. Larger batch sizes can speed up training but can reduce accuracy and increase the risk of overfitting. Conversely, smaller stacks improve accuracy by acting as regularizers and reducing generalization error, although they require more computational resources and increase training time. In this study, we investigate the influence of four different batch sizes on model performance, implementing the *Concatenation 4* technique as a uniform experimental condition to ensure consistency.

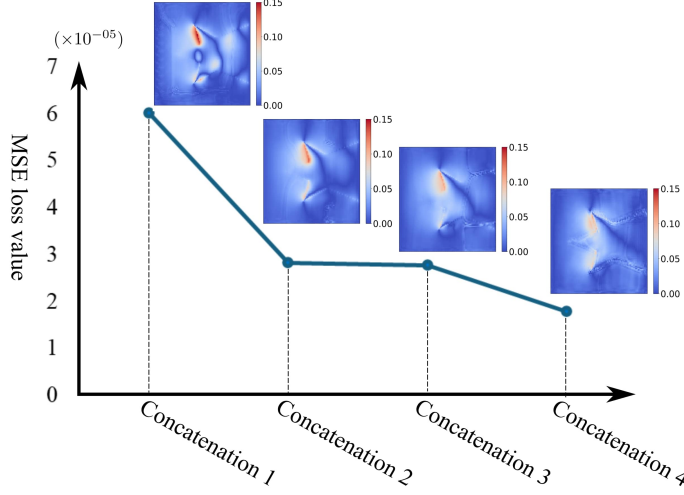


Fig. 13: Exploration of various concatenation strategies within the decoder segment of the enhanced autoencoder. Each panel represents a distinct scenario of concatenation: Concatenation 1 illustrates the model utilizing a singular 13×13 resolution of concatenation. Moving to Concatenation 2, both 13×13 and 26×26 resolutions are concatenated. Concatenation 3 broadens the scope by including 13×13 , 26×26 , and 51×51 resolutions. The most comprehensive strategy, Concatenation 4, applies data from 13×13 , 26×26 , 51×51 , and 101×101 resolutions. The color maps across the panels highlight the absolute error between the network’s predictions and the ground truth calculations.

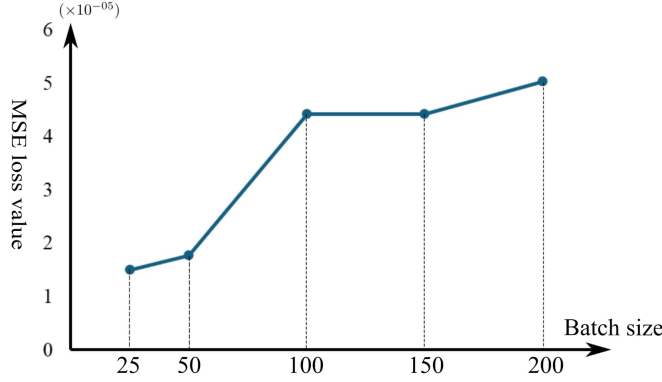


Fig. 14: Investigation of the effects of varying batch sizes on DNNs performance, under a uniform experimental setup employing the Concatenation 4 technique. Throughout this study, the learning rate and optimizer settings remain constant.

Figure 14 depicts the impact of varying batch sizes on the model’s error rate. It is observed that reducing the batch size leads to a decrease in error, highlighting a direct relationship between batch size and model accuracy. However, this reduction in batch size comes at the cost of increased computational resources. After considering the balance between error rates and computational demand, a batch size of 50 is identified as the optimal choice for further investigations. This batch size effectively balances the trade-off between achieving high accuracy and managing computational costs.

5.2 Quality of high-fidelity solutions from studied architectures

Following a comprehensive hyperparameter analysis with a focus on batch size and concatenation strategies, we proceeded to assess the performance of the utilized architectures on out-of-sample test cases, as illustrated in Fig. 6. Beyond examining the direct outputs of these DNNs, which represent 2D

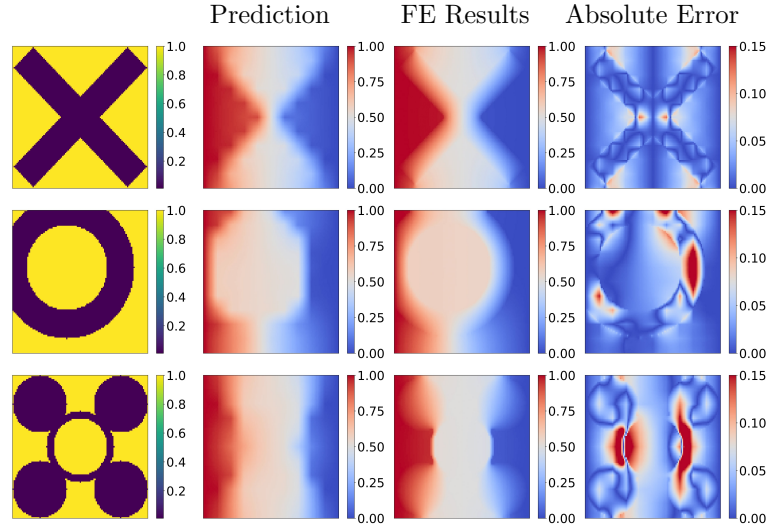


Fig. 15: Prediction results of thermal diffusion using the interpolation approach across three test scenarios. The contour plot in this and other similar figures shows the spatial distribution of the temperature.

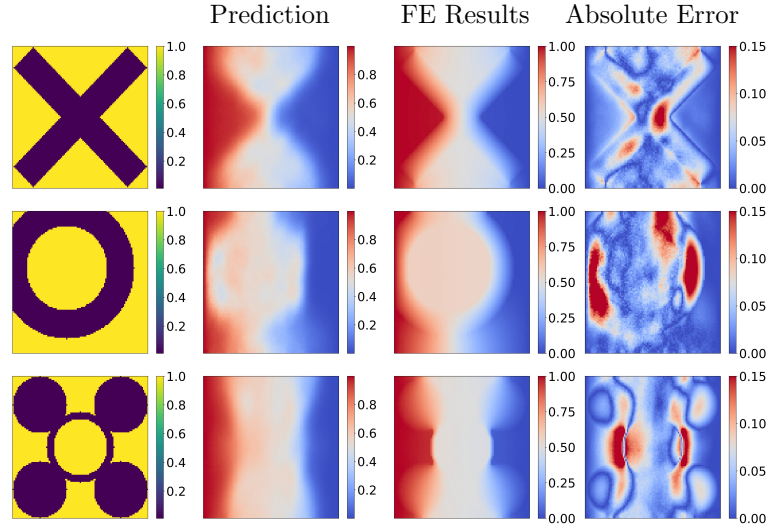


Fig. 16: Prediction results of thermal diffusion using the FFNN across three test scenarios.

temperature distributions, we deepen our comparison by computing the gradients of the DNN outputs. These gradients were then multiplied with spatial heat conductivity coefficient maps to calculate the heat flux. This additional step was crucial for a thorough evaluation, enabling us to determine whether the DNNs used in this study could accurately capture gradient terms. The results of the heat flux calculations are summarized in Appendix A.

