

An inversion problem for optical spectrum data via physics-guided machine learning.

Hwiwoo Park¹, Jun H. Park^{2,†}, and Jungseek Hwang^{1,*}

¹*Department of Physics, Sungkyunkwan University,
Suwon, Gyeonggi-do 16419, Republic of Korea*

²*School of Mechanical Engineering, Sungkyunkwan University,
Suwon, Gyeonggi-do 16419, Republic of Korea*

(Dated: April 4, 2024)

Abstract

We propose the regularized recurrent inference machine (rRIM), a novel machine-learning approach to solve the challenging problem of deriving the pairing glue function from measured optical spectra. The rRIM incorporates physical principles into both training and inference and affords noise robustness, flexibility with out-of-distribution data, and reduced data requirements. It effectively obtains reliable pairing glue functions from experimental optical spectra and yields promising solutions for similar inverse problems of the Fredholm integral equation of the first kind.

Correspondence to [emails: [†]jun.park@skku.edu and ^{*}jungseek@skku.edu].

I. INTRODUCTION

Experimental and theoretical investigations on high-temperature copper-oxide (cuprate) superconductors, since their discovery over 35 years ago [1, 2], have afforded extensive results [3]. Despite these efforts, the microscopic electron-electron pairing mechanism for superconductivity remains elusive. In this regard, researchers have adopted innovative experimental techniques. Particularly, optical spectroscopy has the potential to elucidate the aforementioned pairing mechanisms because it is the only spectroscopic experimental method capable of providing quantitative physical quantities. The absolute pairing glue spectrum measured via optical spectroscopy may serve as a “smoking gun” evidence to address for this problem. The measured spectrum entails information concerning the pairing glue responsible for superconductivity. Extracting this glue function from the measured optical spectra via the decoding approach, which involves an inverse problem, contributes an essential aspect to the elucidation of high-temperature superconductivity.

The decoding-related inversion problem concerning physical systems is expressed as follows:

$$\frac{1}{\tau^{\text{op}}(\omega, T)} = \int_0^\infty d\Omega I^2 \chi(\Omega, T) K(\omega, \Omega, T), \quad (1)$$

which is referred to as the generalized Allen formula [4–7]. Here $1/\tau^{\text{op}}(\omega)$ is the optical scattering rate, $K(\omega, \Omega)$ is the kernel, and $I^2 \chi(\omega)$ is the pairing glue function, which describes interacting electrons by exchanging the force-mediating boson. Further, $\chi(\omega)$ is the boson

spectrum and I denotes the electron–boson coupling constant. The kernel is given as

$$K(\omega, \Omega, T) = \frac{\pi}{\omega} \left[2\omega \coth\left(\frac{\Omega}{2T}\right) - (\omega + \Omega) \coth\left(\frac{\omega + \Omega}{2T}\right) + (\omega - \Omega) \coth\left(\frac{\omega - \Omega}{2T}\right) \right]. \quad (2)$$

This is referred to as the Shulga kernel [5]. The goal was to infer the glue function from the optical scattering rate obtained using optical spectroscopy [7]. Eq. (1) can be expressed in a more general form as follows:

$$y(t) = \int_a^b d\tau x(\tau) k(t, \tau), \quad (3)$$

which is the Fredholm integral equation of the first kind. Here, $k(t, \tau)$ is referred to as the kernel and determined by the underlying physics of the given problem, and $x(\tau)$ and $y(t)$ are physical quantities related to each other through the integral equation. Such inverse problems occur in many areas of physics [6, 8, 9] and are known to be ill-posed [10]. The ill-posed nature arises from the instability of solutions in Eq. (3), where small changes in y can lead to significant changes in x . Consequently, obtaining a solution to the inverse problem becomes challenging, particularly when observations are corrupted by noise.

Conventional approaches for solving inverse problems, expressed in the form of Eq. (3) include singular value decomposition (SVD) [11], least squares fit [12, 13], maximum entropy method (MEM) [6, 7], and Tikhonov regularization [14]. In particular, the MEM can effectively capture the key features of x , whereas the SVD and least-squares fit approaches, respectively, yield non-physical outputs and requiring a priori assumptions on the shape or size of x . At high noise levels, the MEM struggles to determine a unique amplitude [15] or capture the peaks of x [16]. Despite its theoretical advantages in terms of convergence properties and well-defined solutions, Tikhonov regularization poses challenges with regard to the selection of appropriate regularization parameters [17–19].

In recent years, machine learning approaches have frequently been applied to inverse problems of the Fredholm integral of the first kind [16, 20, 21]. Early approaches primarily utilized supervised learning [9, 16, 21] demonstrating that machine learning results were comparable or superior to those of the conventional MEM and more robust against observation noise. Regardless of their successes, they are principally model-agnostic, resembling black box models wherein the physical model is solely used for generating training data but not

leveraged during the inference procedure. Consequently, these methods lack explainability or reliability in their outputs and require substantial training data to achieve competitive results. This typically does not pose any difficulty in conventional machine learning domains, such as image classification, speech recognition, and text generation, wherein vast datasets are readily available, and explainability and reliability may be relatively less critical. In contrast, in scientific domains, a sound theoretical foundation must be provided for the output of the model, and data acquisition must be cost-effective. This is particularly crucial, because the data collection often necessitates expensive experiments or extensive simulations.

