

Supplementary Material A. PHITS Input

Supplementary Material A.1. Material Composition

The compositions of air and concrete employed in this study are presented in Tables [A.1](#) and [A.2](#). In PHITS, materials can be defined using either atomic ratio or mass ratio. The identification of nuclei can be expressed in terms of the number of protons Z and atomic number A using the following formula:

$$\text{Name of Nuclide} = Z \times 1000 + A. \quad (\text{A.1})$$

Nuclide	Composition ¹
7014	78.08
8016	20.95
Density (g/cm ³)	1.21×10^{-3}

¹ atomic ratio of air

Table A.1. Composition of air

Nuclide	Composition ²
1001	0.011
11023	0.043
13027	0.145
16032	0.002
20040	0.197
26056	0.085
8016	1.05
12024	0.034
14028	0.592
19039	0.035
22048	0.006
Density (g/cm ³)	2.2

² mass ratio of concrete

Table A.2. Composition of concrete

Supplementary Material B. Overall performance of the DeepONet models

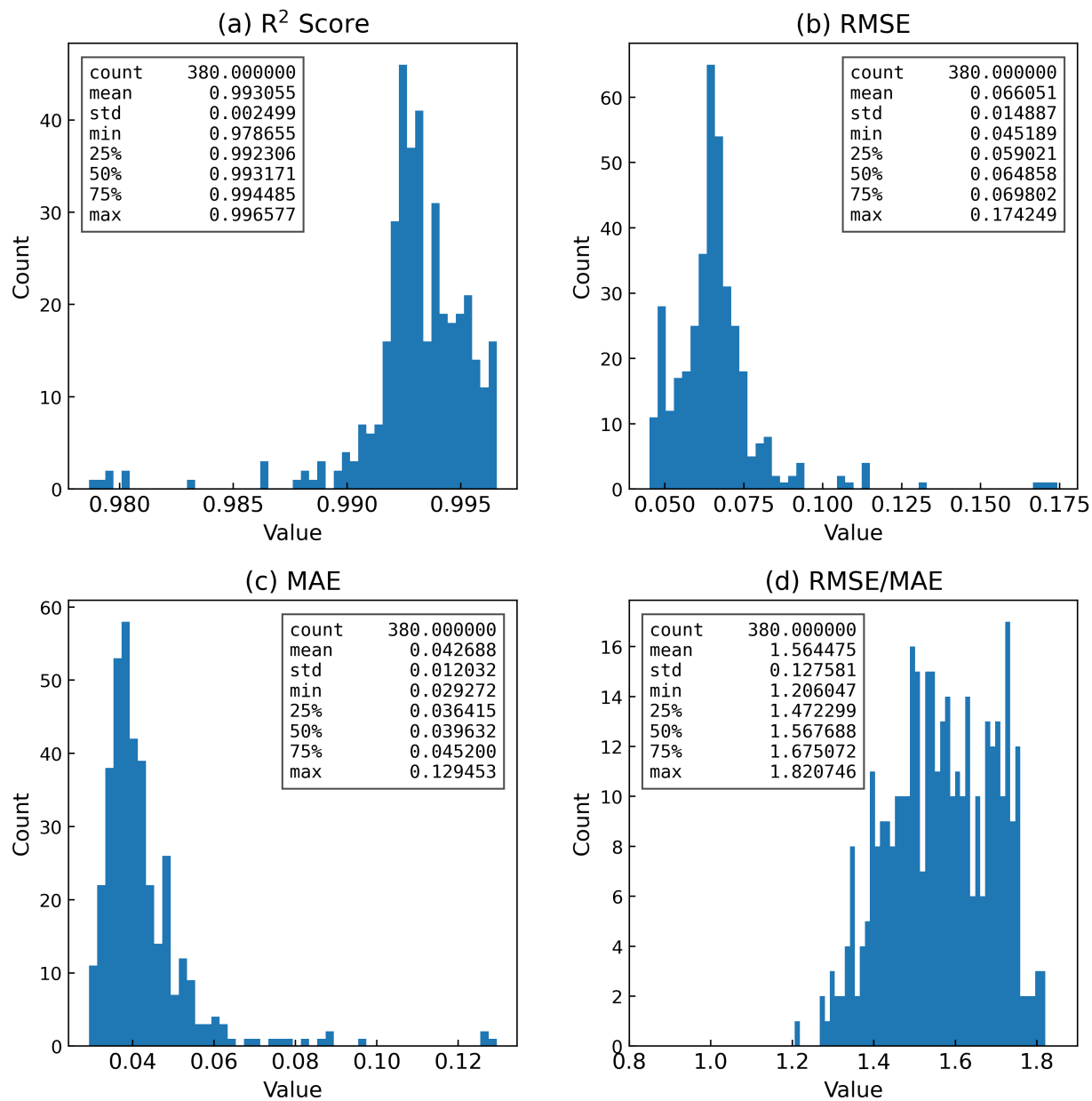


Fig. B.1. DeepONet model trained with the training data Set1. Each pane represents; (a) R^2 score, (b) RMSE, (c) MAE, and (d) RMSE/MAE.

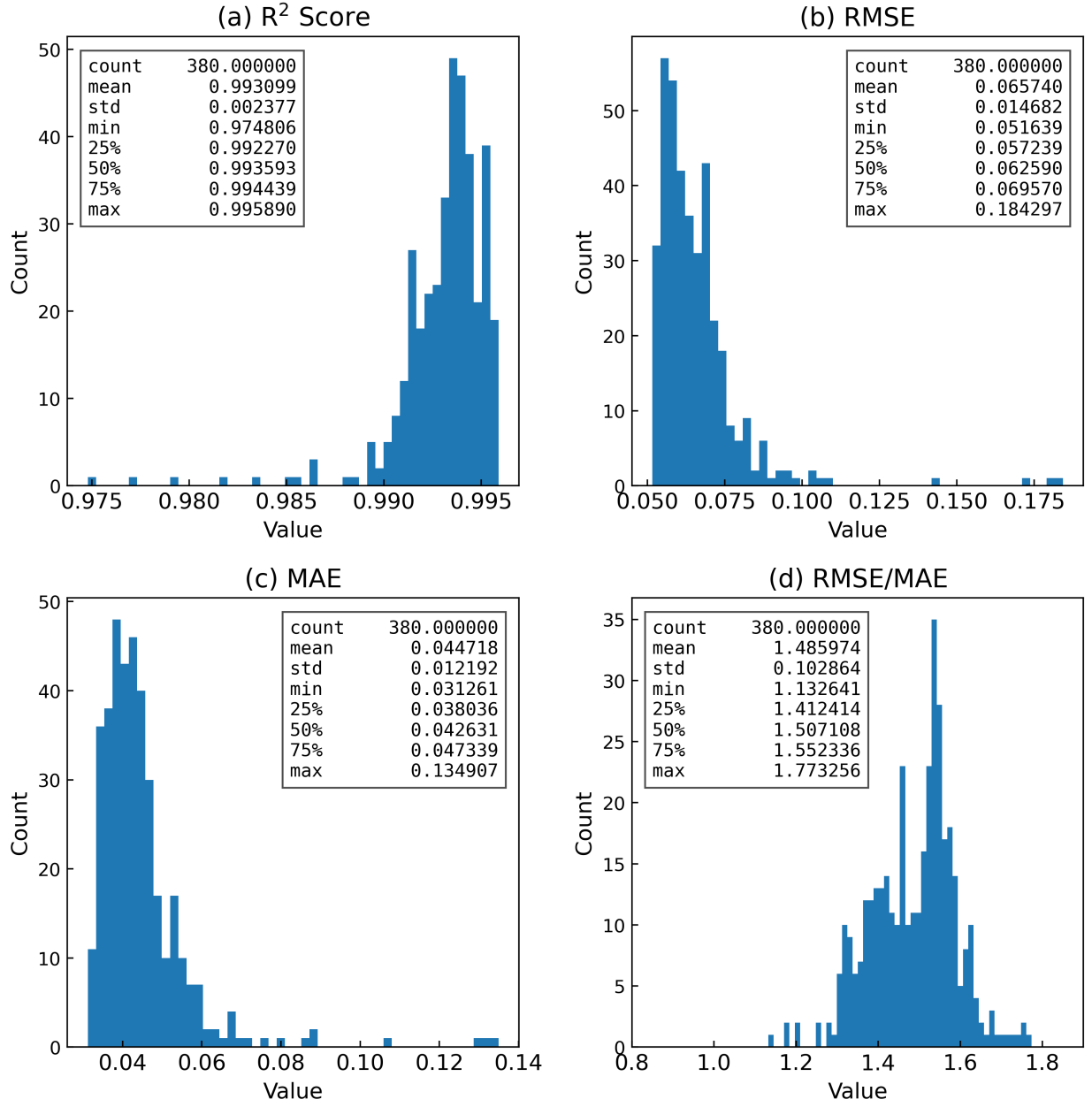


Fig. B.2. DeepONet model trained with the training data Set2. Each pane represents; (a) R^2 score, (b) RMSE, (c) MAE, and (d) RMSE/MAE.

