

Computationally efficient pixelwise deep learning architecture for accurate depth reconstruction for single-photon LiDAR

Yu Zhang, Yiming Zheng

Department of Electrical Engineering, Viterbi School of Engineering, University of Southern California, Los Angeles, United States

Article Info

Article history:

Received Jul 27, 2024

Revised Aug 6, 2025

Accepted Sep 16, 2025

Keywords:

Computational imaging
Computational intelligence
Deep learning
Depth reconstruction
Single-photon LiDAR

ABSTRACT

This work introduces a compact deep learning architecture for depth image reconstruction from time-resolved single-photon histograms. Unlike most deep learning approaches that mainly rely on 3D convolutions, our network is implemented purely with 1D convolutions without assistance from other sensors or pre-processing. Both synthetic and real datasets were used to evaluate the accuracy of our model for challenging signal-to-background ratios (SBRs), ranging from 5:1 to 1:1. Conventional maximum likelihood (ML) and another photon-efficient optimization-based algorithm were adopted for performance comparisons. Results from synthetic data show that our model achieves lower mean absolute error (MAE). Additionally, results from real data indicate that our model exhibits better reconstruction for high-ambient effects and provides better spatial information. Unlike existing 3D deep learning models, we process pixel-wise histograms continuously, rather than splitting the point cloud and stitching them afterward, which saves memory and computational resources, thereby laying a foundation for real-world embedded applications.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



Corresponding Author:

Yu Zhang

Department of Electrical Engineering, Viterbi School of Engineering, University of Southern California
Los Angeles, CA 90089

Email: yz324@usc.edu

1. INTRODUCTION

Single-photon avalanche diodes (SPADs) have been emerging for various applications that rely on single-photon sensitivity, such as single-photon light detection and ranging (LiDAR) using an optimization-based reconstruction method [1], deep learning methods [2], [3], and biomedical signal processing for fluorescence lifetime imaging (FLIM) [4], and non-line-of-sight imaging [5], [6] and cryptography [7], [8]. Researchers proved that using data-driven deep learning (DL) models can accurately reconstruct depth and reflectivity images from the 3D point cloud cubes that include photons' time-of-flight information and spatial information. Further, these DL models are robust for extremely low signal-to-background ratios (SBRs), even less than one. Reconstructing depth information is crucial in autonomous vehicles that need fast and accurate response, even in low-visibility environments. Although data-driven methods reconstruct depth images based on SPAD are emerging, there are still challenges for the DL models. First, most DL models are composed of 3D, consuming enormous computing memory for the computing platform. Even for high-performance graphics processing units (GPUs), the big point cloud with big spatial resolution should be divided into several batches for processing and stitched eventually to obtain a high-resolution depth image. Second, the preparation of training datasets is complex, leveraging image processing tools [9] and large open-source [10], [11] depth training datasets, also consuming a long time (several hours). This work aims to design a computationally efficient, pixel-wise DL model with a compact architecture and training pipeline to address these two bottlenecks.

The contributions of this study are summarized as follows: i) We leveraged existing open-sourced depth image datasets to generate pixel-wise histograms using an analytical mathematical model and trained a deep neural network; ii) We designed a compact 1D U-NET deep neural network for accurate, end-to-end pixel-wise depth in conditions of low SBRs; and iii) We quantitatively compared our model with existing photon-efficient pixel-wise algorithms and achieved better accuracy.

The structure of this study is as follows. Section 2 reviews and summarizes the existing work. Section 3 illustrates the mathematical model of the single-photon LiDAR. Section 4 presents the deep learning architecture, training synthetic data generation, and training details. Section 5 quantitatively evaluates the performance of the DL model and compares it with photon-efficient pixel-wise algorithms and optimization-based methods. Section 6 concludes the study.

2. PRIOR WORK

2.1. Optimization- and statistic-based algorithms

Shin *et al.* [1] first reported a photon-efficient optimization-based method to reconstruct depth images from histograms with extremely low SBR ratios. A signal and noise unmixing optimization algorithm was proposed [12] to accurately split signals from noisy histograms with strong ambient light and accurately reconstruct depth and reflectivity images. The reversible jump Markov chain Monte Carlo (RJ-MCMC). A statistical algorithm [13] was used to perform Bayesian inference for depth and intensity reconstruction of 3D scenes. An improved RJ-MCMC enhanced by a point cloud denoising [14] approach was proposed to achieve real-time 3D reconstruction of moving objects. Koo *et al.* [15] combined a statistical Bayesian algorithm with a deep learning architecture, taking advantage of both accurate inference and model-free properties of statistics and deep learning. A computationally efficient Bayesian algorithm was also proposed [16] for a low-photon-count multispectral LiDAR application.