Various approaches have been explored to incorporate physics in the machine learning process. Projection is applied to the outcomes of supervised learning; this imposes physical constraints on x such as normality, positivity, and compliance with the first and second moments [20]. In a prior study [22], the forward model (Eq. (1)) is added to the loss term for the training and serves as a regularizer. By utilizing the incorporated physics, these approaches can yield comparable results with significantly smaller training datasets as compared to conventional approaches. However, in these approaches, the application of physical constraints occurs at the last stage of the learning process. Consequently, they may not be sufficiently flexible to handle situations wherein the input significantly deviates from the training data. Adaptability to unseen data is crucial because the generated training data may not encompass all possible solutions, potentially leading to a bias in the output toward the training data.

To address this issue, we adopted a recurrent inference machine (RIM) [23]. Compared to other supervised approaches, RIM exhibits exceptional capability to incorporate physical constraints throughout both the learning and inference processes through an iterative process. This implies that physical principles are not applied solely at the last stage of the learning process but are integrated throughout the learning and inference stages; thus, RIM has remarkable flexibility for solving complex inverse problems. Additionally, the RIM framework automates many hand-tuned optimization operations, streamlining the training procedure and achieving improved results [23, 24]. The RIM framework has been applied to various types of image reconstruction ranging from tomographic projections [25] to astrophysics [26]. However, conventional RIM requires modification to achieve competitive results in case of ill-posed inverse problems such as the Fredholm integral equation of the first kind.

Accordingly, we developed the regularized RIM (rRIM). Furthermore, we demonstrated that the rRIM framework is equivalent to iterative Tikhonov regularization [27–29].

The proposed physics-guided approach significantly reduces the amount of training data required, typically by several orders of magnitude. To illustrate this aspect, we quantitatively compared the average test error losses of the rRIM with those of two widely used supervised learning approaches—a fully connected network (FCN) and convolutional neural network (CNN). Owing to the sound theoretical basis of the rRIM, it can adequately explain its outputs; this is a crucial feature in scientific applications and in contrast with purely data-driven black box models. We demonstrated the robustness of the rRIM in handling noisy data by comparing its performance with those of the FCN and CNN models which are well-known for their strong noise robustness compared to MEM, especially under high observational noise conditions [9, 16]. The rRIM exhibited much better robustness than FCN and CNN for a wide range of noise levels. The rRIM exhibits superior flexibility in handling out-of-distribution (OOD) data than do the FCN and CNN approaches. The OOD data refers to data that significantly differs in characteristics, such as shape and distribution, from any of the training data. Finally, we applied the rRIM to the experimental optical spectra of optimally and overdoped $\text{Bi}_2\text{Sr}_2\text{CaCu}_2\text{O}_{8+\delta}$ (Bi-2122) samples. The results were shown to be comparable with those obtained using the widely employed MEM. As a result, rRIM can be a suggestive inversion method to analyze data with significant noise.

II. METHODS

The generalized Allen formula, Eq. (1), can be written as

$$1/\tau^{\text{op}}(\omega, T) = F(I^2\chi(\omega, T)),$$

where $F(\cdot)$ represents an integral (forward) operator. Because of the linearity of an integral operator, this equation can be discretized into the following matrix equation:

$$y = Ax. \tag{4}$$

For notational simplicity, we represent the glue function and optical scattering rate as $x \in \mathbb{R}^n$ and $y \in \mathbb{R}^m$, respectively. A is an $m \times n$ matrix containing the kernel information in Eq. (2). Considering the noise in the experiment, Eq. (4) can be written as

$$y = Ax + \eta, \tag{5}$$

where η follows a normal distribution with zero mean and σ^2 variance, that is, $\eta \sim N(0, \sigma^2 I)$. As the matrix $A \in \mathbb{R}^{m \times n}$ is often under-determined ($m < n$) and/or ill-conditioned (i.e., many of its singular values are close to 0), inferring x from y requires additional assumption on the structure of x . On this basis, the goal of the inverse problem is formulated as follows:

$$\hat{x} = \arg \min_{x \in \mathcal{X}} \|y - Ax\|^2. \quad (6)$$

Notably, the solution to Eq. (6) must satisfy two constraints: it minimizes the norm $\|\cdot\|$ and belongs to the data space $\mathcal{X}$, which represents the space of proper glue functions. The norm can be chosen based on the context of the problem; in this study, we selected l_2 norm.

The inverse problem expressed by Eq. (6), with the prior information on x , can be considered a maximum a posteriori (MAP) estimation [30],

$$\hat{x} = \arg \max_x \{\log p(y|x) + \log p(x)\}.$$

The first term on the right-hand side represents the likelihood of the observation y given x , which is determined using the noisy forward model (Eq. (5)). Further, the second term represents prior information regarding the solution space $\mathcal{X}$.

