

Dynamic coherent diffractive imaging with a physics-driven untrained learning method: supplement

DONGYU YANG,^{1,2,6} JUNHAO ZHANG,^{1,2,6} YE TAO,^{1,2} WENJIN LV,^{1,2} SHUN LU,³ HAO CHEN,² WENHUI XU,^{4,5}  AND YISHI SHI^{1,2,*} 

¹*School of Optoelectronics, University of Chinese Academy of Sciences, Beijing 100049, China*

²*Center for Materials Science and Optoelectronics Engineering, University of Chinese Academy of Sciences, Beijing 100049, China*

³*Institute of Computing Technology, Chinese Academy of Sciences, Beijing 100049, China*

⁴*Department of Electrical and Electronic Engineering, Southern University of Science and Technology, Shenzhen 518055, China*

⁵*Harbin Institute of Technology, Harbin 150001, China*

⁶*These authors contributed equally.*

*optsys@gmail.com

This supplement published with The Optical Society on 16 September 2021 by The Authors under the terms of the [Creative Commons Attribution 4.0 License](https://creativecommons.org/licenses/by/4.0/) in the format provided by the authors and unedited. Further distribution of this work must maintain attribution to the author(s) and the published article's title, journal citation, and DOI.

Supplement DOI: <https://doi.org/10.6084/m9.figshare.16595132>

Parent Article DOI: <https://doi.org/10.1364/OE.433507>

Dynamic coherent diffractive imaging with a physics-driven untrained learning method: supplemental document

Network architecture

The Deep CDI network architecture is shown in Fig. 1 and convolutional blocks detailed in Fig. S1.

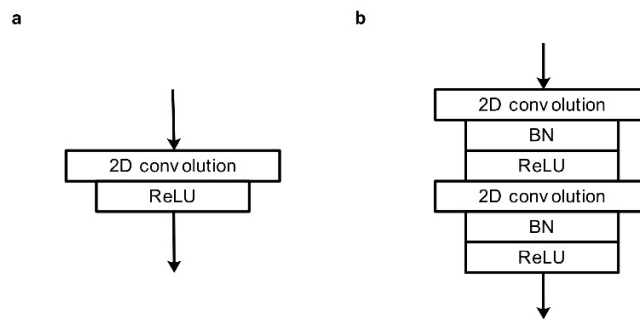


Fig. S1 (a) Single convolutional layer (b) Double convolutional layer. ReLU: rectified linear unit. BN: batch normalization

Single convolutional layer: 2D Convolution with kernel size = 1 and stride = 1.

Double convolutional layer: 2D Convolution with kernel size = 3, stride = 1 and padding = 1.

Max pooling: 2D MaxPool kernel size = 2.

Transposed convolution: 2D ConvTranspose with kernel size = 2 and stride = 2.

Diffraction data of static optical experiment

The diffraction data at corresponding position in Fig. 6a are shown in Supplementary Fig. 2.

The diffraction data are saved as standard TIFF file with a resolution of 512×512 and a pixel depth of 8 bits per pixel.

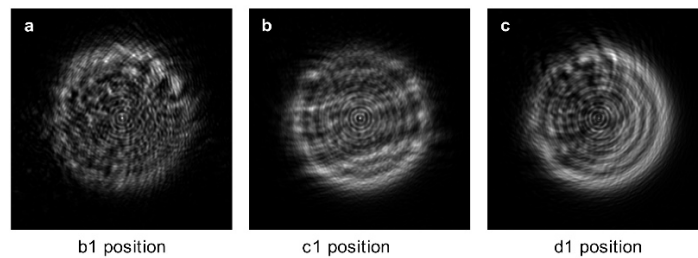


Fig. S2 The static optical experiment diffraction data at corresponding position.

Sample preparation

Small intestine section. In our static experiment, the small intestine section slide is purchased from Saiensi Co. Ltd. Zhejiang (Province) China (Item No.65).

The live rotifer. In our dynamic experiment, the live rotifer were cultured in purified water with chlorella at 20°C for 8h of light a day. A tripette (2.5 μ l) is used to transfer the live rotifer from culture medium to a glass coverslip for imaging.

Generalization performance of Deep CDI by trained with a single diffraction pattern

The usage of networks in Deep CDI is different from that of the conventional end-to-end approach. In the end-to-end approach, the network is supposed to represent a universal function that maps the data in the object space into the image space, which is achieved by training with a large amount of labeled data. In Deep CDI, it is the interplay between the physical constraints and the network that allows the complex field to be reconstructed. The first physical constraint - support region constraint is independent of input data, while the second physical constraint - free propagation constraint at the detector plane relates to the input diffraction measurement itself. Therefore, the network trained by one diffraction measurement cannot be used to directly predict output from another different measurement.