2.2. Deep learning algorithms

Deep learning is becoming prevalent in feature extraction in computer vision [17], [18]. Deep neural networks have been extensively leveraged in depth reconstruction for SPAD arrays equipped with time-correlated single-photon counting (TCSPC). A sensor fusion [19] 3D deep neural architecture was first introduced to merge high-resolution intensity and low-resolution depth images to enhance the spatial feature extraction during the training. The captured raw data in this work was widely adopted for subsequent work in this field [20], [21]. Another fusion architecture was reported to merge monocular depth images with 3D point cloud convolution modules to enhance depth image reconstruction. Two different architectures were investigated for non-fusion architectures that only leverage the point cloud from the SPAD array without other features from other sensors, where results indicated that the non-fusion architecture could achieve comparable accuracy to fusion-based architectures. A 3D convolutional architecture with pixel-wise residual shrinkage [3] was reported to redefine the optimization target as a classification for each histogram, achieving high reconstruction accuracy. Study [22] presented an edge-enhanced architecture, embedding attention modules in their 3D convolutional architecture, to improve the edge reconstruction. Sparsity in the point cloud was investigated [23] to accelerate the inference of a 3D architecture, achieving real-time high-resolution depth reconstruction. SPAD is also used for sensing through fog [24].

Existing statistical and optimization methods are high-latency and unsuitable for embedded hardware in single-photon LiDAR systems. Moreover, despite the fast forward propagation of DL models, 3D tensor processing of point clouds remains computationally intensive on hardware. This work bridges the gap between DL and computationally efficient methods by leveraging 1D histogram processing, resulting in a compact DL architecture and simplifying the synthetic data generation process.

3. PROBLEM DEFINITION

The active single-photon imaging systems have been reported in existing studies [1], [3], [12], [21]. SPAD arrays with TCSPC-based LiDAR systems can be well approximated and modelled using known optical and sensor parameters. We aim to reconstruct depth information from the histogram of each pixel, which is subject to an inhomogeneous Poisson process [22]. Therefore, the histogram can be estimated as (1).

$$\Phi(t) = \eta \cdot (\epsilon(t) + n_b) + n_d, \quad (1)$$

where $\eta \in (0, 1)$, indicating the quantum efficiency of the sensor. n_b and n_d are the background noise and dark-count noise. $\epsilon(t)$ is the signal flux reflected from the target, which can be modelled as (2),

$$\epsilon(t) = \alpha \cdot s\left(t - \frac{2D}{c}\right) \quad (2)$$

where α is the attenuation factor, D is the distance from the sensor to the target, and c is the speed of light. Therefore, the histogram represents that the reflected photons of N illuminations can be modelled as (3),

$$h(t) \sim \text{Poisson}(N)\Phi(t) \quad (3)$$

By following the equation for analytically generating synthetic training datasets, we employ the datasets to train a deep neural network, which is discussed in the next section. This pixel-wise processing DL alleviates computational complexity compared with 3D-based DL architecture, simplifying the feature extraction from complex 3D latent space to 1D latent space while maintaining accuracy.

4. DEEP LEARNING ARCHITECTURE

Inspired by previous U-NET-like 3D architectures [2], [20], [21], we proposed a similar topology, but only 1D convolution was used. Batches of histograms were fed into the U-NET for training. Key modules in the network (down-sampling, concatenation, convolution, up-sampling) are colored in Figure 1 and indicated in the black dashed box. Batch normalization modules were used to improve the training stability and convergence speed. The ground truth depth images in the training datasets are from the NYUV2 datasets [13]. A texture filtering [12] algorithm was used to alleviate the imperfection of depth information due to the Kinect camera. Histograms were generated using (1), (2), and (3). In the last layer, the multi-channel feature is processed into a single-channel feature and processed by an $\text{argmax}(\cdot)$ to find the peak index, thereby calculating the distance. Notably, unlike previous training datasets of 3D architecture that generate huge point clouds (tens of gigabytes) from multiple scenes in the datasets, our model only generates histograms from one scene, where the training datasets are just 15.1 MB. Similarly, our training speed is approximately ten of times faster than previous 3D U-NET architectures [2], [20], and [21]. While generating histograms in the training datasets, we defined optical parameters presented in Tables 1 to 5.