Fig. 15 and Fig. 16 present the prediction results from the interpolation and FFNN architectures, respectively. As shown in the figures, both the interpolation and feedforward neural network techniques frequently fail to accurately capture detailed edge patterns. This case is clearly visible in Fig. 17. This discrepancy is due to several factors. Interpolation methods, including linear, quadratic, and cubic approaches, rely on mathematical formulas to interpolate new data points from a known set of points. Although useful, these methods are basic and localized. They may not fully capture the complexities of patterns that are highly variable or non-linear, often smoothing out results and failing to represent sharp edges and complex details. On the other hand, feedforward neural networks also confront a unique set of challenges. By processing input data as a one-dimensional array (in flattening operation), FFNNs lose critical spatial information that is essential for detecting complex patterns, such as edges.

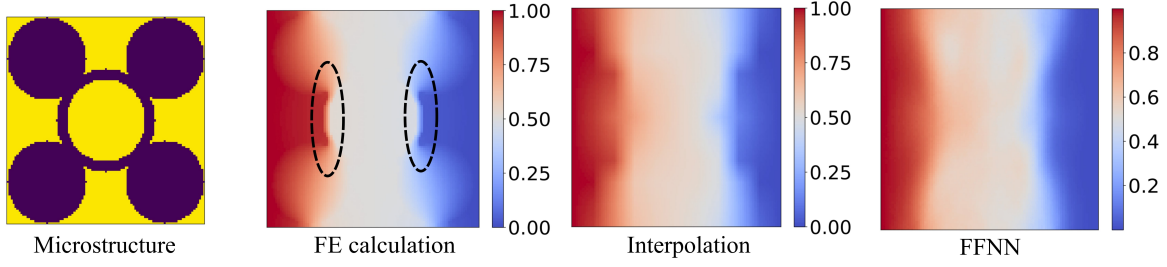


Fig. 17: A comparative analysis of interpolation and FFNN upscaling methods applied to a two-dimensional microstructure model. The image labeled FE calculation shows the result of high-fidelity finite element analysis, in which sharp interfaces within the microstructure are highlighted, as indicated by the dashed ellipses. The Interpolation and FFNN images show the upscaled results of the low-fidelity solution of the microstructure under consideration. As can be seen, compared to the high-fidelity FE calculation, the remarkable discrepancies in detail preservation, especially at sharp interfaces, show the limitations of the interpolation and FFNN approaches.

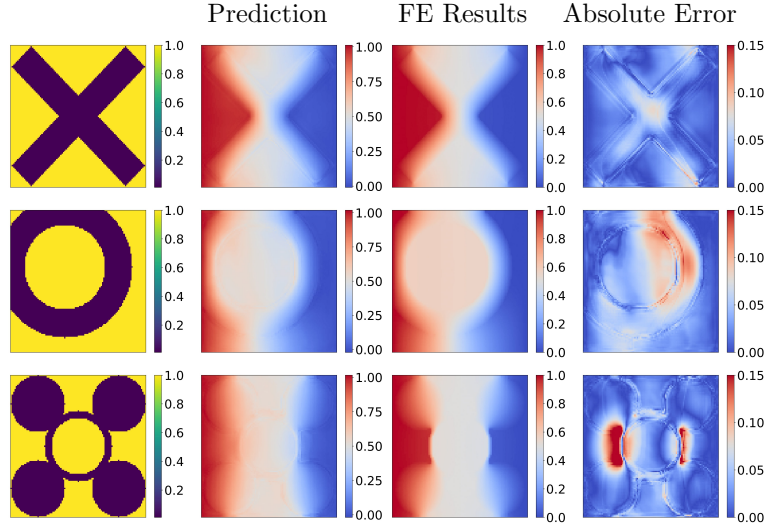


Fig. 18: Prediction results of the thermal diffusion using the MEA-Type 1 architecture, which incorporates twelve convolutional layers in the encoder section are demonstrated in three test cases.

The lack of convolutional or pooling operations further reduces FFNNs' capabilities in identifying and reconstructing significant high-level features necessary for high-resolution output.

In contrast, architectures such as MEA, which contain convolutional layers similar to the U-Net model, are characterized not only by the recognition of spatial hierarchies and dependencies but also by the reinforcement of the model through the concatenation of microstructures of various resolutions. The convolutional layers in MEA can use special filters to recognize edges and textures, regardless of their position in the input. In addition, the concatenation of microstructures helps to preserve detailed information, preventing the loss of high-resolution detail that typically occurs during downsampling.

Fig. 18 showcases the predictive capabilities of the MEA architecture in comparison to the FE solutions for the heat conductivity problem. The error column in Fig. 18 reveals surprisingly effective performance by the network on severely OOD test cases. This accuracy is evident in the absolute error color map, highlighting the robustness and adaptability of the MEA architecture in accurately modeling complex patterns even when confronted with data that significantly deviates from the distribution observed during training. This capability originates from the use of convolutional layers, which not only recognize spatial data but also enhance the model by concatenating microstructures throughout the decoder segment. These layers use specialized filters to detect edges and textures regardless of their position in the input, while the concatenation process enhances the decoder segment to preserve highly detailed information during upsampling.

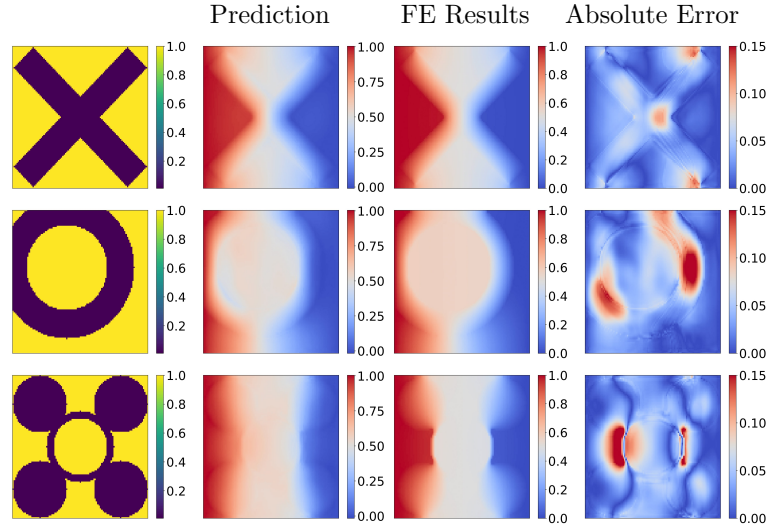


Fig. 19: Prediction results of the thermal diffusion using the MEA-Type 2 architecture, which incorporates just one convolutional layer in the encoder section are demonstrated in three test cases.