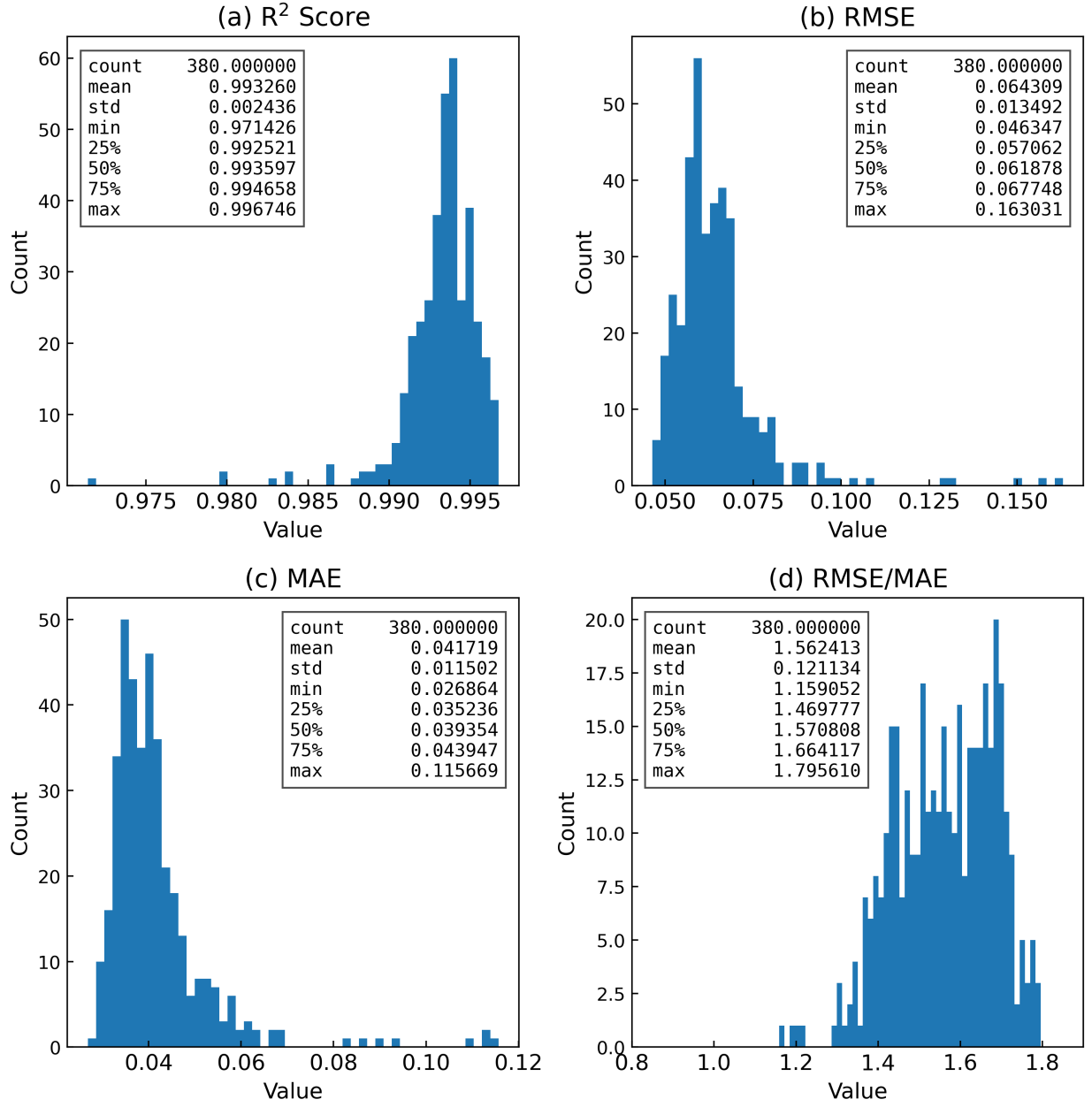


Fig. B.3. DeepONet model trained with the training data Set3. Each pane represents; (a) R^2 score, (b) RMSE, (c) MAE, and (d) RMSE/MAE.