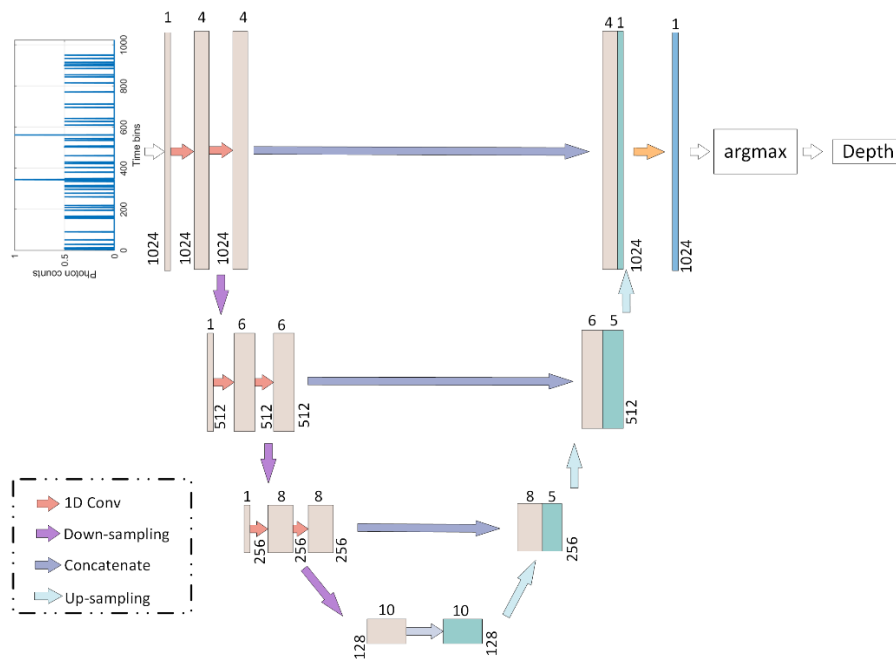


Figure 1. The inference pipeline of the 1D U-NET architecture for pixel-wise depth image

Table 1. Pre-defined parameters to generate synthetic histograms

Parameter	Value
Number of time bins	1024
Temporal resolution	19.53 ps
Spatial resolution	64×64
SBRs	[0.01, 5]
Laser FWHM	117.18 ps
Laser FWHM peak index	5

Table 2. Downsample the details of the deep neural network

Down sample			
Layer name	Output shape	Activation function	
Conv-down sample, K (9), S (2), P (4) + Batch normalization	(256, 1, 1, 512)	ReLU	
Conv-down sample, K (7), S (2), P (3) + Batch normalization	(256, 1, 1, 256)	ReLU	
Conv-down sample, K (5), S (2), P (2) + Batch normalization	(256, 1, 1, 128)	ReLU	

Table 3. Convolutional layers details of the deep neural network

Conv-same spatial dimension			
Layer name	Output shape	Activation function	
Conv, K (9), S (1), P (4) + Batch normalization	(256, 1, 4, 1024)	ReLU	
Conv, K (9), S (1), P (4) + Batch normalization	(256, 1, 4, 1024)	ReLU	
Conv, K (9), S (1), P (4) + Batch normalization	(256, 1, 6, 512)	ReLU	
Conv, K (9), S (1), P (4) + Batch normalization	(256, 1, 6, 512)	ReLU	
Conv, K (7), S (1), P (3) + Batch normalization	(256, 1, 8, 256)	ReLU	
Conv, K (7), S (1), P (3) + Batch normalization	(256, 1, 8, 256)	ReLU	
Conv, K (5), S (1), P (2) + Batch normalization	(256, 1, 10, 128)	ReLU	
Conv, K (5), S (1), P (2) + Batch normalization	(256, 1, 10, 128)	ReLU	