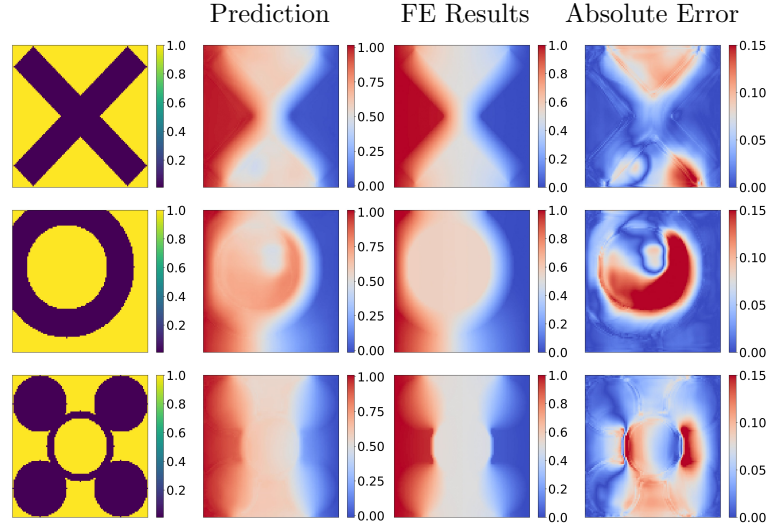


Fig. 20: Prediction results of the thermal diffusion using standard U-Net architecture section are demonstrated on three test cases.

In an alternative version of MEA-Type 1, only one convolutional layer is used in the encoder segment to increase the number of channels from 1 to 128. After training this MEA architecture (MEA-Type 2), its performance was evaluated using the same test cases that were used for the MEA-Type 1 model with twelve convolutional layers. The results, shown in Fig. 19, show mixed results. In some cases, MEA-Type 2 shows acceptable prediction accuracy, but in other cases its performance deteriorates significantly, especially in direct comparison with MEA-Type 1. This discrepancy in accuracy is most evident when comparing error maps, where the model with twelve convolutional layers in the encoder consistently performs better than the single-layer alternative.

The observed superiority of the MEA network with multiple convolutional layers in the encoder segment over MEA architecture with a single encoder layer can be attributed to an understanding of how deep learning models process and learn from data. Deep convolutional networks learn hierarchically structured features. Initial convolutional layers capture basic aspects like edges and textures, while deeper layers synthesize these into more complex patterns. A single convolutional layer, regardless of its channel depth, lacks this hierarchical structure, potentially limiting its capacity to model complex dependencies and patterns within the data. Besides, the addition of layers increases a network's capacity, allowing it to learn a richer set of features. More convolutional operations in the encoder segment

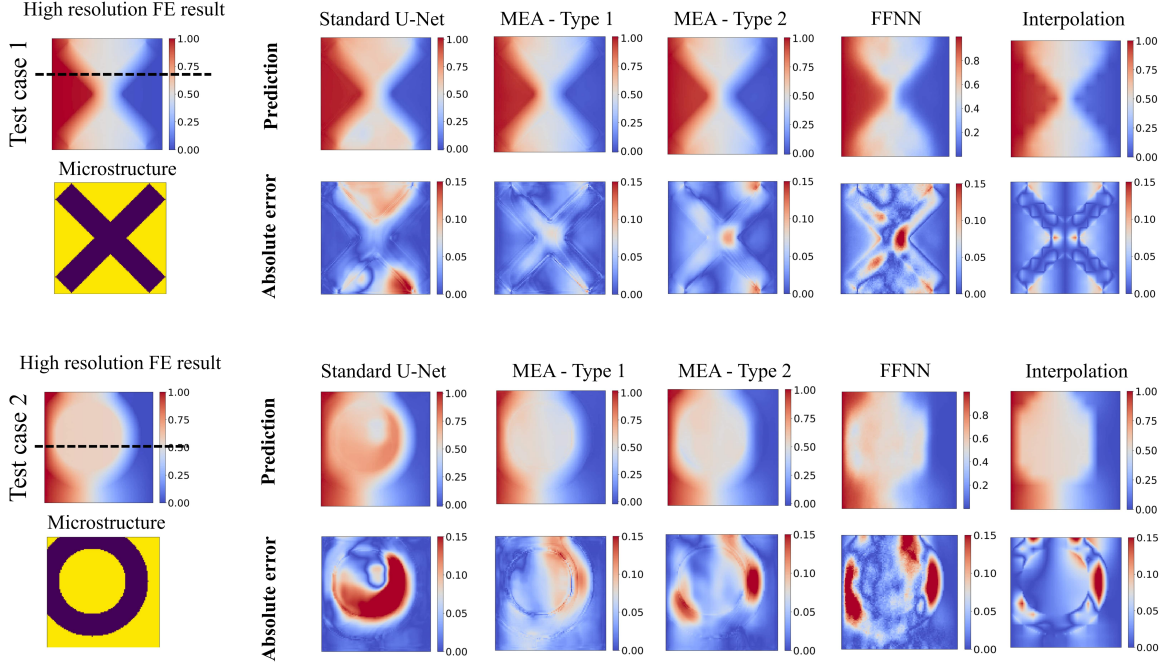


Fig. 21: Comparison of the prediction results of the architectures on a single test case. MEA-Type 1 represents an MEA architecture with twelve convolutional layers in the encoder section, while MEA-Type 2 features a model with just one convolutional layer in the encoder segment.

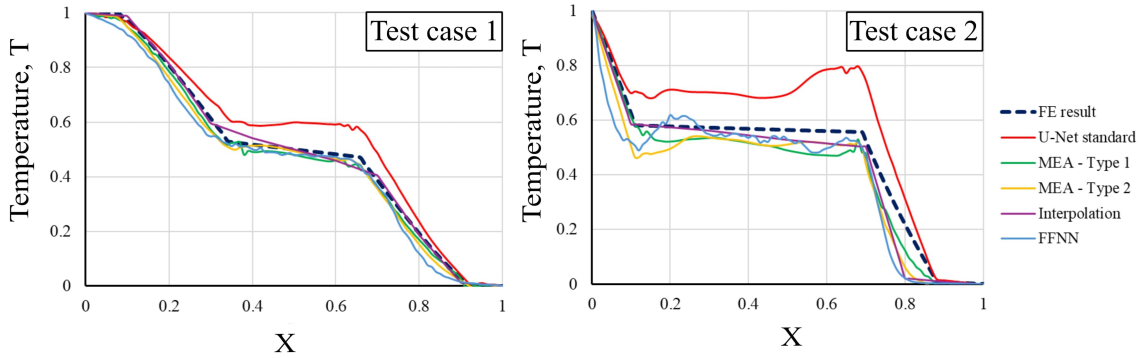


Fig. 22: Two-dimensional temperature plot on the cross-section of test cases shown in Fig. 21. The position of the cross-section is indicated by a black dashed line in Fig. 21. The MEA-Type 1 architecture demonstrates accurate predictions when compared to the FE simulation results.

enable the network to capture a broader and more intricate understanding of the data. Hence, this architectural depth enhances the network’s robustness to variations and noise in input data, especially critical for accurately predicting OOD test cases.