To illustrate this, we trained two Deep CDI models, L-model and H-model, by using the LFW-face data and the Hela cell image, respectively, with the same hyper-parameters and physical parameters. And then, the trained L-model is used to predict the Hela cell, while the trained H-model is used to predict the LFW-face data. The results are presented in Fig. S3. One can clearly see from Fig. S3 (d, e) and (k, l) that both trained models fail to predict the amplitude and phase information from another measurement, as we expected. When dealing with different measurements, the network needs to be retrained or trained from scratch.

Generalization performance of Deep CDI on a time series of diffraction patterns

For a dynamic process, we have demonstrated that the time-consuming reconstruction can be shortened by employing an initial training step with only a fraction of measurements. A large number of measurements is not involved in the training step and is directly used to predict the corresponding outputs. Therefore, for these measurements, it is interesting to compare the reconstruction quality of the trained network (dynamic training) to the one achieved when the network is trained with each one of these frames separately (single diffraction pattern training).

In this section, we analyzed the effect of the two training strategies on the quality of reconstruction images. Three diffraction patterns at different time points, i.e., $t = 1s$, $t = 30s$, $t = 60s$, are selected as examples to examine the performance. For the single diffraction pattern training, the network is trained from scratch using each of the three diffraction patterns separately. For the dynamic training, the results are directly obtained by inputting the diffraction patterns into the trained network in Section 3.3. The reconstruction results by these two strategies are shown in Fig. S4. These two methods are almost the same on the visual effects. As no ground truth can be provided here, we directly compare the similarity of the reconstruction results from the two strategies using SSIM. The SSIM index values of the reconstruction amplitudes between the two methods are 0.9810, 0.9817, and 0.9818 for Fig S4 (a1, a2), (b1, b2) and (c1, c2), respectively. The SSIM index values of the reconstruction phases between the two methods are 0.9823, 0.9818, and 0.9827. Therefore, the performance of the dynamic training strategy is satisfying.