The solution of the MAP estimation can be obtained using a gradient based recursive algorithm expressed as follows:

$$x_{t+1} = x_t + \gamma_t \nabla l(x_t), \quad (7)$$

where $l(x) := \log p(y|x) + \log p(x)$ and γ_t is a learning rate. Determining the appropriate learning rate is challenging. In contrast, the RIM formulates Eq. (7) as

$$x_{t+1} = x_t + g_t(\nabla l(x_t); \phi), \quad (8)$$

where g_t is a deep neural network with learnable parameters ϕ . More explicitly, RIM utilizes a recurrent neural network (RNN) structure as follows (Fig. 1)

$$\begin{aligned} x_{t+1} &= x_t + g_t, \\ \begin{bmatrix} g_t \\ s_{t+1} \end{bmatrix} &= m_\phi(\nabla_t, s_t). \end{aligned} \quad (9)$$

Given the observation y , the inference begins with a random initial value of x_0 . At each time step t , the gradient information $\nabla_t := \nabla_x \log p(y|x)|_{x=x_t}$ is introduced into the update

network m_ϕ along with a latent memory variable s_t . The network’s output g_t is combined with the current prediction to yield the next prediction x_{t+1} . The memory variable s_{t+1} , another output of m_ϕ , acts as a channel for the model to retain long term information, facilitating effective learning and inference during the iterative process [31].

In the training process, the loss is calculated by comparing the training data x (shown in blue in Fig. 1) with the model prediction x_t at each time step. The accumulated loss is calculated formulas follows:

$$\mathcal{L}(\phi) = \sum_{t=1}^{N_t} w_t (x_t(\phi) - x)^2. \quad (10)$$

Here, w_t represents a positive number (in this study, we set $w_t = 1$). The parameters ϕ indicate that x_t is obtained from the neural network m_ϕ and are updated using the backpropagation through time (BPTT) technique [32]. The solid edges in Fig. 1 illustrates the pathways for the propagation of the gradient $\partial\mathcal{L}/\partial\phi$ to update ϕ , whereas the dashed edges indicate the absence of gradient propagation. The trained model can perform an inference to estimate x_t given the input y without referencing the training data x .

In the conventional RIM, the gradient of the log-likelihood is given as

$$\nabla_t = \frac{1}{\sigma^2} A^T (y - Ax_t),$$

where, A^T is the transpose of matrix A (see the Supplemental Material for its derivation.) Owing to the inherent ill-posed nature of the aforementioned equation, it does not yield competitive results as compared to other machine learning approaches. To address this issue, we utilized the equivalence between the RIM framework and iterative Tikhonov regularization, as detailed in the Supplemental Material. Specifically, we utilized the following gradient derived from the preconditioned Landweber iteration [29] to formulate the rRIM algorithm:

$$\nabla_t = (A^T A + h^2 I)^{-1} A^T (y - Ax_t) \quad (11)$$

Here, h is a regularization parameter. Notably, the rRIM demonstrates unique flexibility in determining the appropriate regularization parameter, a task that is typically challenging.

A few important aspects are noteworthy. First, the noisy forward model in Eq. (5) plays a guiding role in both learning and inference throughout the iterative process, as shown in Eqs. (11), which provides key physical insight into the model. Second, the gradient of log-prior information in Eq. (7) is implicitly acquired through the gradient of loss, $\partial\mathcal{L}/\partial\phi$,

which incorporates iterative comparisons between x_t and the training data x . Third, the learned optimizer (Eq. (8)) yields significantly improved results compared to the vanilla gradient algorithm, Eq. (7) [24]. Finally, the total number of time steps, N_t , serves as another regularization parameter that influences overfitting and underfitting. In the case of the rRIM, the output was robust to variations in this parameter.

The iterative Tikhonov regularization algorithm functions as an optimization scheme, while the rRIM effectively addresses optimization problems using a recurrent neural network model. Specifically, rRIM optimizes the iterative Tikhonov regularization algorithm by replacing hand-designed update rules with learned ones, offering flexibility in selecting regularization parameters for iterative Tikhonov regularization. This enables automatic or lenient choices. As a result, the rRIM optimizes the iterative Tikhonov regularization and serves as an efficient and effective implementation of this method. As a result, the rRIM optimizes the iterative Tikhonov regularization and serves as an efficient and effective implementation of this method.

III. RESULTS AND DISCUSSIONS

We generated a set of training data $\{x_n, y_n\}_{n=1}^N$, where x and y represent $I^2\chi$ and $1/\tau^{\text{op}}$, respectively. Generating a robust training dataset that accurately represents the data space is challenging, especially in the context of inverse problems where the true characteristics of solutions are not readily available. In our study, based on existing experimental data, a parametric model employing Gaussian mixtures with up to 4 Gaussians was used to create diverse x values to simulate experimental results. Substituting the generated x values into Eq. (4) yields the corresponding y values. The temperature was set at 100 K. Datasets of varying sizes, $N \in \{100, 1000, 10000, 100000\}$ were generated and split into training, validation, and test data sets with 0.8, 0.1, and 0.1 ratios, respectively. Additionally, noisy samples were created by adding Gaussian noise with different standard deviations, $\sigma \in \{0.00001, 0.0001, 0.001, 0.01, 0.1\}$. The noise amplitude added to each sample y was determined by multiplying σ with the maximum value of that sample. To ensure training stability, we scaled the data by dividing the y values by 300 (see the Supplemental Material for a more detailed description). It is worth noting that, depending on the characteristics of both the data and the model, a more diverse dataset could potentially span a broader range

of the data space and uncover additional solutions.