Table 4. De-convolutional layers details of the deep neural network

Up sample			
Layer name	Output shape	Activation function	
ConvTrans., K (9), S (2), P (4) + Batch normalization	(256, 1, 5, 256)	ReLU	
ConvTrans, K (9), S (2), P (4) + Batch normalization	(256, 1, 5, 512)	ReLU	
ConvTrans, K (9), S (2), P (4) + Batch normalization	(256, 1, 5, 1024)	ReLU	
ConvTrans, K (9), S (2), P (4) + Batch normalization	(256, 1, 5, 1024)	ReLU	

Table 5. De-convolutional layers details of the deep neural network

Refine			
Layer name	Output shape	Activation function	
Conv, K (9), S (2), P (4) + Batch normalization	(256, 1, 1, 1024)	-	

The deep learning model is implemented using PyTorch and runs on an NVIDIA RTX A1000 GPU. The learning rate is set to 10^{-5} , and RMSprop is the optimizer. Kullback-Leibler (KL) divergence is employed as the loss function to evaluate the model's performance. An early stopping mechanism is incorporated with a patience of 20 epochs to prevent overfitting. The dataset comprises 50,000 histograms for training, with an additional 5,000 for validation during training. Signal-to-background ratios (SBRs) for the training datasets are set to 5, 2.5, 1, 0.5, 0.2, 0.1, 0.05, and 0.01, consistent with other photon-efficient architectures [2], [4], [22]. The training and validation losses are shown in Figure 2. This setup ensures robust training and evaluation of the deep learning model, leveraging PyTorch's capabilities and harnessing the computational power of the NVIDIA RTX A1000 GPU for efficient processing.

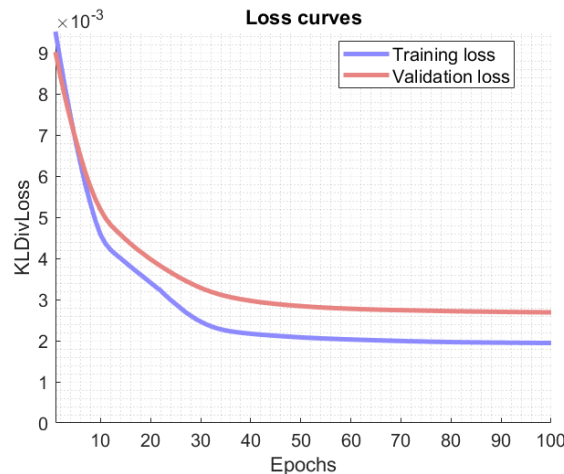


Figure 2. KL divergence loss curve of training and validation

5. QUANTITATIVE EVALUATION

This section assesses the precision of depth reconstruction achieved by our deep learning architecture. It juxtaposes it with conventional maximum likelihood (ML) methods and Shin *et al.* optimization-based algorithm. Synthetic datasets are meticulously simulated, ensuring a comprehensive evaluation framework.

5.1. Synthetic datasets evaluation

We used depth images in the Middlebury datasets [11] as the ground truth (GT) depth image and generated synthetic histograms using known optical parameters for our network’s evaluation. The datasets were also leveraged by other SPAD-based depth image reconstruction using deep learning [2], [3], [21]. As shown in Figure 3, our network is robust for low SBRs. We also compared our network with ML and Shin *et al.* methods, which are also pixel-wise. Shin *et al.* algorithm also used a pixel-averaging method to enhance the accuracy of spatial dimensions. The comparison across seven different SBRs from seven scenes. The results are shown in Table 6.