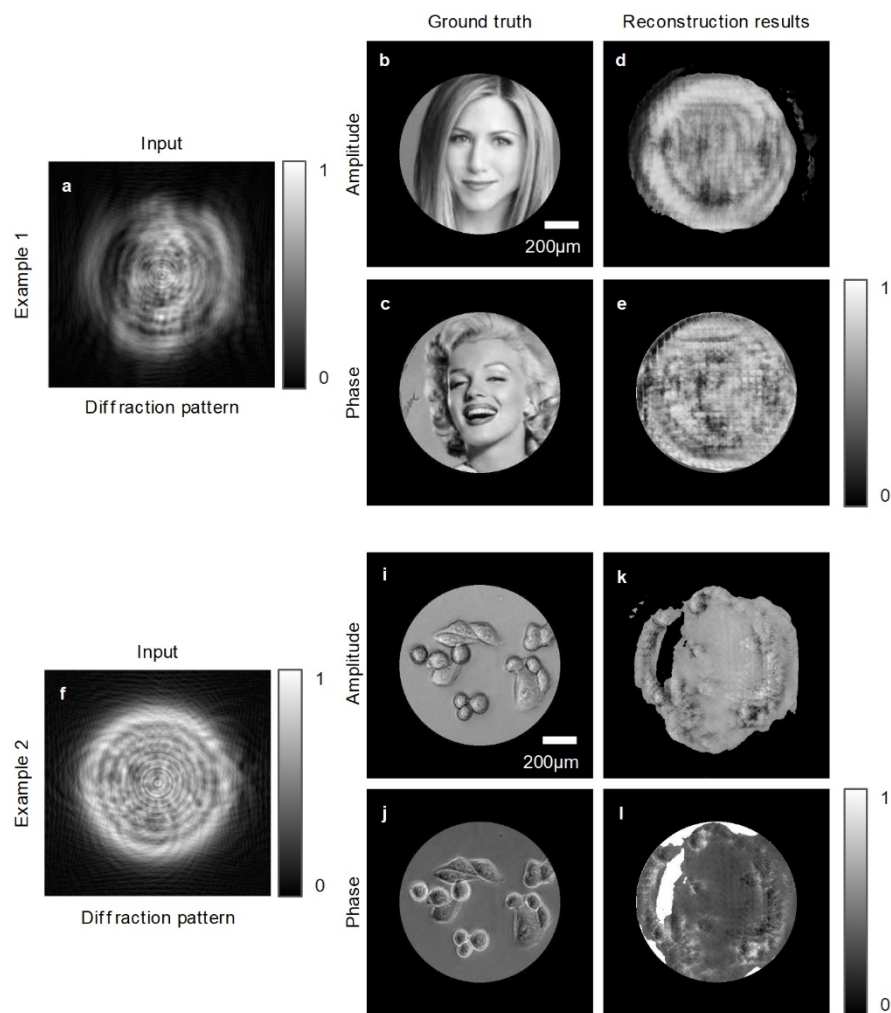


Fig. S3 Generalization performance of Deep CDI by trained with single diffraction pattern. (a, f) the diffraction patterns. (b, c) and (i, j) are the ground truth of a and f, respectively. (d, e) are the reconstruction results by H-model from diffraction pattern f. (k, l) are the reconstruction results by L-model from diffraction pattern a.

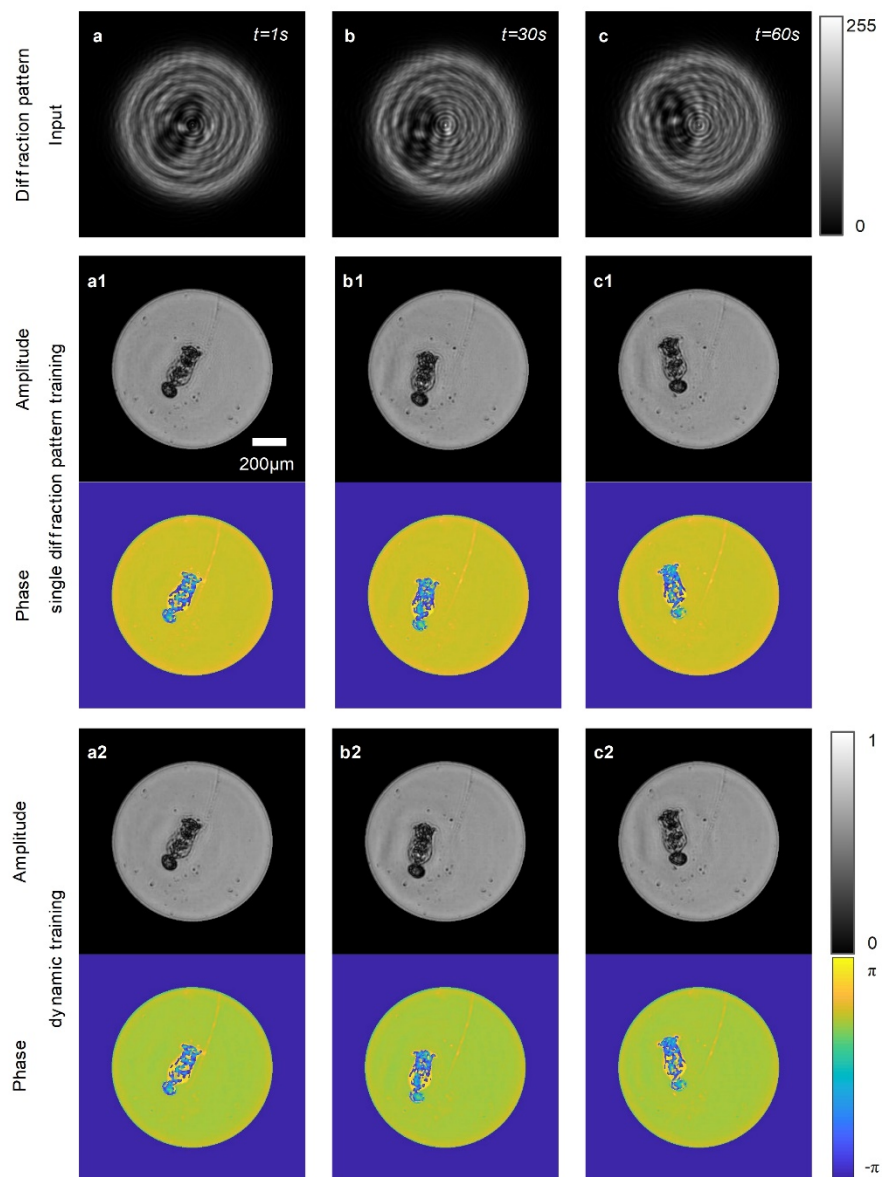


Fig. S4 Comparison of single diffraction pattern training and dynamic training. (a, b, c) are the diffraction patterns at different time points. (a1, b1, c1) and (a2, b2, c2) are the reconstruction results by single diffraction training and dynamic training, respectively