We trained three models - FCN, CNN, and rRIM - on each dataset; the model architectures are presented in the Supplemental Material. All models were trained using the Adam optimizer [33]. As regards the evaluation metrics, the mean-squared error (MSE) loss,

$$\text{MSE} = \frac{1}{N} \sum_{n=1}^N \left(x_n^{\text{training data}} - x_n^{\text{model output}} \right)^2,$$

was used for the FCN and CNN, and the accumulated error loss (Eq. (10)) for the rRIM. Hyperparameter tuning was performed by using a validation set to obtain an optimal set of parameters for each model. Additionally, early stopping, as described in [32], was applied to mitigate overfitting. For comparison, we calculated the same MSE loss using test datasets for all three models. Kernel smoothing was applied to reduce noise in the inference results. In the rRIM, we set $N_t = 15$ and $h = 0.01$. For reference, we initially trained the FCN, CNN, and rRIM by using a noiseless dataset.

We conducted ten independent training runs for each model, each with different batch setups, and calculated the average test losses. Fig. 2(a) illustrates the trend of the average test losses for various training set sizes N . Across all three models, we observed a consistent pattern wherein the losses decreased as N increased. Notably, the rRIM outperformed both the FCN and CNN across the entire range of N in terms of error size and reliability. Additionally, our findings demonstrate that the CNN yields superior results compared to the FCN, as previously shown in [9].

Figs. 2(c) and (d) illustrate the inference process in the rRIM for the test data following training with the noiseless dataset when $N = 1000$. At each time step, the updated gradient information obtained from Eq. (11) is used to generate a new prediction. Fig. 2(c) shows the inference steps for $t = 1, 2, 3$, and $15 (= N_t)$. The corresponding y values are shown in Fig. 2(d); the intermediate results are omitted because of the absence of significant changes beyond a few initial time steps.

We evaluated the performance of the rRIM with noisy data by following a procedure similar to that employed for the noiseless case. For each of the three models, we conducted 10 independent trials with varying noise levels. The average test losses for the $N = 1000$ data set are shown in Fig. 2(b). Notably, the rRIM exhibits significantly lower test losses than do the other models up to a certain noise threshold, beyond which its performance begins to deteriorate. This behavior can be attributed to the inherently ill-posed nature of the

problem. In contrast to the model-agnostic nature of the FCN and CNN, the rRIM adopts an iterative approach that involves repeated application of the forward model. This iterative process can lead to error accumulation, influenced by factors such as the total number of time steps (N_t), the norm of A , and the noise intensity [34]. When the noise intensity is low, these factors may have a minimal impact on the overall error. However, as the noise intensity increases, their influence becomes more pronounced, potentially resulting in significant error accumulation. Similar patterns are observed for different training set sizes, as detailed in the Supplemental Material. Although it does not accommodate extremely high noise levels, the rRIM is suitable for a wide range of practical applications with moderate noise levels. It is worth noting that the superiority of FCN and CNN in noise robustness over MEM, as demonstrated in previous studies [9, 16], was observed at much lower noise levels than in the present study.

We compared the inference results of the rRIM with those of the FCN and CNN in Figs. 3(a), (b), and (c) by using the $N = 1000$ with $\sigma = 10^{-3}$ dataset. Inferences were made on the test data samples that were not part of the training process. Evidently, the rRIM accurately captures the true data height, whereas the FCN and CNN models tend to overshoot (a) and undershoot (c) it. Only the rRIM accurately replicated the shape of the peak (Fig. 3(b)).

The ability of an algorithm to handle OOD data is crucial for its credibility, particularly in real-world scenarios where the solutions often lie beyond the scope of the dataset. To demonstrate the flexibility of the rRIM, we generated three sample datasets that differed significantly from the training datasets. The first dataset is a simple Gaussian distribution with a large variance, whereas the other two are square waves with different heights and widths (solid lines in Figs. 3(d), (e), and (f)). We inferred these data by using the rRIM, FCN, and CNN (the results are shown using dashed lines). These models were trained using a noiseless dataset of size $N = 100000$. A Comparison of the results with the FCN and CNN models clearly indicates that the rRIM effectively captures the location, height, and width of the peaks in the provided data. Similar patterns are observed for trapezoidal and triangular waves, with detailed results provided in the Supplemental Material.

Finally, we applied the rRIM to real experimental data consisting of optically measured spectra, $1/\tau^{\text{op}}(\omega)$ from one optimally doped sample ($T_c = 96$ K) and two overdoped samples ($T_c = 82$ and 60 K) of Bi-2212, denoted as OPT96, OD82, and OD60, respectively (indicated

using different colors in Fig. 4). Here T_c is the superconducting critical temperature. These measurements were performed at $T = 100$ K [7]. These experimental spectra were fed into the rRIM for inference. Fig. 4(a) presents a comparison of the results of the rRIM (dashed line) with previously reported MEM results (dotted line) [7]. Fig. 4(b) shows a comparison of the reconstructions of the optical spectra from the results of the rRIM (dashed line) and MEM (dotted line) by using Eq. (4) with the experimental results (solid line). The rRIM results were comparable to those of MEM.

IV. CONCLUSIONS

In this study, we devised the rRIM framework and demonstrated its efficacy in solving inverse problems involving the Fredholm integral of the first kind. By leveraging the forward model, we achieved superior results compared to pure supervised learning with significantly smaller training dataset sizes and reasonable noise levels. The rRIM shows impressive flexibility when handling OOD data. Additionally, we showed that the rRIM results were comparable to those obtained using MEM. Remarkably, we established that the rRIM can be interpreted as an iterative Tikhonov regularization procedure known as the preconditioned Landweber method [27, 28]. This characteristic indicates that the rRIM is an interpretable and reliable approach, making it suitable for addressing scientific problems.

