

End-to-End Learning Framework for Space Optical Communications in Non-Differentiable Poisson Channel

Abdelrahman Elfikky, Morteza Soltani, *Member, IEEE*, Zouheir Rezki, *Senior Member, IEEE*

Abstract—This letter studies an autoencoder (AE) model for point-to-point space optical communications (SOC) in the presence of an intensity modulation and direct detection setting wherein the optical link is modeled by a Poisson channel model. While the Poisson channel is a realistic channel model in SOC, its non-differentiable nature poses challenges when applied to deep learning. In this letter, a novel non-gradient-based optimization framework has been applied to estimating the channel gradient, addressing the non-differentiability of Poisson channels. Furthermore, our AE incorporates normalization layers in both encoders and decoders for input data standardization and efficient training convergence. The incorporation of the hyper-tuning optimization algorithm alongside the AE enhances the performance of the standard AE. This improvement stems from the effective minimization of the error cost function through gradient descent. The numerical results demonstrate that the proposed AE outperforms both state-of-the-art learning frameworks and model-based schemes in bit error rate (BER) performance within the Poisson channel.

Index Terms—Free space optics, end-to-end learning, Poisson channel, non-differentiable channels.

I. INTRODUCTION

A. Background:

Optical wireless communication (OWC), in comparison to traditional RF technology, offers several advantages, such as higher data rates, enhanced cost-effectiveness, and increased frequency availability [1]. OWC can be integrated into space optical communications (SOC), enabling efficient communication between satellites and ground stations [2]. In SOC, the laser beam undergoes collection by the receiving telescope and is subsequently directed towards the detector after covering a specific distance [2]. Unlike OWC, the laser transmitter has high photon capability to facilitate effective downlink and uplink communication in SOC. Furthermore, laser transmitters in SOC necessitate specific characteristics, including a narrow bandwidth and low modulation rates [2].

B. Related state-of-the-art

Investigations by researchers are ongoing to integrate machine learning frameworks into SOC. This integration has

brought about notable advancements in areas such as channel equalization, modulation classification, coding, and decoding [3]–[5]. Auto-encoders (AEs) have emerged as highly effective end-to-end (E2E) solutions, particularly in point-to-point SOC scenarios with differentiable channels. The study in [3] showcased the ability of AEs to optimize both transmitter and receiver components, achieving satisfactory bit error rate (BER) performance in point-to-point channels. The authors in [5] designed an AE model for MAC channel communication links between two GEO satellites and a common receiver in a ground station. Furthermore, in [6], [7], an AE model employing multi-decoder structure is designed for point-to-point communications in weak turbulence fading channels in SOC. In the study conducted by [4], the authors implement the standard AE within a Log-normal fading channel to address weak turbulence in Optical Wireless Communication (OWC).

When training AEs, both the channel model and all the layers of the AE must be differentiable. This poses a challenge, as some channel models are non-differentiable. In SOC, the received optical signal is often very weak, leading to the consideration of photon counting statistics. The Poisson distribution accurately represents the probability distribution of the number of photon detections in a given time period. One downside of this accuracy is that it cannot be implemented as an AE channel on its own due to its non-differentiability, as the Poisson channel distribution is a probability mass function and hence discrete. In [8], instead of using the Poisson distribution directly, the authors approximate it with a differentiable Gaussian distribution. However, this approach is limited to high optical power situations and lacks generality across different power ranges.

C. Challenge and Contribution

The Poisson model demonstrates proficiency in SOC settings by aptly characterizing the statistical patterns of discrete photon arrival, providing valuable insights into the likelihood of observing a particular photon count within a specified time frame. In contrast, conventional continuous models may fail to accurately capture the photon counting process which is by nature, discrete. However, a significant portion of deep learning (DL) literature tends to avoid considering the Poisson channel [4]–[6] due to its non-differentiable nature, making it impractical for calculating gradients during the backpropagation (BP) process. Hence, the challenge in this problem lies in utilizing the Poisson channel within SOC without

This research received support from the National Science Foundation with Grant No. 2114779. (Corresponding author: Abdelrahman Elfikky)

Abdelrahman Elfikky and Zouheir Rezki are with the Department of Electrical and Computer Engineering, University of California, Santa Cruz, California, USA (e-mail: {afikky,zrezki}@ucsc.edu).

