

Learning-based model using unpaired datasets for super-resolution confocal microscopy

Carlos Trujillo^{a,1}, Lauren Thompson^b, Omar Skalli^b, and Ana Doblas^{c,2}

^aApplied optics group, School of Applied Science and Engineering, Universidad EAFIT, Medellín, Colombia

^bDepartment of Biological Sciences, University of Memphis, Memphis, TN 38152, USA

^cDepartment of Electrical and Computer Engineering University of Massachusetts Dartmouth, MA 02747, USA

¹catrujilla@eafit.edu.co, ²adoblas@umassd.edu,

Abstract: One of the major drawbacks of confocal microscopy is its limited spatial resolution. This work assesses the performance of an unpaired learning-based model to provide confocal images with improved resolution. © 2024 The Author(s)

1. Introduction

Confocal Scanning Microscopes (CSMs) are widely used tools that provide valuable morphological and functional information about cells and tissues. The hallmark of confocal microscopy over widefield microscopy is its ability to produce optically sectioned multi-color images in which different organelles within the biological specimen are stained using multiple dyes, enabling colocalization studies. Although the theoretical spatial resolution of a confocal system can surpass the diffraction limit [1], the diffraction limit gives the minimum resolvable distance in confocal images, restricting the applicability of confocal microscopy to nano-scale quantitative analysis. Despite the success of super-resolution microscopic techniques in achieving impressive resolution power, their high price and reduced access have hampered their applicability within the biological and biomedical community. Super-resolution confocal microscopy has been demonstrated using computational approaches based on deconvolution algorithms and deep learning models. Whereas the maximum improvement in the resolution capability was $1.52\times$ using deconvolution [2], deep learning (DL) frameworks have led to a resolution improvement of $1.3\times$ using a U-Net model [3] and $2.64\times$ using a cross-modality training of a conditional generative adversarial network (cGAN) [4]. The latter cross-modality framework converts confocal images to STED ones [4]. A limitation of these DL models is the creation of a paired dataset, requiring paired images of the field of view for two different imaging conditions (native versus improved resolution). Here, we explore an unpaired DL framework for resolution enhancement in confocal microscopy.

2. Unpaired DL framework

The experimental acquisition of a diverse confocal dataset with low-resolution (LR) and high-resolution (HR) image pairs with the same spatial orientation poses many challenges, including the search for the same cell among the hundreds on a microscope slide with two different microscope objective lenses. To circumvent this experimental challenge, we propose to train a cycleGAN model for lateral resolution enhancement in confocal microscopy. Figure 1 shows the architecture for the cycleGAN model, which trains two cGANs models with each other. In

other words, the principle of the cycleGAN model is cycle consistency, enforcing that an image translated from domain A to B should be able to be accurately translated back from B to A domains [5]. In the proposed cycleGAN, the discriminator models (patchGANs) comprise several convolutional layers, each with 64, 128, 256, and 512 filters successively, utilizing a filter size of 4×4 and a stride of 2×2 for downsampling. Leaky ReLU activations with an alpha value of 0.2 follow each convolutional layer, and instance normalization is applied. The generator models (cGANs) comprise convolutional layers with 64, 128, 256, and 512 filters, where the first layer uses a filter size of 7×7 , and the subsequent layers utilize a filter size of 3×3 . Additionally, each generator incorporates nine residual blocks, each with 512 filters. Convolutional transpose layers are employed for upsampling, with filter sizes of 3×3 and strides of 2×2 . A composite model coordinates the training process, with a learning rate of 0.0002 and a beta value of 0.5 for the Adam optimizer. The loss function of the composite model integrates one adversarial loss and three

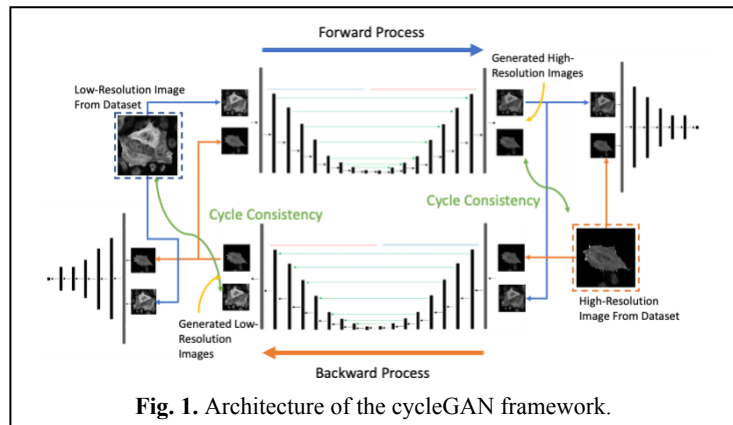


Fig. 1. Architecture of the cycleGAN framework.

perceptual losses (Identity loss, Forward cycle consistency loss, and Backward cycle consistency loss), with a weight ratio of 1:5:5:5, thereby ensuring a balanced integration of their respective contributions during the training process.