Although our approach outperforms other supervised learning-based methods in many respects, it has certain limitations that warrant further investigation. First, we fixed the temperature in the kernel, limiting the applicability of our approach to experimental results at the same temperature. Extending this method for applicability in a wide range of temperatures would require a suitable training set. Another challenge concerns the generation of a robust training dataset that accurately reflects the data space required for a specific problem. This challenge is common to any machine learning approach to inverse problems that entails training data generation using the given forward model. Although the rRIM can partially address this issue by incorporating prior information through iterative comparison with the training data, more robust and innovative solutions are required. Accordingly, we posit that deep generative models [35, 36] must be explored in this regard. While we have shown rRIM’s ability to handle OOD data, a comprehensive quantification of this capability has not been included in the current study. Acknowledging the significance of rRIM’s

ability to manage OOD data, we intend to delve deeper into this aspect in future research endeavors. Finally, uncertainty evaluation must be considered for assessing the reliability of the output; this aspect was not addressed in the proposed approach.

REFERENCES

- [1] Bednorz, J. G. & Muller, A. *Z. Phys. B* **64**, 189 (1986).
- [2] Wu, M. K. *et al.* Superconductivity at 93 k in a new mixed-phase Y-Ba-Cu-O compound system at ambient pressure. *Phys. Rev. Lett.* **58**, 908 (1987).
- [3] Plakida, N. *High-temperature cuprate superconductors* (Springer-Verlag GmbH Berlin, Heidelberg, 2010).
- [4] Allen, P. B. Electron-phonon effects in the infrared properties of metals. *Phys. Rev. B* **3**, 305 (1971).
- [5] Shulga, S. V., Dolgov, O. V. & Maksimov, E. G. Electronic states and optical spectra of htsc with electron-phonon coupling. *Phys. C Supercond. its Appl.* **178**, 266 (1991).
- [6] Schachinger, E., Neuber, D. & Carbotte, J. P. Inversion techniques for optical conductivity data. *Phys. Rev. B* **73**, 184507 (2006).
- [7] Hwang, J., Timusk, T., Schachinger, E. & Carbotte, J. P. Evolution of the bosonic spectral density of the high-temperature superconductor $\text{Bi}_2\text{Sr}_2\text{CaCu}_2\text{O}_{8+\delta}$. *Phys. Rev. B* **75**, 144508 (2007).
- [8] Wazwaz, A.-M. The regularization method for fredholm integral equations of the first kind. *Computers and Mathematics with Applications* **61**, 2981–2986 (2011).
- [9] Yoon, H., Sim, J. H. & Han, M. J. Analytic continuation via domain knowledge free machine learning. *Phys. Rev. B* **98**, 245101 (2018).
- [10] Vapnik, V. N. *Statistical Learning Theory* (John Wiley and Sons, 1998).
- [11] Dordevic, S. V. *et al.* Extracting the electron-boson spectral function $\alpha^2F(\omega)$ from infrared and photoemission data using inverse theory. *Phys. Rev. B* **71**, 104529 (2005).
- [12] Hwang, J. *et al.* *a*-axis optical conductivity of detwinned ortho-ii $\text{YBa}_2\text{Cu}_3\text{O}_{6.50}$. *Phys. Rev. B* **73**, 014508 (2006).
- [13] van Heumen, E. *et al.* Optical determination of the relation between the electron-boson coupling function and the critical temperature in high- T_c cuprates. *Phys. Rev. B* **79**, 184512

- (2009).
- [14] Ito, K. & Jin, B. *Inverse problems: Tikhonov theory and algorithms*, vol. 22 (World Scientific, 2014).
 - [15] Hwang, J. Intrinsic temperature-dependent evolutions in the electron-boson spectral density obtained from optical data. *Sci. Rep.* **6**, 23647 (2016).
 - [16] Fournier, R., Wang, L., Yazyev, O. V. & Wu, Q. S. Artificial neural network approach to the analytic continuation problem. *Phys. Rev. Lett.* **124**, 056401 (2020).
 - [17] Calvetti, D., Morigi, S., Reichel, L. & Sgallari, F. Tikhonov regularization and the l-curve for large discrete ill-posed problems. *Journal of computational and applied mathematics* **123**, 423 (2000).
 - [18] Reichel, L., Sadok, H. & Shyshkov, A. Greedy tikhonov regularization for large linear ill-posed problems. *International Journal of Computer Mathematics* **84**, 1151 (2007).
 - [19] De Vito, E., Fornasier, M. & Naumova, V. A machine learning approach to optimal tikhonov regularization i: affine manifolds. *Analysis and Applications* **20**, 353 (2022).
 - [20] Arsenault, L.-F., Neuberg, R., Hannah, L. A. & Millis, A. J. Projected regression method for solving fredholm integral equations arising in the analytic continuation problem of quantum physics. *Inverse Problems* **33**, 115007 (2017).
 - [21] Park, H., Park, J. H. & Hwang, J. Electron-boson spectral density functions of cuprates obtained from optical spectra via machine learning. *Phys. Rev. B* **104**, 235154 (2021).
 - [22] Nguyen, H. V. & Bui-Thanh, T. Tnet: A model-constrained tikhonov network approach for inverse problems. *arXiv preprint arXiv:2105.12033* (2021).
 - [23] Putzky, P. & Welling, M. Recurrent inference machines for solving inverse problems. *arXiv preprint arXiv:1706.04008* (2017).
 - [24] Andrychowicz, M. *et al.* Learning to learn by gradient descent by gradient descent. In *Proceedings of the 30th International Conference on Neural Information Processing Systems*, NIPS’16, 3988 (Curran Associates Inc., Red Hook, NY, USA, 2016).
 - [25] Adler, J. & Oktem, O. Solving ill-posed inverse problems using iterative deep neural networks. *Inverse Probl.* **33**, 124007 (2017).
 - [26] Morningstar, W. R. *et al.* Data-driven reconstruction of gravitationally lensed galaxies using recurrent inference machines. *The Astrophysical Journal* **883**, 14 (2019).
 - [27] Landweber, L. An iteration formula for fredholm integral equations of the first kind. *American*