The predictive performance of the standard U-Net, as implemented in this study, is showcased in Fig. 20. This model processes high-resolution microstructures as inputs to produce corresponding high-fidelity solutions as outputs. Examination of the error maps reveals that the standard U-Net outperforms the MEA architecture in certain test cases (e.g., row 2 in Fig. A1 and Fig. A3), suggesting relatively better prediction accuracy under specific conditions. However, in a wider range of OOD test scenarios, the standard U-Net provides unsatisfactory results in some test cases compared to the MEA models (see Fig.18 and Fig.20). These cases show that the standard U-Net does not have the same degree of generalizability as the MEA models with the same amount of training data.

Figure 21 presents a comprehensive comparison of the predictions made by various architectures developed in this study, across two test cases. It details the spatial distribution and magnitude of errors

for each architecture, highlighting the superior performance of the MEA architectures, particularly MEA-Type1. This advantage is evident in the spatial 2D plot of absolute error values, further corroborated by the statistical analysis in Table 2. This table analyzes the prediction performance across all depicted test cases (see Figure 6), starting with the calculation of the mean error for each model on the test cases, followed by computing the average of these mean errors for each model across the six test cases.

Fig. 22 shows the predicted temperature values that are plotted along the black lines represented in Fig. 21. Besides, the finite element results for each test case are also indicated in Fig. 22 with dashed lines, showing that less sophisticated architectures like the interpolation model yield results close to the FE results. However, as seen in Fig. 17, the interpolation model fails to accurately capture the sharp interfaces crucial in fields such as fluid mechanics simulations, rendering it a less suitable choice.

Conversely, the standard U-network model, although capable of identifying sharp interfaces, did not provide promising results in *Test case 2*, as shown in Figure 22. However, according to Table 2, the overall performance of the model is comparable to that of MEA-Type 2, although it requires a significantly longer evaluation time.

Furthermore, Table 3 presents the total number of parameters for each architecture as well as the evaluation time for a single test case. The test case is uniform across all models to ensure comparability. As observed in Table 3, the MEA architectures contain fewer parameters compared to the standard U-Net model, which significantly reduces evaluation time.

Table 2: Mean error of the studied architectures on test cases shown in Fig. 6. The mean error is calculated by summing the error values for each test case and then dividing by the total number of elements of the corresponding test case.

Test cases	Standard U-Net	MEA - Type 1	MEA - Type 2	FFNN	Interpolation
Test 1	0.0377	0.0226	0.0289	0.0375	0.023
Test 2	0.011	0.0193	0.0203	0.0239	0.023
Test 3	0.0056	0.0279	0.0289	0.0313	0.0178
Test 4	0.0528	0.0439	0.0431	0.04638	0.0316
Test 5	0.0301	0.0220	0.0262	0.03	0.0175
Test 6	0.0535	0.0302	0.0356	0.0449	0.0465
Average	0.0318	0.0277	0.0305	0.0357	0.0267

Table 3: Comparison of different neural network architectures for obtaining the high-fidelity solution map of a single test case. **The FE computational cost on the high resolution parametric space of the same test case is equal to 3.619 seconds.**

	Standard U-Net	MEA - Type 1	MEA - Type 2	FFNN	Interpolation
Number of parameters	1,509,145	1,226,918	1,082,102	56,143,201	-
Evaluation time (s)	3.298	0.013	0.006	0.0419	0.000995

5.3 Training data impact: Standard U-Net vs. MEA approach

As shown in Fig. 23, we analyze the impact of the training data set on two best models in this study, the standard U-Net and the MEA approach. The analysis shows that the MEA architecture clearly outperforms the standard U-Net, especially with limited training data. However, as the amount of training data increases, the performance of the standard U-net converges relatively quickly to that of the MEA model. Therefore, it can be claimed that the MEA model is computationally advantageous when the data available for training is limited. Moreover, the integration possibility of MEA model with other methods such as FOL and FEM to obtain the low-fidelity solution facilitates the development of a multi-fidelity approach, an advantage that the standard U-Net is missing. Figure 24 shows the prediction results of both models trained with only 1,000 data points, highlighting the superior performance of the MEA model compared to the standard U-Net under the conditions of a small dataset. A more comprehensive comparison of the models is shown in Figure A6 in Appendix A. It can be argued

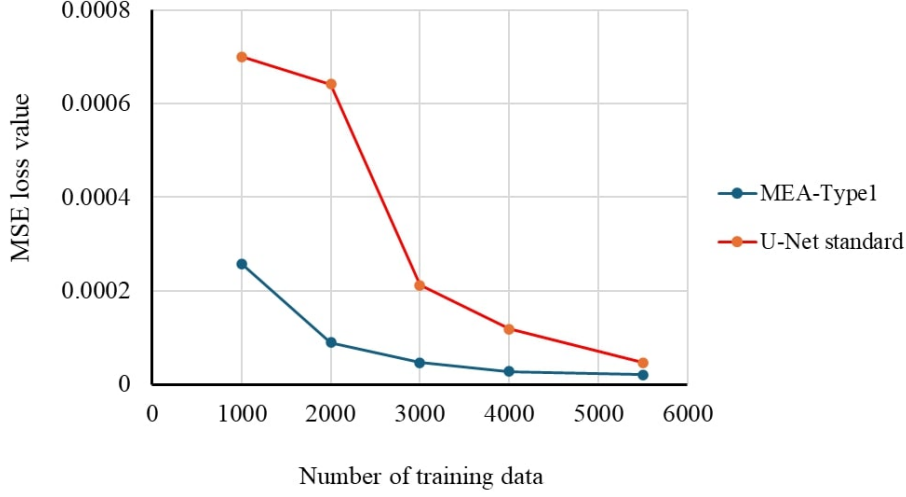


Fig. 23: Effect of the amount of training data on the validation loss of standard U-Net and MEA architecture.