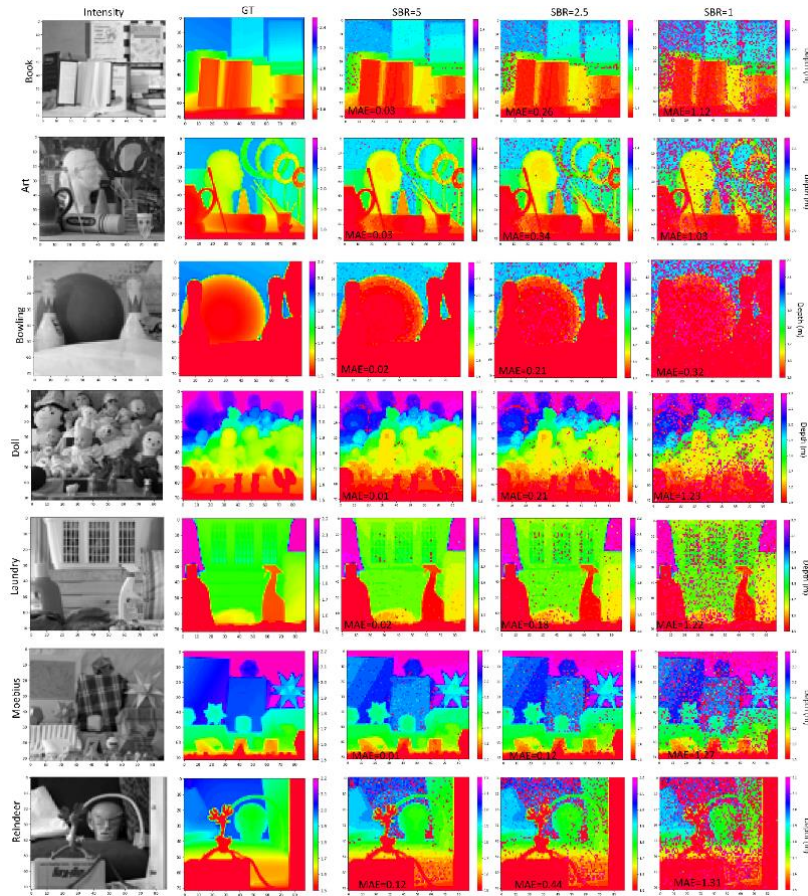


Figure 3. Synthetic datasets. The point cloud is simulated using pre-defined optical parameters. Different SBRs are 1, 2.5, and 2. MAEs of each reconstructed image versus the GT images are indicated in each image

Table 6. Accuracy comparisons among three pixel-wise reconstruction algorithms in synthetic training datasets

Algorithm	SBR	Book	Art	Bowling	Doll	Moebius	Reindeer
ML	5	2.59	2.52	2.13	2.55	2.48	2.56
	2.5	4.78	4.66	4.32	4.70	4.53	4.57
	1	6.87	6.79	6.32	6.73	6.54	6.69
Shin <i>et al.</i>	5	2.25	2.14	2.07	2.21	2.11	2.54
	2.5	4.53	4.44	4.03	4.58	4.46	4.53
	1	6.44	6.23	6.02	6.57	6.52	6.45
1D UNET	5	0.03	0.03	0.02	0.01	0.02	0.12
	2.5	0.26	0.34	0.21	0.21	0.18	0.44
	1	1.12	1.03	0.32	1.23	1.22	1.31

5.2. Captured datasets evaluation

Apart from the evaluation of synthetic datasets, we also investigated the performance of captured datasets [20]. We also compared our deep neural network with ML and Shin *et al.* method. As shown in Figure 4, we present the reflectivity images to reference spatial information. The intensity images were retrieved from the scanned point cloud, and all histograms were taken at the temporal dimension. Regarding ML's performance, the reconstruction depth images contain numerous NaN values represented by white pixels in the depth images due to the low photon counts. Our approach achieved a comparable reconstruction to Shin's method. Notably, Shin's method is sometimes susceptible to intense ambient light. For example, the bulb was not reconstructed robustly in the lamp scene. And our method achieved better visualization of the bulb. Also, as Shin's algorithms involve a spatial averaging process, the spatial depth might be worse if pixel-wise depth information is not recovered accurately. Future work can employ more advanced neural networks, such as a graph neural network (GNN) [25] for point cloud analysis [23].