- Journal of Mathematics* **73**, 615 (1951).
- [28] Yuan, D. & Zhang, X. An overview of numerical methods for the first kind fredholm integral equation. *SN Applied Sciences* **1**, 1178 (2019).
 - [29] Neumaier, A. Solving ill-conditioned and singular linear systems: A tutorial on regularization. *SIAM Review* **40**, 636 (1998).
 - [30] Figueiredo, M. A. T., Nowak, R. D. & Wright, S. J. Gradient projection for sparse reconstruction: Application to compressed sensing and other inverse problems. *IEEE Journal of Selected Topics in Signal Processing* **1**, 586 (2007).
 - [31] Cho, K. *et al.* Learning phrase representations using RNN encoder-decoder for statistical machine translation. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 1724 (Association for Computational Linguistics, Doha, Qatar, 2014).
 - [32] Goodfellow, I., Bengio, Y. & Courville, A. *Deep Learning* (MIT Press, 2016).
 - [33] Kingma, D. P. & Ba, J. Adam: A method for stochastic optimization. In Bengio, Y. & LeCun, Y. (eds.) *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings* (2015).
 - [34] Groetsch, C. W. *Inverse problems in the mathematical sciences*, vol. 52 (Springer, 1993).
 - [35] Whang, J., Lei, Q. & Dimakis, A. G. Compressed Sensing with Invertible Generative Models and Dependent Noise. *CoRR* **abs/2003.08089** (2020).
 - [36] Ongie, G. *et al.* Deep learning techniques for inverse problems in imaging. *IEEE J. Sel. Areas Inf. Theory* **1**, 39 (2020).

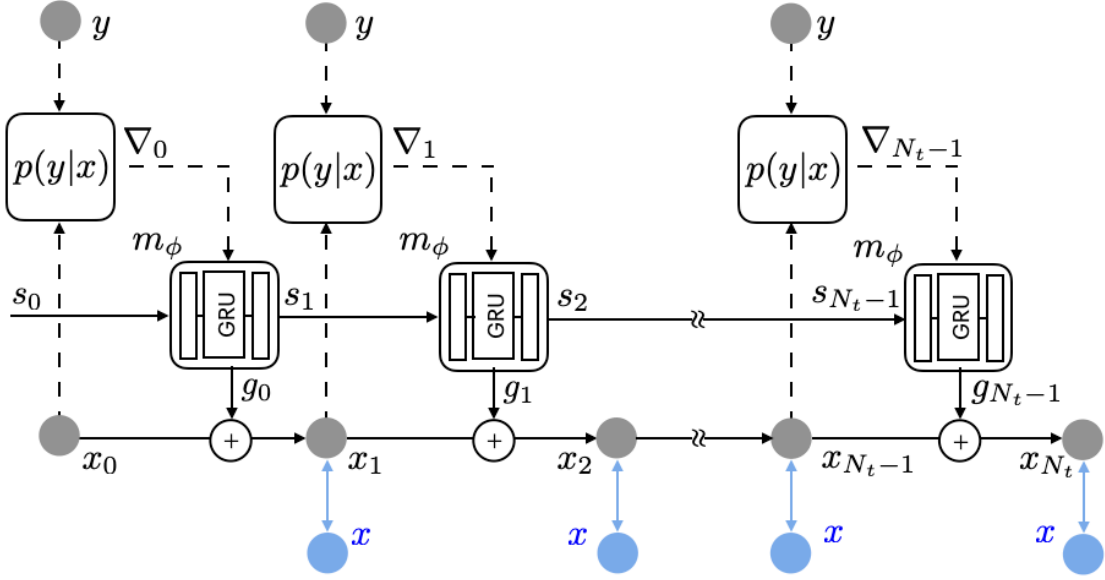


Figure 1. (Color online) Schematic model of rRIM. Shown in blue is only used for training. Further details concerning the update network is provided in the Supplemental Material (adopted and modified from [23]).