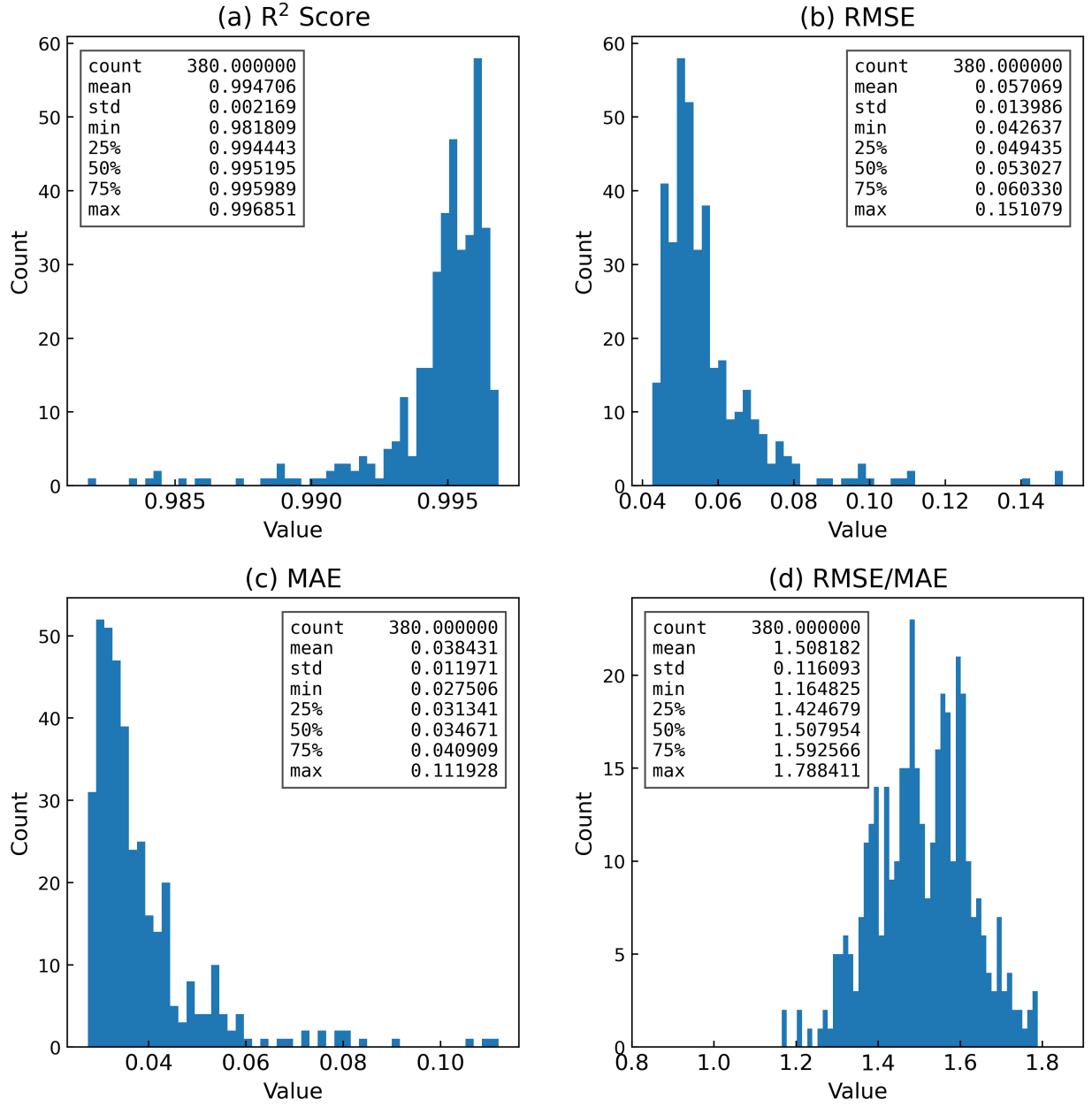


Fig. B.4. DeepONet model trained with the training data Set4. Each pane represents; (a) R^2 score, (b) RMSE, (c) MAE, and (d) RMSE/MAE.