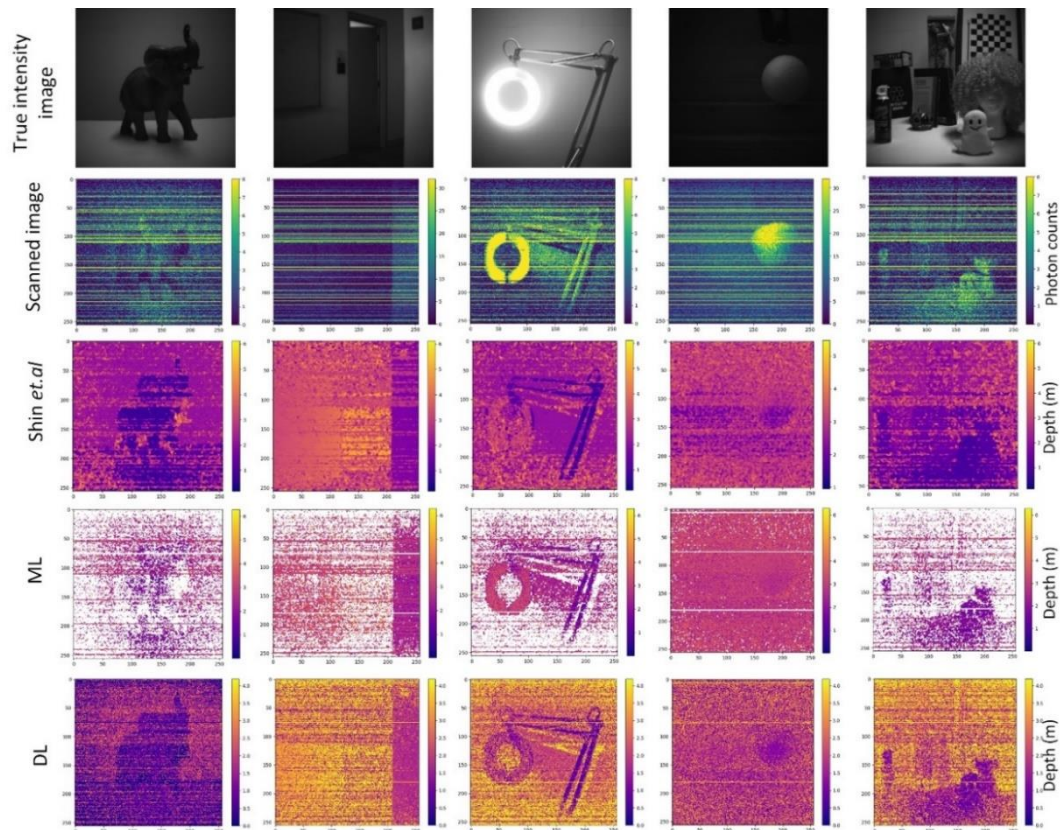


Figure 4. Reconstructed depth images of the captured point cloud. There are five scenes: an elephant doll, a hallway, a lamp, a ball on a staircase, and stuff on a table. Images of reflectivity, intensity, and reconstruction are depicted

6. CONCLUSION

This work presents a compact and accurate 1D depth image reconstruction from histograms of SPAD arrays with low SRBs. Compared with the previous 3D deep neural networks that require tens of gigabytes of training datasets of point clouds, our network only requires a 15.1 MB histogram training dataset. Additionally, the 3D networks consume hours to finish training, whereas our 1D architecture only requires 12 minutes. Similarly, for inference, the high spatial resolution point cloud for 3D networks should be divided into small portions, for example, 1/8 spatial resolution, to infer partial depth images in multiple batches and stitch the depth images afterwards. 3D convolutions consume huge GPU memory and cannot be processed in one batch. However, our 1D pixel-wise architecture does not have the memory overflow issue due to lightweight 1D convolutions, making it easier to implement on embedded hardware in vehicles or drones for practical applications. Compared with conventional machine learning and other photon-efficient algorithms, our methods show higher accuracy for synthetic datasets. As for the evaluation of captured datasets, our network is more robust against ambient light. The limitation of this work is that no spatial

information is extracted during DL training due to the pixel-wise processing nature. A potential approach to address this could involve incorporating a low-cost RGB image to provide 2D spatial structural details during training, compensating for the lack of spatial resolution.

FUNDING INFORMATION

Authors state no funding involved.

AUTHOR CONTRIBUTIONS STATEMENT

This journal uses the Contributor Roles Taxonomy (CRediT) to recognize individual author contributions, reduce authorship disputes, and facilitate collaboration.

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Yu Zhang	✓	✓	✓	✓	✓	✓	✓		✓	✓			✓	
Yiming Zheng		✓		✓		✓			✓	✓	✓			

C : Conceptualization	I : Investigation	Vi : Visualization
M : Methodology	R : Resources	Su : Supervision
So : Software	D : Data Curation	P : Project administration
Va : Validation	O : Writing - Original Draft	Fu : Funding acquisition
Fo : Formal analysis	E : Writing - Review & Editing	

CONFLICT OF INTEREST STATEMENT

The author declares that there is no conflict of interest.