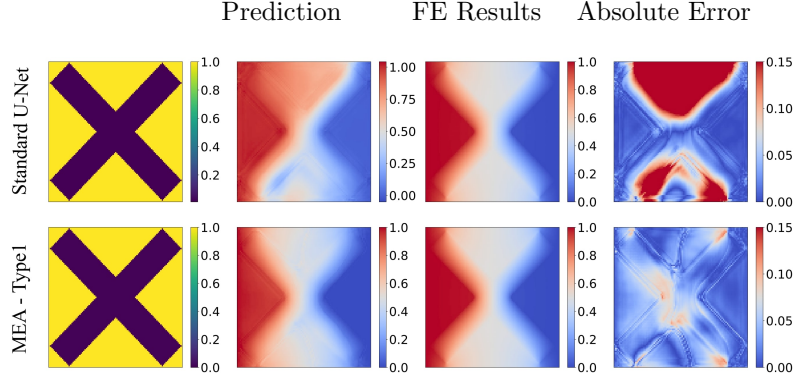


Fig. 24: Prediction results of the standard U-Net and the MEA-Type 1 model, each trained with only 1000 data points on a single test case. As the absolute error maps show, the MEA model outperforms the standard U-Net, especially when the number of available data sets is limited. A more comprehensive comparison of the models is shown in Fig. A6 across three OOD test cases.

that the combination of a smaller number of parameters, increased computational speed and improved generalizability due to the concatenations features highlights the advantages of the MEA over the conventional U-Net architecture where the amount of data available for training is limited. Considering the desired accuracy level and the available resources, the MEA-Type 2 model can also be preferred due to its significantly shorter evaluation time compared to MEA type 1 (see Table 3). However, it should be noted that this efficiency is accompanied by a slight increase in errors.

6 Conclusion

In this paper, we introduce the Microstructure-Embedded Autoencoder (MEA), a novel deep learning method developed to upscale low-fidelity solutions into high-fidelity ones. The MEA architecture not only effectively upscales low-resolution solutions to high-resolution ones while preserving critical details, especially discontinuities such as sharp interfaces in the solution field, but also achieves this with significantly lower training data requirements. This reduction in required training data is a major advantage and makes MEA model an efficient and practical solution in this area.

We argue that this method can be effectively integrated into various numerical solvers, including recent developments such as physics-informed neural operators. This capability allows users to perform computations and train neural operators on a reduced space with fewer inputs. Interestingly, we

achieved higher accuracy with this approach compared to some existing methods, such as interpolation and the standard U-Net architecture, despite using low amount of training data. The main reason for this improvement is the integration of physical information from the microstructure topology and the inclusion of a novel physically informed operator learning algorithm, both of which are not present in the classical U-Net architecture.

In this work, the proposed method is applied to the thermal diffusion problem within a heterogeneous material and involves three main steps. First, it condenses a high-resolution parametric space representing the spatial distribution of heat conductivity coefficients for two-phase composite materials into low-resolution grids. Second, it solves the boundary value problem on the coarsest grid using a pre-trained physics-informed finite operator model. Finally, it enhances the low-fidelity solution through the MEA architecture to achieve a high-fidelity output. Our findings demonstrate that MEA significantly reduces computational time, costing 280 times less than the FEM method for similar tasks. This reduction in computational cost represents its efficiency and suitability for problems in computational mechanics where both speed and accuracy are crucial. For a summary and comparison of the investigated methods in this paper see also Table 4.

In future developments, it is worthwhile to extend the methodology to a 3D setting where the advantages and acceleration of the method become even more evident. Additionally, exploring the application of the method to other physical problems, such as mechanical deformation or transient response of structures over time, would be beneficial. Coupling the method with other numerical schemes (e.g. FEM, FFT, and FDM), including alternative neural operators (e.g. DeepOnet and FNO), is also a straightforward and valuable avenue for further investigations.

Table 4: Comparison of different methods to obtain a high-fidelity solution map from a given high-fidelity microstructure map. MEA-Type 1 represents an MEA architecture with twelve convolutional layers in the encoder section, while MEA-Type 2 features a model with just one convolutional layer in the encoder segment.

Method	Interpolation approach	Feed Forward Neural Network	MEA, Type 1 (this work)	MEA, Type 2 (this work)	Classical U-Net
Computational / training cost	Extremely cheap	Moderate	Cheap	Cheap	Rather high
Implementation complexity	Very easy	Easy	Moderate	Moderate	Moderate
Possible disadvantages	Poor accuracy for sharp solution	Poor accuracy for sharp solution	Network complexity	Network complexity	Network complexity / training cost
Main advantage	Easy implementation	Easy network architecture	High accuracy even for sharp solution using moderate amount of data	High accuracy even for sharp solution using moderate amount of data	Accuracy (upon feeding enough data)

Acknowledgments

The authors are grateful to the Zentrum für Digitalisierungs- und Technologieforschung der Bundeswehr (dtec.bw) for their financial support. The author Shahed Rezaei would like to thank the Deutsche Forschungsgemeinschaft (DFG) for the funding support provided to develop the present work in the project Cluster of Excellence “Internet of Production” (project: 390621612).

Declarations

The authors declare no conflict of interest

Data availability

All the source codes and datasets to reproduce results in this study will be openly available on GitHub at <https://github.com/RasoulNajafi/Microstructure-Embedded-Autoencoder-MEA-/tree/main>.

References

- [1] Peherstorfer, B., Willcox, K., Gunzburger, M.: Survey of multifidelity methods in uncertainty propagation, inference, and optimization. *Siam Review* **60**(3), 550–591 (2018)
- [2] Fernández-Godino, M.G.: Review of multi-fidelity models. arXiv preprint arXiv:1609.07196 (2016)
- [3] Le Gratiet, L., Garnier, J.: Recursive co-kriging model for design of computer experiments with multiple levels of fidelity. *International Journal for Uncertainty Quantification* **4**(5) (2014)
- [4] Perdikaris, P., Venturi, D., Royset, J.O., Karniadakis, G.E.: Multi-fidelity modelling via recursive co-kriging and gaussian–markov random fields. *Proceedings of the Royal Society A: Mathematical, Physical and Engineering Sciences* **471**(2179), 20150018 (2015)
- [5] Giles, M.B.: Multilevel monte carlo path simulation. *Operations research* **56**(3), 607–617 (2008)
- [6] Bierig, C., Chernov, A.: Approximation of probability density functions by the multilevel monte carlo maximum entropy method. *Journal of Computational Physics* **314**, 661–681 (2016)
- [7] Minisci, E., Vasile, M.: Robust design of a re-entry unmanned space vehicle by multi-fidelity evolution control. In: *Proceedings of the 13th Annual Conference on Genetic and Evolutionary Computation*, pp. 689–696 (2011)
- [8] He, L., Qian, W., Zhao, T., Wang, Q.: Multi-fidelity aerodynamic data fusion with a deep neural network modeling method. *Entropy* **22**(9), 1022 (2020)
- [9] Zhang, X., Xie, F., Ji, T., Zhu, Z., Zheng, Y.: Multi-fidelity deep neural network surrogate model for aerodynamic shape optimization. *Computer Methods in Applied Mechanics and Engineering* **373**, 113485 (2021)
- [10] Chen, J., Gao, Y., Liu, Y.: Multi-fidelity data aggregation using convolutional neural networks. *Computer methods in applied mechanics and engineering* **391**, 114490 (2022)
- [11] Song, D.H., Tartakovsky, D.M.: Transfer learning on multifidelity data. *Journal of Machine Learning for Modeling and Computing* **3**(1) (2022)
- [12] Meng, X., Karniadakis, G.E.: A composite neural network that learns from multi-fidelity data: Application to function approximation and inverse pde problems. *Journal of Computational Physics* **401**, 109020 (2020)
- [13] Chakraborty, S.: Transfer learning based multi-fidelity physics informed deep neural network. *Journal of Computational Physics* **426**, 109942 (2021)
- [14] Aliakbari, M., Mahmoudi, M., Vadasz, P., Arzani, A.: Predicting high-fidelity multiphysics data from low-fidelity fluid flow and transport solvers using physics-informed neural networks. *International Journal of Heat and Fluid Flow* **96**, 109002 (2022)
- [15] Zhang, Y., Gong, Z., Zhou, W., Zhao, X., Zheng, X., Yao, W.: Multi-fidelity surrogate modeling for temperature field prediction using deep convolution neural network. *Engineering Applications of Artificial Intelligence* **123**, 106354 (2023)
- [16] Shu, D., Li, Z., Farimani, A.B.: A physics-informed diffusion model for high-fidelity flow field reconstruction. *Journal of Computational Physics* **478**, 111972 (2023)
- [17] Ronneberger, O., Fischer, P., Brox, T.: U-net: Convolutional networks for biomedical image segmentation. In: *Medical Image Computing and Computer-assisted intervention–MICCAI 2015: 18th International Conference, Munich, Germany, October 5–9, 2015, Proceedings, Part III* 18, pp. 234–241 (2015). Springer
- [18] Gupta, A., Bhaduri, A., Graham-Brady, L.: Accelerated multiscale mechanics modeling in a deep