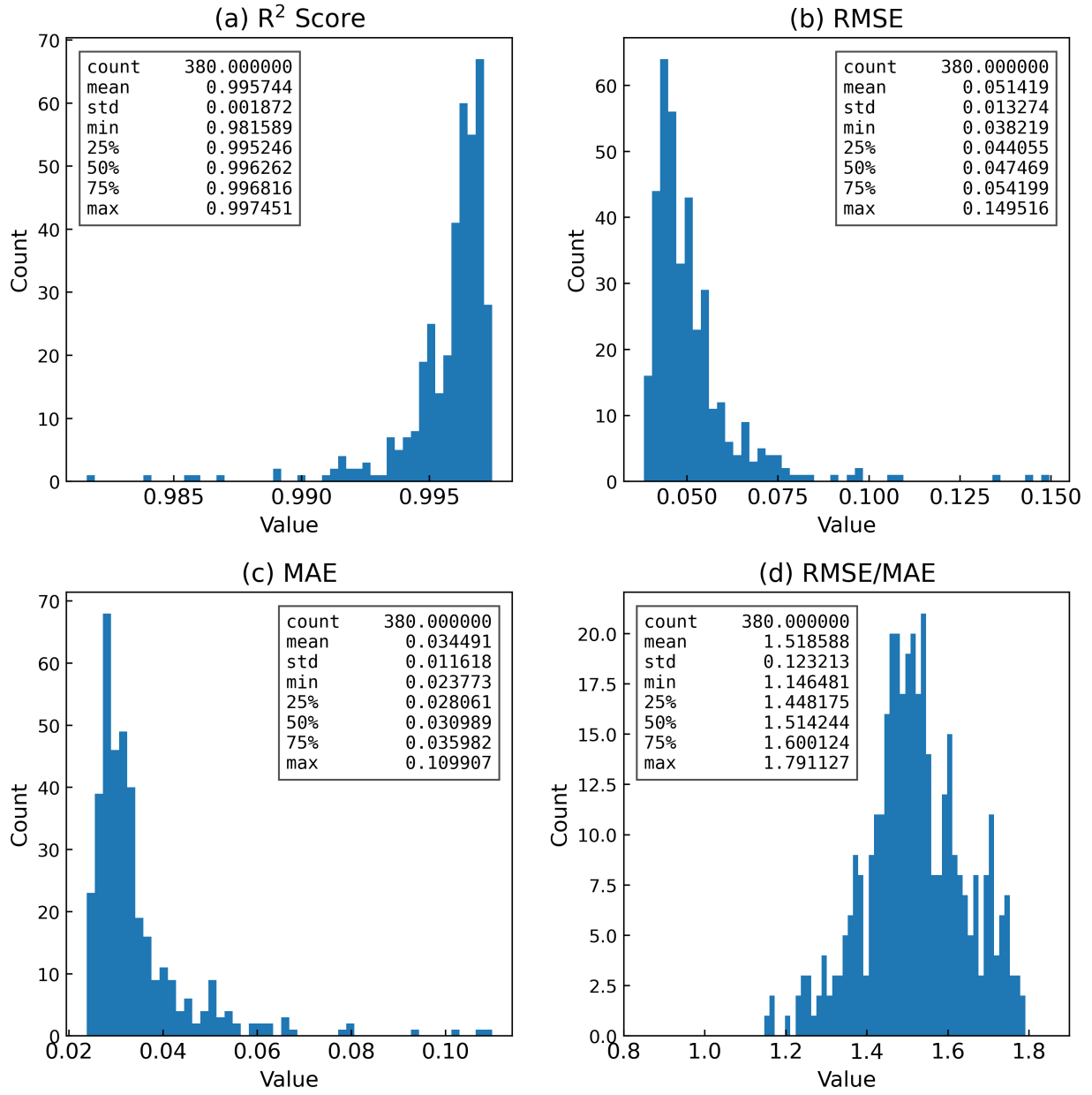


Fig. B.5. DeepONet model trained with the training data Set5. Each pane represents; (a) R^2 score, (b) RMSE, (c) MAE, and (d) RMSE/MAE.