DATA AVAILABILITY

The data supporting this study's findings are available from the corresponding author, YZ, upon reasonable request.

REFERENCES

[1] D. Shin *et al.*, "Photon-efficient imaging with a single-photon camera," *Nature Communications*, vol. 7, no. 1, Jun. 2016, doi: 10.1038/ncomms12046.

[2] Z. Zang, D. Xiao, and D. Day-Uei Li, "Non-fusion time-resolved depth image reconstruction using a highly efficient neural network architecture," *Optics Express*, vol. 29, no. 13, p. 19278, Jun. 2021, doi: 10.1364/oe.425917.

[3] G. Yao, Y. Chen, Y. Liu, X. Hu, and Y. Pan, "Robust photon-efficient imaging using a pixel-wise residual shrinkage network," *Optics Express*, vol. 30, no. 11, p. 18856, May 2022, doi: 10.1364/oe.452597.

[4] D. Xiao, Y. Chen, and D. D. U. Li, "One-dimensional deep learning architecture for fast fluorescence lifetime imaging," *IEEE Journal of Selected Topics in Quantum Electronics*, vol. 27, no. 4, pp. 1–10, Jul. 2021, doi: 10.1109/JSTQE.2021.3049349.

[5] M. O'Toole, D. B. Lindell, and G. Wetzstein, "Confocal non-line-of-sight imaging based on the light-cone transform," *Nature*, vol. 555, no. 7696, pp. 338–341, Mar. 2018, doi: 10.1038/nature25489.

[6] X. Su, Y. Hong, J. T. Ye, F. Xu, and X. Yuan, "Model-guided iterative diffusion sampling for NLOS reconstruction," *IEEE Journal of Selected Topics in Quantum Electronics*, vol. 30, no. 1, pp. 1–11, Jan. 2024, doi: 10.1109/JSTQE.2024.3350971.

[7] P. Keshavarzian *et al.*, "A 3.3-Gb/s SPAD-based quantum random number generator," *IEEE Journal of Solid-State Circuits*, vol. 58, no. 9, pp. 2632–2647, Sep. 2023, doi: 10.1109/JSSC.2023.3274692.

[8] A. Tontini, L. Gasparini, N. Massari, and R. Passerone, "SPAD-based quantum random number generator with an nth-order rank algorithm on FPGA," *IEEE Transactions on Circuits and Systems II: Express Briefs*, vol. 66, no. 12, pp. 2067–2071, Dec. 2019, doi: 10.1109/TCSSII.2019.2909013.

[9] J. Jeon, S. Cho, X. Tong, and S. Lee, "Intrinsic image decomposition using structure-texture separation and surface normals," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, vol. 8695 LNCS, no. PART 7, Springer International Publishing, 2014, pp. 218–233, doi: 10.1007/978-3-319-10584-0_15.

[10] N. Silberman, D. Hoiem, P. Kohli, and R. Fergus, "Indoor segmentation and support inference from RGBD images," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, vol. 7576 LNCS, no. PART 5, Springer Berlin Heidelberg, 2012, pp. 746–760, doi: 10.1007/978-3-642-33715-4_54.

[11] D. Scharstein and R. Szeliski, "A taxonomy and evaluation of dense two-frame stereo correspondence algorithms," *International Journal of Computer Vision*, vol. 47, no. 1–3, pp. 7–42, Apr. 2002, doi: 10.1023/A:1014573219977.

[12] J. Rapp and V. K. Goyal, "A few photons among many: Unmixing signal and noise for photon-efficient active imaging," *IEEE Transactions on Computational Imaging*, vol. 3, no. 3, pp. 445–459, Sep. 2017, doi: 10.1109/tci.2017.2706028.

[13] J. Tachella *et al.*, "Bayesian 3D reconstruction of complex scenes from single-photon lidar data," *SIAM Journal on Imaging Sciences*, vol. 12, no. 1, pp. 521–550, Jan. 2019, doi: 10.1137/18M1183972.