learning framework. *Mechanics of Materials* **184**, 104709 (2023)

- [19] Sepasdar, R., Karpatne, A., Shakiba, M.: A data-driven approach to full-field damage and failure pattern prediction in microstructure-dependent composites using deep learning. *arXiv preprint arXiv:2104.04485* (2021)
- [20] Yan, H., Xie, W., Gao, B., Yang, F., Meng, S.: A multi-task learning model for fast prediction of mechanical behavior of ud-cfrp composites under transverse tension. *Composite Structures* **324**, 117555 (2023)
- [21] Wang, S., Wang, H., Perdikaris, P.: Learning the solution operator of parametric partial differential equations with physics-informed deepnets. *Science Advances* **7**(40), 8605 (2021)
- [22] Li, Z., Zheng, H., Kovachki, N., Jin, D., Chen, H., Liu, B., Azizzadenesheli, K., Anandkumar, A.: Physics-informed neural operator for learning partial differential equations. *arXiv:2111.03794* (2023) [arXiv:2111.03794](https://arxiv.org/abs/2111.03794)
- [23] Faroughi, S., Darvishi, A., Rezaei, S.: On the order of derivation in the training of physics-informed neural networks: case studies for non-uniform beam structures. *Acta Mechanica* **234** (2023)
- [24] Rezaei, S., Moeineddin, A., Kaliske, M., Apel, M.: Integration of physics-informed operator learning and finite element method for parametric learning of partial differential equations. *arXiv:2401.02363* (2024) [arXiv:2401.02363](https://arxiv.org/abs/2401.02363)
- [25] Liu, X.-Y., Zhu, M., Lu, L., Sun, H., Wang, J.-X.: Multi-resolution partial differential equations preserved learning framework for spatiotemporal dynamics. *Communications Physics* (2024)
- [26] Raissi, M., Perdikaris, P., Karniadakis, G.E.: Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational Physics* **378**, 686–707 (2019)
- [27] Kontolati, K., Goswami, S., Karniadakis, G.E., Shields, M.D.: Learning in latent spaces improves the predictive accuracy of deep neural operators (2023)
- [28] Guo, H., Zhuang, X., Chen, P., Alajlan, N., Rabczuk, T.: Analysis of three-dimensional potential problems in non-homogeneous media with physics-informed deep collocation method using material transfer learning and sensitivity analysis. *Engineering with Computers* (2022)
- [29] Rezaei, S., Asheri, A., Xu, B.-X.: A consistent framework for chemo-mechanical cohesive fracture and its application in solid-state batteries. *Journal of the Mechanics and Physics of Solids* **157**, 104612 (2021)
- [30] Goswami, S., Bora, A., Yu, Y., Karniadakis, G.E.: In: Rabczuk, T., Bathe, K.-J. (eds.) *Physics-Informed Deep Neural Operator Networks*, pp. 219–254. Springer, Cham (2023)
- [31] Rezaei, S., Harandi, A., Moeineddin, A., Xu, B.-X., Reese, S.: A mixed formulation for physics-informed neural networks as a potential solver for engineering problems in heterogeneous domains: Comparison with finite element method. *Computer Methods in Applied Mechanics and Engineering* **401**, 115616 (2022)
- [32] Bathe, K.-J.: *Finite Element Procedures*, 1st edn. Prentice Hall, Englewood Cliffs, NJ (1996)
- [33] Hughes, T.J.R.: *The Finite Element Method: Linear Static and Dynamic Finite Element Analysis*, 1st edn. Prentice Hall, Englewood Cliffs, NJ (1987)
- [34] Bradbury, J., Frostig, R., Hawkins, P., Johnson, M.J., Leary, C., Maclaurin, D., Necula, G., Paszke, A., VanderPlas, J., Wanderman-Milne, S., Zhang, Q.: JAX: Composable Transformations of Python+NumPy programs. [http://github.com/google/jax](https://github.com/google/jax)
- [35] Haghighat, E., Juanes, R.: *Sciann: A keras/tensorflow wrapper for scientific computations and*

physics-informed deep learning using artificial neural networks. *Computer Methods in Applied Mechanics and Engineering* **373**, 113552 (2021)

[36] Schmidhuber, J.: Deep learning. *Scholarpedia* **10**(11), 32832 (2015)

Appendix A Additional Results

In this section, we present the prediction results of the implemented MEA architectures and the standard U-Net across all six test cases. In addition, the predictive capabilities of the investigated models in capturing the gradient terms in heat flux problems are analyzed, and the prediction results of the heat flux are compared with the FE simulation results.