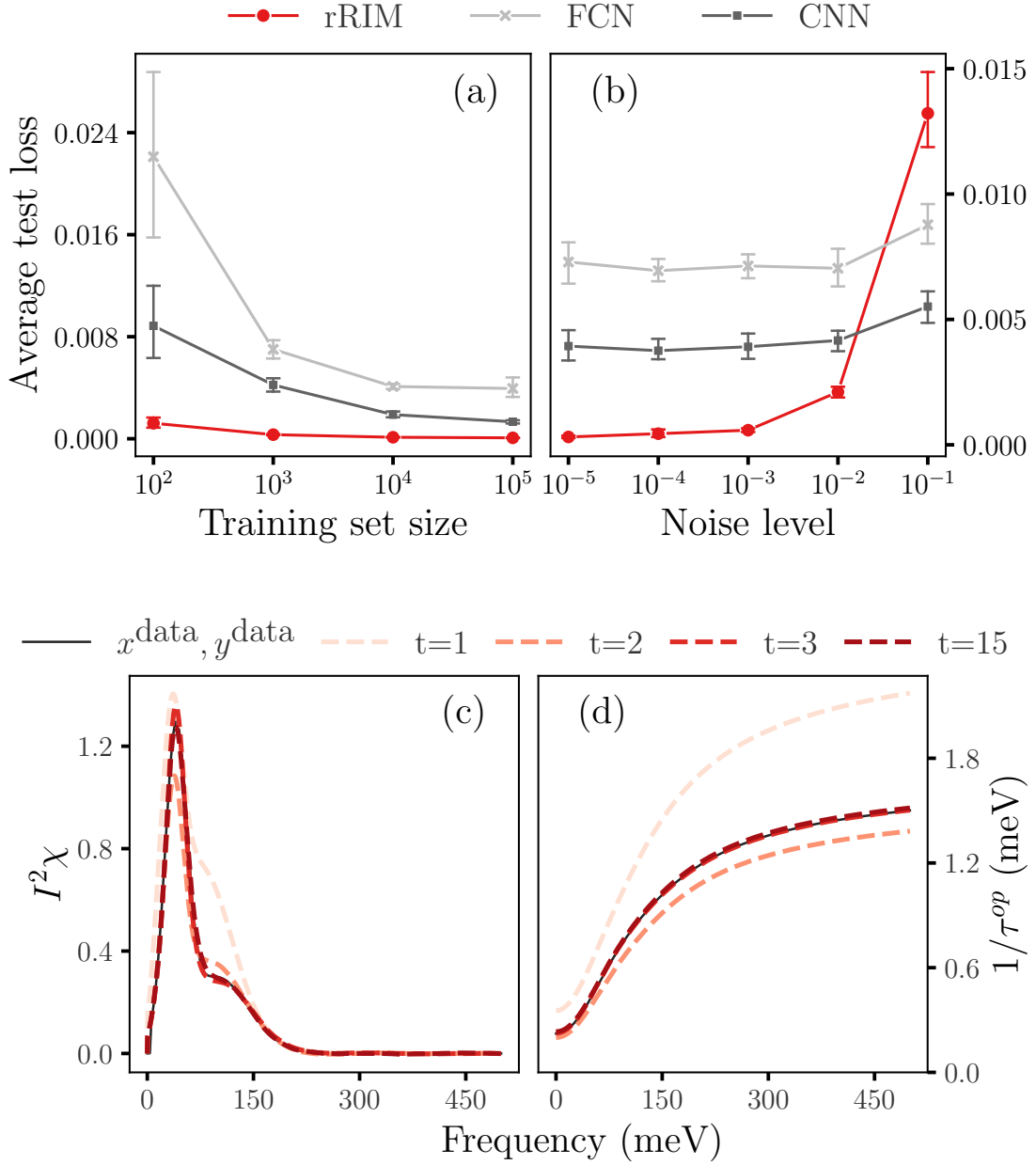


Figure 2. (Color online) Comparison of average test losses for rRIM, FCN, and CNN (a) for different training set sizes N and (b) for different noise levels for $N = 1000$. (c) Inference steps of rRIM for noiseless data for a selection of initial and final predictions (dotted lines) alongside with the true values of x (solid line) and (d) their corresponding y values, which are scaled down by 300.