[14] J. Tachella *et al.*, "Real-time 3D reconstruction from single-photon lidar data using plug-and-play point cloud denoisers," *Nature Communications*, vol. 10, no. 1, Nov. 2019, doi: 10.1038/s41467-019-12943-7.

- [15] J. Koo, A. Halimi, and S. McLaughlin, "A Bayesian based deep unrolling algorithm for single-photon lidar systems," *IEEE Journal on Selected Topics in Signal Processing*, vol. 16, no. 4, pp. 762–774, Jun. 2022, doi: 10.1109/JSTSP.2022.3170228.
- [16] J. Tachella, Y. Altmann, M. Marquez, H. Arguello-Fuentes, J.-Y. Tourneret, and S. McLaughlin, "Bayesian 3D reconstruction of subsampled multispectral single-photon lidar signals," *IEEE Transactions on Computational Imaging*, vol. 6, pp. 208–220, 2020, doi: 10.1109/tci.2019.2945204.
- [17] J. Li, B. Wang, H. Ma, L. Gao, and H. Fu, "Visual feature extraction and tracking method based on corner flow detection," *IECE Transactions on Intelligent Systematics*, vol. 1, no. 1, pp. 3–9, May 2024, doi: 10.62762/tis.2024.136895.
- [18] F. Wang and S. Yi, "Spatio-temporal feature soft correlation concatenation aggregation structure for video action recognition networks," *IECE Transactions on Sensing, Communication, and Control*, vol. 1, no. 1, pp. 60–71, Oct. 2024, doi: 10.62762/tsc.2024.212751.
- [19] Z. Sun, D. B. Lindell, O. Solgaard, and G. Wetzstein, "SPADnet: deep RGB-SPAD sensor fusion assisted by monocular depth estimation," *Optics Express*, vol. 28, no. 10, p. 14948, May 2020, doi: 10.1364/oe.392386.
- [20] X. Zhao, X. Jiang, A. Han, T. Mao, W. He, and Q. Chen, "Photon-efficient 3D reconstruction employing a edge enhancement method," *Optics Express*, vol. 30, no. 2, p. 1555, Jan. 2022, doi: 10.1364/oe.446369.
- [21] J. Peng, Z. Xiong, X. Huang, Z. P. Li, D. Liu, and F. Xu, "Photon-efficient 3D imaging with a non-local neural network," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, vol. 12351 LNCS, Springer International Publishing, 2020, pp. 225–241, doi: 10.1007/978-3-030-58539-6_14.
- [22] D. L. Snyder and M. I. Miller, "Self-exciting point processes," in *Random Point Processes in Time and Space*, Springer New York, 1991, pp. 287–340, doi: 10.1007/978-1-4612-3166-0_6.
- [23] W. Shi and R. Rajkumar, "Point-GNN: Graph neural network for 3D object detection in a point cloud," in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, Jun. 2020, pp. 1708–1716, doi: 10.1109/CVPR42600.2020.00178.
- [24] Z. Zang and D. D. Uei Li, "Object classification through heterogeneous fog with a fast data-driven algorithm using a low-cost single-photon avalanche diode array," *Optics Express*, vol. 32, no. 19, p. 33294, Sep. 2024, doi: 10.1364/OE.527244.
- [25] W. Zhang and Q. Hong, "Modeling brain functional networks using graph neural networks: A review and clinical application," *IECE Transactions on Intelligent Systematics*, vol. 1, no. 2, pp. 58–68, Sep. 2024, doi: 10.62762/tis.2024.680959.

BIOGRAPHIES OF AUTHORS



Yu Zhang    received his B.Eng. degree in electrical and electronic engineering from the University of Kent, United Kingdom, in 2024. He is currently pursuing the M.S. degree in electrical engineering at the Viterbi School of Engineering, University of Southern California, Los Angeles, United States, since January 2025. He can be contacted at email: yz324@usc.edu.



Yiming Zheng    received a B.Eng. degree in rail transit signal and control from Beijing Union University, Beijing, China, and a bachelor's degree in management in technological systems from Ural State University of Railway Transport, Yekaterinburg, Russia, both in 2023. He is currently pursuing the M.S. degree in electrical engineering at the Viterbi School of Engineering, University of Southern California, Los Angeles, United States, since January 2025. He can be contacted at email: yzheng52@usc.edu.