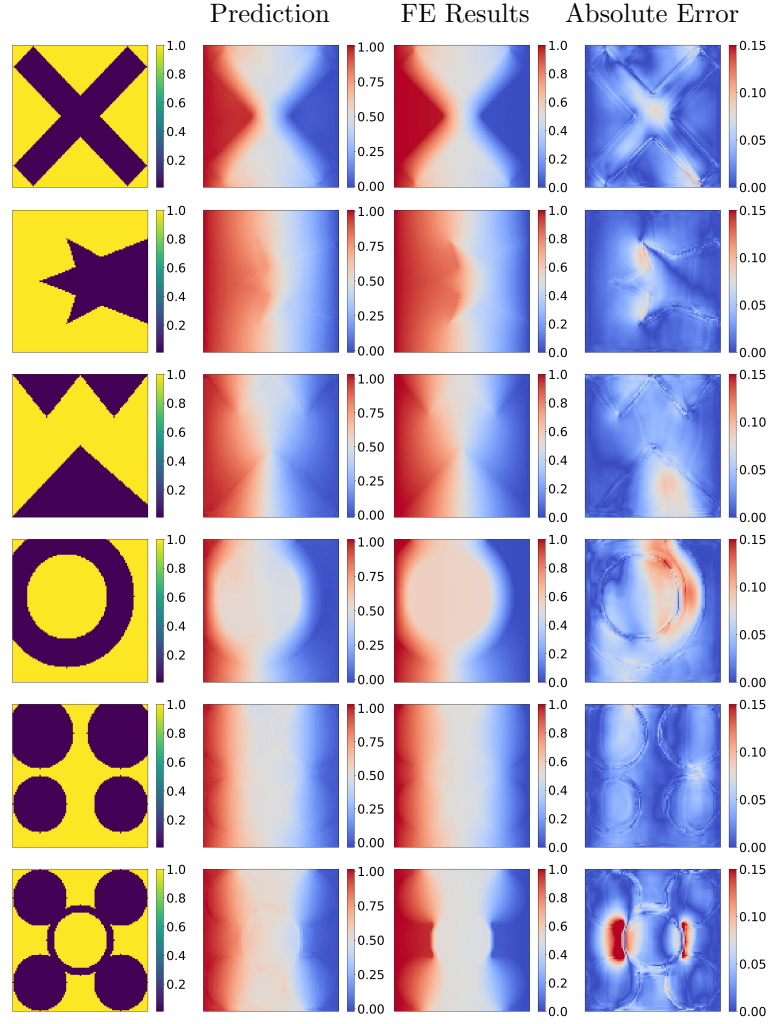


Fig. A1: Prediction results of the MEA-Type1 architecture, which incorporates twelve convolutional layers in the encoder section to convert low-fidelity solutions of the steady state heat transfer problem into high-fidelity ones, are demonstrated on six test cases.

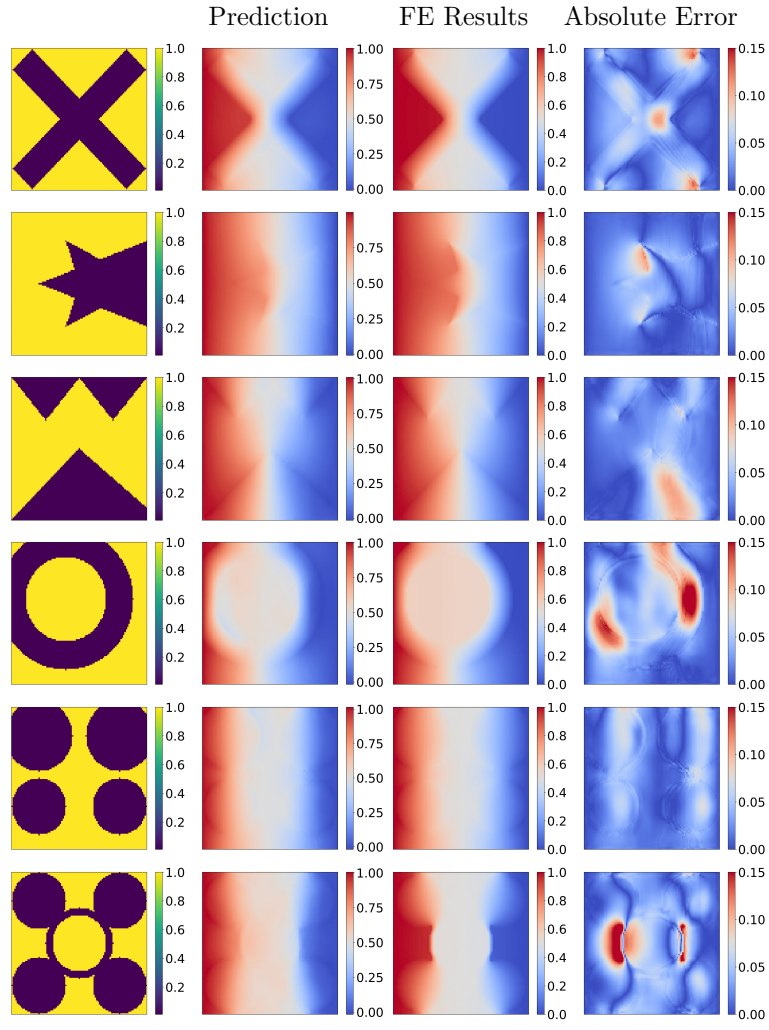


Fig. A2: Prediction results of enhanced MEA-Type2 architecture, which incorporates just one convolutional layer in the encoder section are demonstrated on six test cases.

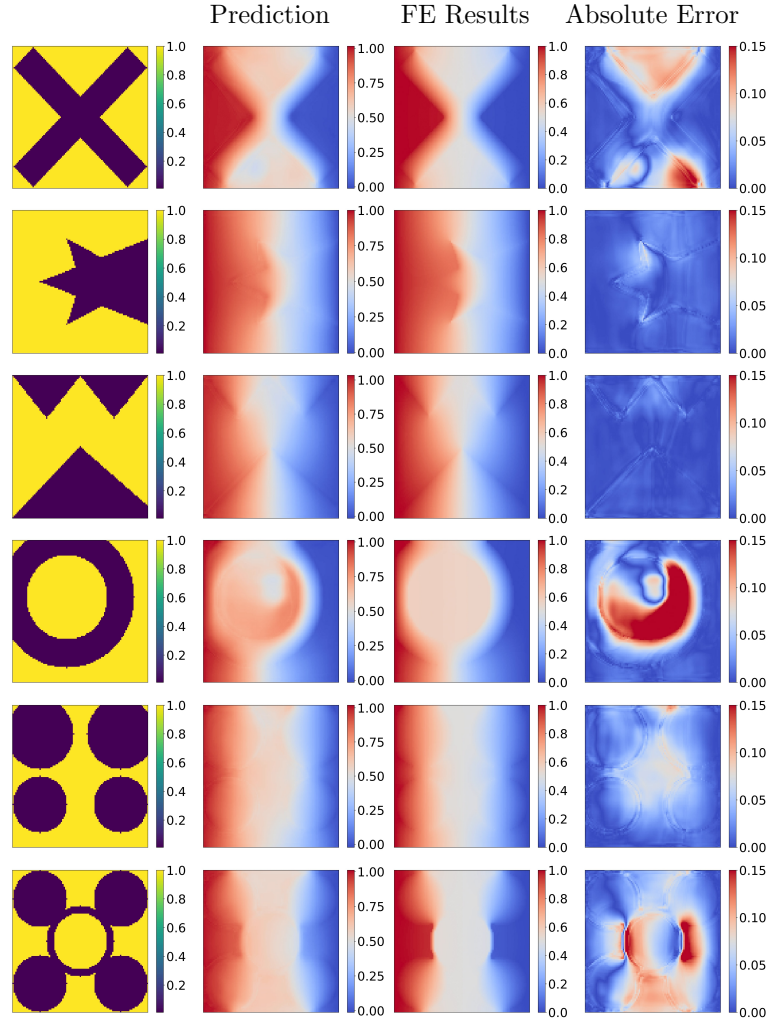


Fig. A3: Prediction results of the standard U-Net architecture section are demonstrated on six test cases.