3. Results

An experimental confocal dataset was acquired using the Nikon Ti-E A1rSi confocal microscope at the University of Memphis Integrated Microscopy Center. The experimental biological samples U-373 MG human Glioblastoma cells stained by immunofluorescence for the cytoskeletal protein vimentin. We recorded 11 different fields of view of the sample with a $60\times/1.4$ NA microscope objective lens. Each field of view has 2048×2048 pixels². The pixel and optical resolution were 0.10 and 0.23 μm . The HR confocal image was generated by recording a 3D volume and computationally estimating the extended depth of field (EDF) image. We also recorded 22 LR images with a $10\times/0.3$ NA microscope objective lens with a field of view of 512×512 pixels². The HR pixel and optical resolution were 0.45 and 1.03 μm , respectively. To match pixel resolution between the LR and HR images, the LR images were scaled $4.5\times$, generating LR images of 2304×2304 pixels². The field of view was matched between the image sets by cropping areas of 512×512 . Then, we performed dataset augmentation on the LR and HR image sets by randomly rotating our images by 0, 90, 180, and 270 degrees and flipping them vertically and horizontally. Finally, we have a dataset of 585 LR and HR images, respectively. We used an 80-20 training/validation split.

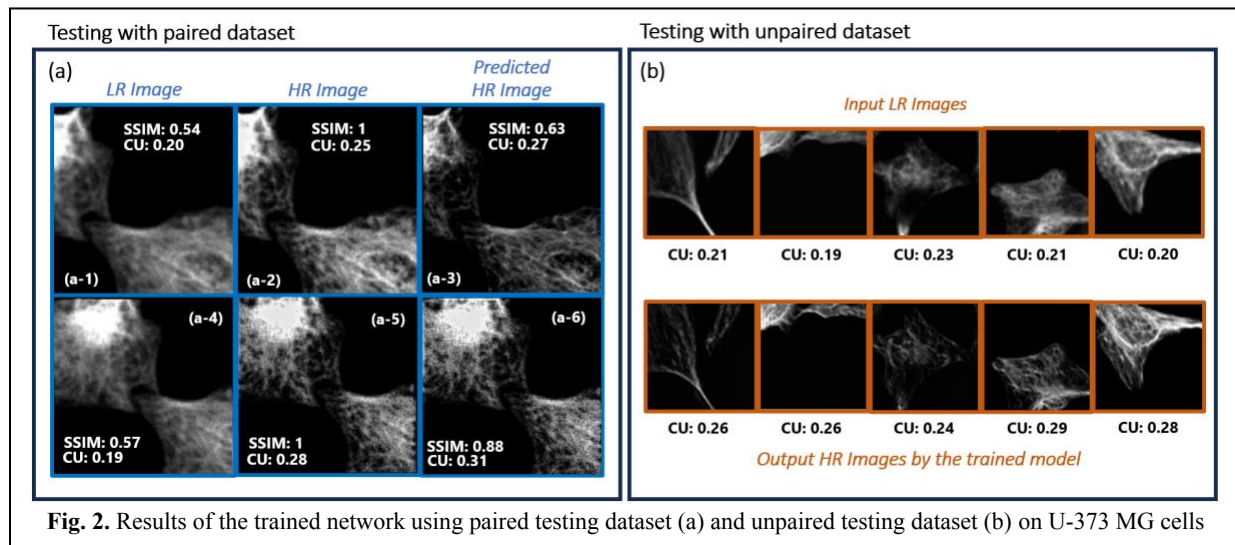


Fig. 2. Results of the trained network using paired testing dataset (a) and unpaired testing dataset (b) on U-373 MG cells

The trained model translating in creating HR images from LR ones is assessed using a paired dataset of 68 images. This paired dataset was compiled by acquiring images of the same two cells for both microscope objective lenses. The model was trained for 100 epochs. The Structural Similarity Index Measure (SSIM) was used to evaluate the performance of the trained generator model. Figure 2a shows the performance of the trained model using paired testing datasets. The predicted HR images for both fields of view are consistently closer to the ground-truth HR images, as the SSIM value increases from the input LR image to the predicted HR image. Although the predicted HR image for the first field of view has a relatively SSIM value (i.e., 0.63), certain cells' structures, which are initially imperceptible in the LR image, become discernible akin to the ground-truth HR image. A similar pattern is observed in the second field of view. Increasing the resolution limit involves an increase in the cutoff spatial frequency. In other words, the size of the compact support in the images' spectrum should be enlarged. For this reason, we proposed the cutoff (CU) metric that quantifies the bandwidth of the images' spectrum. Images with a higher resolution power should yield a higher CU value. Figure 2a also reports the estimated CU values, demonstrating that the predicted HR images contain more high spatial frequencies. Finally, Fig. 2b shows the performance of the trained network using an unpaired testing dataset. Again, the spatial bandwidth in the predicted HR images has been increased (i.e., increase of the CU values), validating the capability of unpaired cycleGAN models to super-resolution confocal microscopy.

4. References

1. M. Martinez-Corral, and G. Saavedra, Progress in Optics 53, 1-67 (2009).
2. J. B. de Movel, S. le Calvez, and M. Ulfendahl, Biophysical Journal 80, 2455 – 2470 (2001).
3. W. Wang, B. Wu, B. Zhang, J. Ma, and J. Tan, Optics Letters 46(19), 4932 (2021).
4. H. Wang, Y. Rivenson, Y. Jin, et al., Nature Methods 16, 103–110 (2019).
5. J.-Y. Zhu, T. Park, P. Isola, and A. A. Efros, 2017 IEEE International Conference on Computer Vision (ICCV), 2242-2251(2017).