Supplementary Material C. Network Architectures of FCN and CNN

In this study, we conducted a comprehensive comparison between DeepONet, and two conventional neural network methods: Fully-Connected Neural Networks (FCN) and Convolutional Neural Networks (CNN). All FCNs and CNNs were implemented using the PyTorch 2.0 deep learning framework, ensuring a standardized and efficient modeling process.

Supplementary Material C.1. FCN

The FCNs used in this study were constructed with a simple architecture consisting of an input layer, a hidden layer, and an output layer. The input layer had 2 neurons, the hidden layer contained 1024 neurons, and the output layer consisted of 1 neuron. ReLU activation functions were employed throughout the network. During training, the mean squared error (MSE) was used as the loss function, and the model parameters were optimized using the Adam optimization algorithm with a learning rate of 0.001. The training process utilized the mini-batch gradient descent method with a batch size of 20, iterated 1000 epochs.

```
FCN(
    (activation): ReLU()
    (linears): ModuleList(
      (0): Linear(in_features=2, out_features=1024, bias=True)
      (1): Linear(in_features=1024, out_features=1, bias=True)
    )
)
```

Supplementary Material C.2. CNN

CNNs were designed with a more complex architecture, including an input layer, a 1-dimensional convolutional layer, two linear layers, and an output layer. Dropout layers were incorporated with a dropout rate of 0.05 to prevent overfitting. The input and output layers each had 2 neurons, while the convolutional layer had 1 input channel and 1024 output channels with a kernel size of 1. Outputs from the convolutional layer were passed through linear layers, generating 2048 outputs before reaching the final output layer. Tanh activation functions were used in the network. The CNN model was also trained using MSE as the loss function, with the Adam optimization algorithm and a learning rate of 0.001. Training employed the mini-batch gradient descent method with a batch size of 50, repeated for 1000 epochs.

```
CNN(
    (activation): Tanh()
    (cnn1): Conv1d(1, 1024, kernel_size=(1,), stride=(1,))
    (flat): Flatten(start_dim=1, end_dim=-1)
    (dropout): Dropout(p=0.05, inplace=False)
    (lin1): Linear(in_features=2048, out_features=2048, bias=True)
    (lin2): Linear(in_features=2048, out_features=1, bias=True)
)
```

Supplementary Material D. Accuracy Metric for Regression Task

A regression model entails predicting a numerical output, y , based on input data X . Evaluating the accuracy of such a model involves assessing how well the predicted output, y_{pred} , matches the true output, y_{true} . This section delves into the metrics employed in this research to gauge the accuracy of the model.

Supplementary Material D.1. R -squared (R^2) Score

The R^2 score is a statistical indicator used to assess the quality of fit of a regression model. It sheds light on how well the target variable's true values and the anticipated values of the model match each other. This is how the R^2 score is determined:

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_{true,i} - y_{pred,i})^2}{\sum_{i=1}^n (y_{true,i} - \bar{y}_{true})^2} \quad (D.1)$$

where $y_{pred,i}$ represents the predicted value of the target variable for the i -th data point, $y_{true,i}$ represents the observed value of the target variable for the i -th data point, $\bar{y}_{true}$ is the mean of the true values of the target variable, and n is the number of data points.

The range of the R^2 score is between 0 and 1. A higher score denotes a better prediction accuracy, with 0 denoting that predictions are inaccurate, and 1 denoting a perfect fit. It's crucial to remember that a high score does not always indicate a causal connection or the absence of overfitting.

Supplementary Material D.2. Root-Mean-Squared Error (RMSE)

The RMSE is a statistic used to assess the predictive performance of a regression model by measuring the average size of the discrepancies between predicted and observed values. It indicates how well the projected values match the actual values. The RMSE is determined as follows:

$$RMSE = \sqrt{\frac{\sum_{i=1}^n (y_{true,i} - y_{pred,i})^2}{n}} \quad (D.2)$$

The RMSE is provided in exactly the same units as the target variable, making it simpler to read and compare to the original data scale. A smaller RMSE suggests that the results predicted by the model are more accurate, implying a better match. However, like with any statistic, RMSE should be seen in context and in conjunction with other assessment metrics in order to have a thorough picture of the model's performance.