To further evaluate the performance of the MEA architectures and compare them with the standard U-Net, we conducted an analysis focusing on their ability to capture gradient terms after predicting spatial temperature distributions. This analysis involves calculating the gradient of the predicted temperature first and then multiplying it with the spatial heat conductivity coefficients to determine the heat flux. By comparing Fig. A4 which illustrates the results from the MEA-Type1 architectures with the results from the standard U-Net shown in Figure A5, it is clear that the result of both methods deviates from the finite element calculations. However, the initial focus of this study is on predicting the primary variable, which in this case is the temperature field. Future work will need to expand on this to improve predictions for the spatial derivatives of the primary variable, as demonstrated in [31].

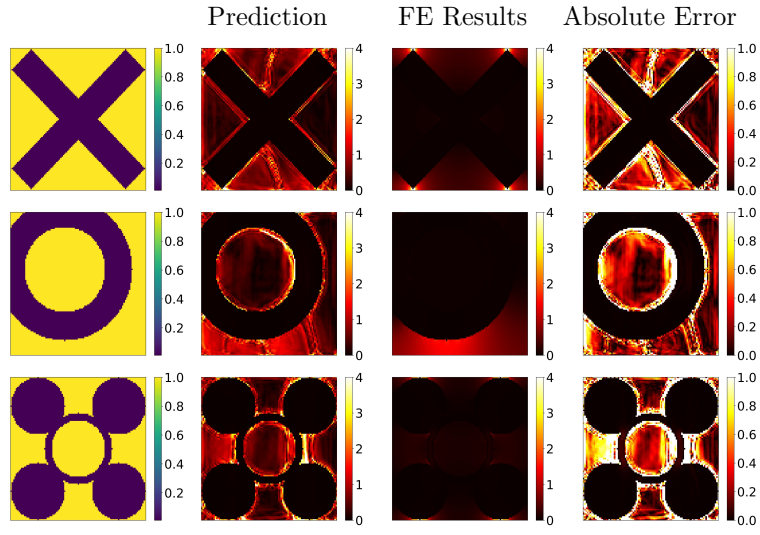


Fig. A4: Prediction results of heat flux using the MEA-Type1 model, featuring a twelve convolutional operation in the encoder section.

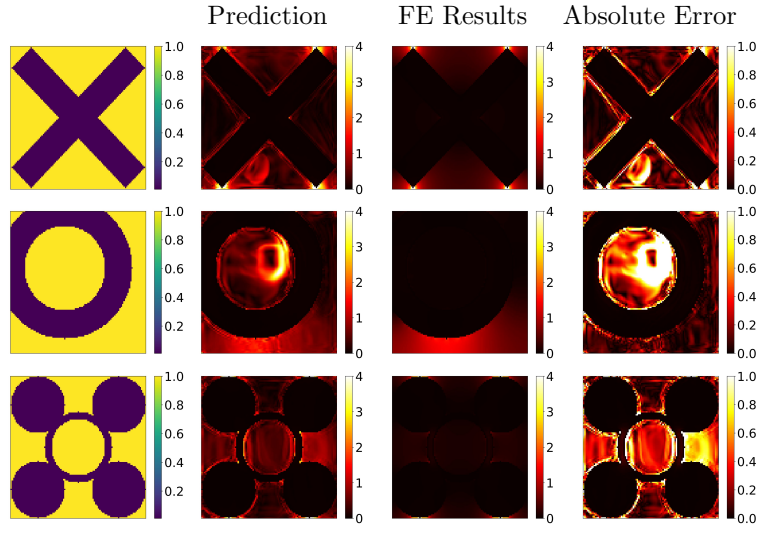


Fig. A5: Prediction results of heat flux using the standard U-Net architecture across three test scenarios.

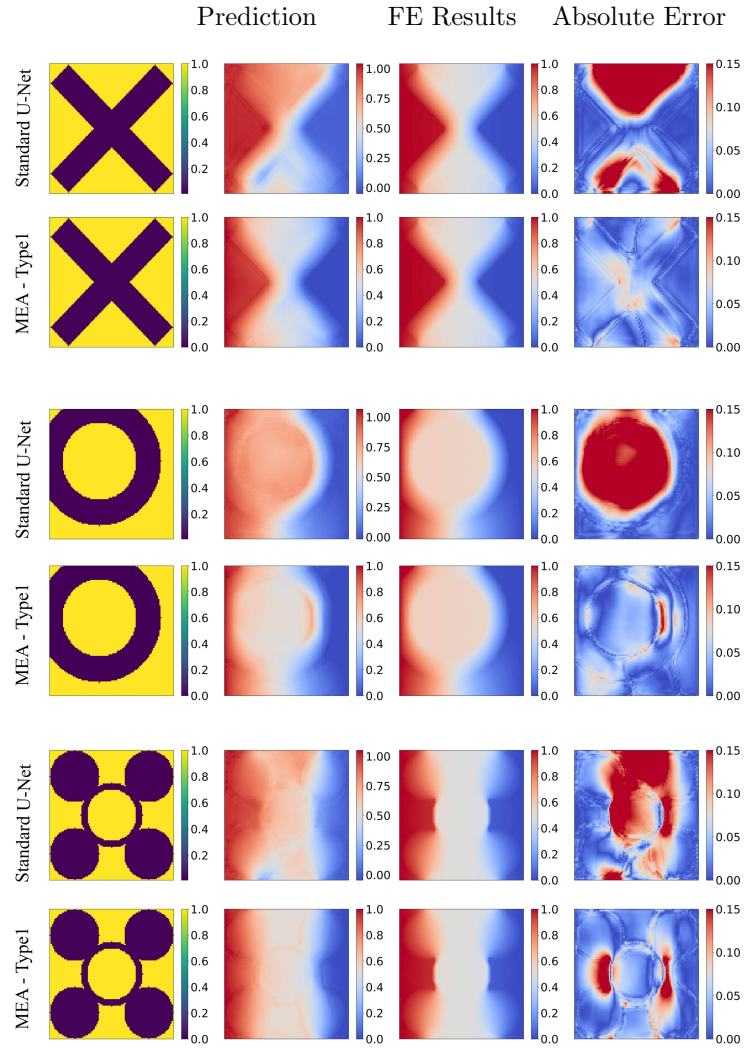


Fig. A6: Prediction results of the standard U-Net and MEA-Type 1 model for three test cases. Both models are trained with 1,000 data points. As can be observed, the MEA architecture shows better performance in terms of accuracy compared to the standard U-Net when the amount of training data points is considerably small.