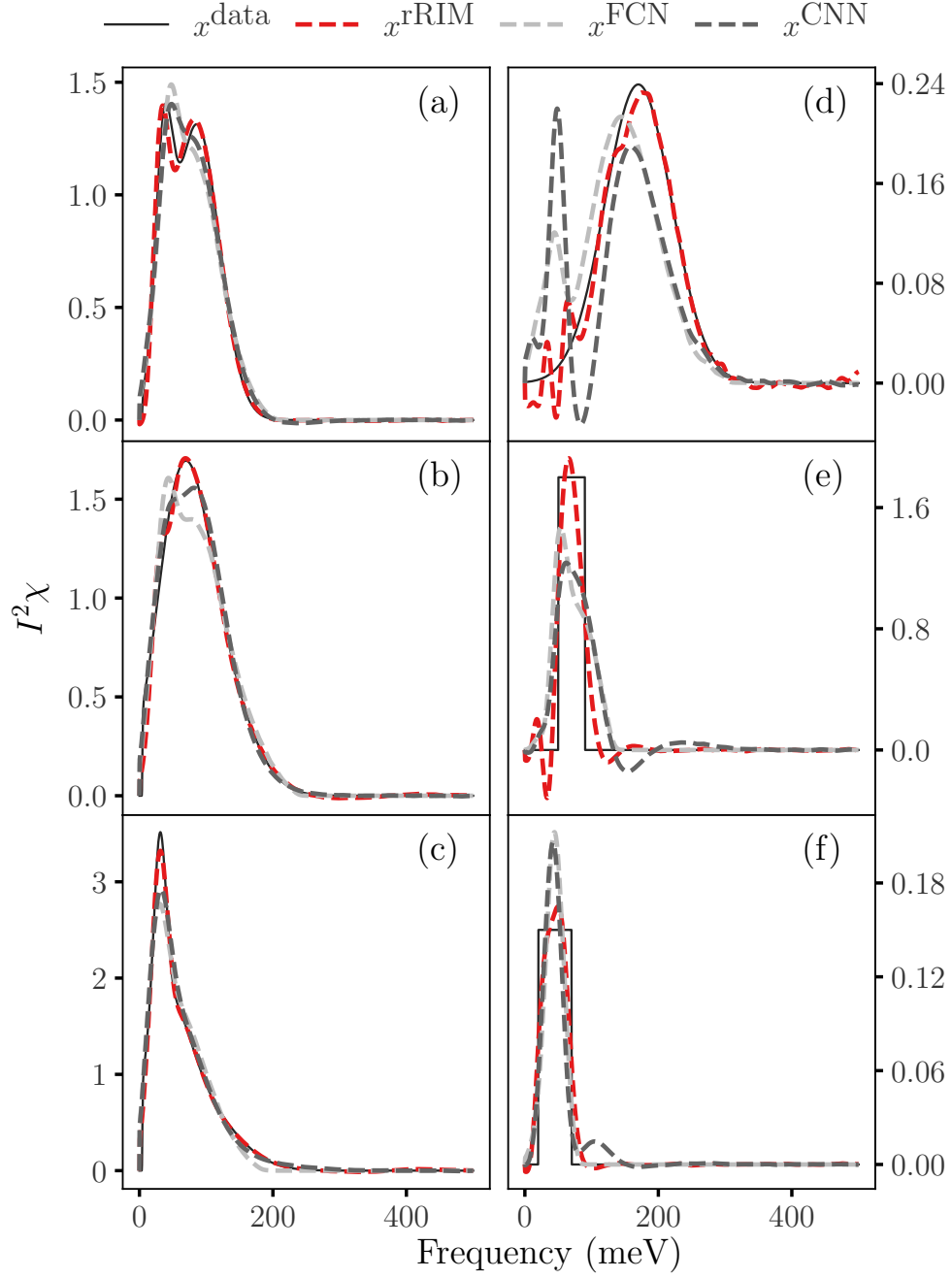


Figure 3. (Color online) Comparison of prediction capabilities of rRIM, FCN, and CNN; for noisy test data samples: (a), (b), and (c); for OOD data samples: (d), (e), and (f).

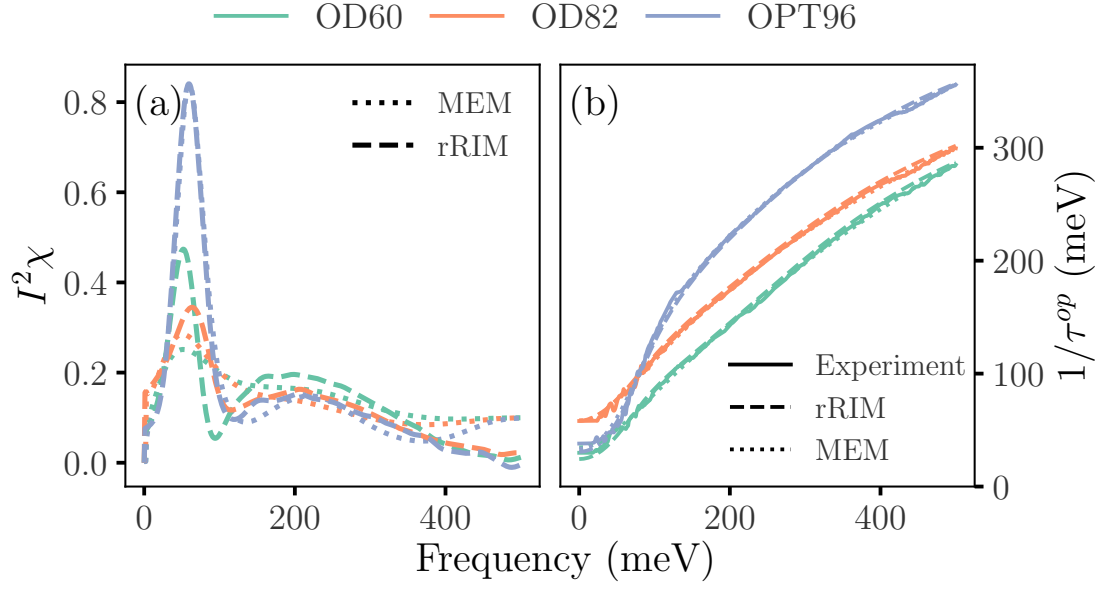


Figure 4. (Color online) (a) Inference results of rRIM (dashed) with those of MEM (dotted) for experimental data. (b) Experimentally measured data (solid) with rRIM (dashed) and MEM (dotted) reconstruction. Samples (OD60, D82, and OPT96) are differentiated using the same color code in both figures.

Author Contribution JH and JP wrote the main manuscript. HP and JP performed the calculations for getting the data in the paper. All authors reviewed the manuscript.

Acknowledgements This paper was supported by the National Research Foundation of Korea (NRFK Grants No. 2017R1A2B4007387 and No. 2021R1A2C101109811).

Competing Interests The authors declare that they have no competing financial interests.

Availability of Data and Materials The datasets used and/or analysed during the current study available from the corresponding author on reasonable request.

Correspondence Correspondence and requests for materials should be addressed to Jungseek Hwang (email: jungseek@skku.edu) and Jun H. Park (jun.park@skku.edu).