Supplementary Material D.3. Mean Absolute Error (MAE)

Another statistic used to assess the accuracy of a regression model's predictions is mean absolute error (MAE). It measures the difference between anticipated and true values in the same way that RMSE does, but instead of squaring the differences, it takes the absolute value of the discrepancies. Because it does not give more weight to higher mistakes, MAE is more resistant to outliers. The MAE is determined as follows:

$$MAE = \frac{\sum_{i=1}^n |y_{true,i} - y_{pred,i}|}{n} \quad (D.3)$$

The MAE is expressed in the identical units as the target variable, making it straightforward to read and compare to the original data scale. A lower MAE suggests that the model's predictions are

more accurate. When compared to RMSE, which squares the errors, MAE is advantageous when we wish to punish huge mistakes less. MAE, like R2 score and RMSE, should be used in conjunction with other assessment metrics to provide a complete picture of the model's performance.

Supplementary Material D.4. Ratio of RMSE to MAE

As mentioned above, both RMSE and MAE have the same units as the original data. This section shows how these values are related. To make the equation more readable, the absolute value of the residual between the predicted and true value, e_i , is expressed as follows:

$$e_i = |y_{true,i} - y_{pred,i}|. \quad (D.4)$$

Therefore, Equations D.2 and D.3 can be modified by the followings:

$$RMSE^2 = \frac{\sum_{i=1}^n e_i^2}{n}, \quad (D.5)$$

$$MAE^2 = \frac{(\sum_{i=1}^n e_i)^2}{n^2}. \quad (D.6)$$

Let's consider the concept of variance in the context of a random variable X . Variance, denoted as $Var(X)$, measures the extent to which individual data points in a dataset deviate from the mean of that dataset. It is a statistical measure of the dispersion or spread of the data points. The variance of a random variable X is calculated using the following formula:

$$\begin{aligned} Var(X) &= E(X^2) - (E(X))^2 \\ &= \frac{\sum_{i=1}^n x_i^2}{n} - \frac{(\sum_{i=1}^n x_i)^2}{n^2} \end{aligned} \quad (D.7)$$

where x_i represents the value of the random variable X for the i -th data point. If the variable X is replaced with the residual error, the equation can be arranged as following:

$$RMSE^2 - MAE^2 = Var(e). \quad (D.8)$$

Also, the e_i that only takes values greater than or equal to 0, mean value of e is equal to MAE, the ratio of RMSE to MAE can be expressed as

$$\frac{RMSE}{MAE} = \sqrt{1 + \frac{Var(e)}{MEAN(e)^2}}. \quad (D.9)$$

Let us assume that the error follows a normal distribution with mean 0 and standard deviation σ . The distribution of the absolute value of the error ($=e_i$) is then the distribution of the absolute value of the normal distribution. The probability density function, f , is defined as follows:

$$f = \frac{2}{\sqrt{2\pi}\sigma} \exp\left(-\frac{e^2}{2\sigma^2}\right). \quad (D.10)$$

Using the probability density function, the mean and variance can be computed as:

$$\begin{aligned}
\text{MEAN}(e) &= \int_0^\infty e \cdot f de \\
&= \int_0^\infty e \cdot \frac{2}{\sqrt{2\pi}\sigma} \exp\left(-\frac{e^2}{2\sigma^2}\right) de \\
&= \sqrt{\frac{2}{\pi}}\sigma
\end{aligned} \tag{D.11}$$

$$\begin{aligned}
\text{Var}(e) &= \int_0^\infty (e - \text{MEAN}(e))^2 \cdot f de \\
&= \int_0^\infty (e - \sqrt{\frac{2}{\pi}}\sigma)^2 \cdot \frac{2}{\sqrt{2\pi}\sigma} \exp\left(-\frac{e^2}{2\sigma^2}\right) de \\
&= \left(1 - \frac{2}{\pi}\right)\sigma^2
\end{aligned} \tag{D.12}$$

Finally, the ratio of RMSE to MAE can be expressed by

$$\begin{aligned}
\frac{\text{RMSE}}{\text{MAE}} &= \sqrt{1 + \frac{\left(1 - \frac{2}{\pi}\right)\sigma^2}{\left(\sqrt{\frac{2}{\pi}}\sigma\right)^2}} \\
&= \sqrt{\frac{\pi}{2}} \approx 1.253.
\end{aligned} \tag{D.13}$$