Morteza Soltani is with the Qualcomm Inc., San Diego, California, USA (e-mail: msoltani@qti.qualcomm.com).

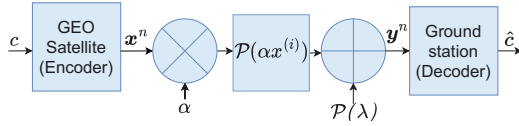


Fig. 1: System model for point-to-point SOC over the discrete-time memoryless Poisson channel.

approximating to Gaussian models, as well as addressing the BP problem. Our primary contributions address these challenges in the following manner:

- 1) We presented an autoencoder (AE) model that avoids approximation of the discrete Poisson channel during forward propagation. We employ a covariance matrix adaptation evolution strategy (CMA-ES) integrated with the proposed AE to estimate the gradients of the Poisson channel in an efficient manner.
- 2) The AE is based on batch normalization (BN) layers in the encoder, and layer normalization (LN) in the decoder. Normalization layers act as a form of regularization and ensures that the network has a consistent distribution of data during training.
- 3) The proposed AE outperformed the standard AE when the CMA-ES algorithm was applied alone in terms of BER. Additionally, the performance surpassed that of the standard AE with a non-gradient descent approach based on particle swarm optimization (PSO). The AE model also showed significant BER improvement compared to the convolutional codes and uncoded modulation schemes.

II. SYSTEM MODEL

We model a geostationary satellite with a laser transmitter, and the detector is located at the ground station as shown in Fig. 1. The received optical power at the photodetector is weak within the SOC and individual photons gain prominence. Accordingly, the SOC channel is represented using the discrete Poisson channel, which accurately depicts discrete photon arrival, showing the likelihood of observing a certain photon count in a timeframe. On the other hand, continuous models do not capture this discrete nature. The emitted light intensity follows the positivity, and peak power constraint $0 \leq x^{(i)} \leq A$, for $i \in [1, n]$. The average optical power constraint $\mathbb{E}[x] \leq \mathcal{E}$. The discrete-time memoryless optical wireless channel employing intensity modulation and direct detection (IM-DD) is modeled by a Poisson distribution as follows [1]:

$$p_{Y|X}(y^{(i)} | x^{(i)}) = e^{-(\alpha x^{(i)} + \lambda)} \frac{(\alpha x^{(i)} + \lambda)^{y^{(i)}}}{y^{(i)}!}, \quad (1)$$

where $y^{(i)} \in \mathbb{N}$ is the channel output, α is a scalar value that represents the channel gain, and $\lambda \geq 0$ is the dark current rate of the photodetector. We note that the first random variable (rv) is Poisson-distributed with probability law $p(\alpha x^{(i)})$, and the second rv is also Poisson-distributed with probability law $p(\lambda)$. Since these two rvs are independent, then the output is also another rv whose probability law is $p(\alpha x^{(i)} + \lambda)$.

While the Poisson channel is an accurate model in SOC, it faces a challenge related to non-differentiability, which can complicate its implementation in DL modes. In the following section, we elaborate on how to address the non-differentiable issue in DL models.

III. PROPOSED AE MODEL OVER THE POISSON CHANNEL

In this section, we describe the structure of the proposed AE, including the normalization layers. Then, we discuss how we utilize a CMA-ES in conjunction with the proposed AE to effectively estimate the gradients of the Poisson channel.

1. The proposed AE: First, the AE can be described as an unsupervised neural network (NN) that auto-learns how to compress the data efficiently via an encoding process. In addition to compressing data, the AE learns how to recreate the original data from the compressed form. The AE system can be expressed by the pair (k, n) , where k and n are the number of message bits and the codeword length, respectively. The channel code rate is described as $R = k/n$. As shown in Fig. 2, the proposed AE consists of both a transmitter and a receiver NN, which are jointly optimized to streamline the learning process. AE can learn to perform both encoding and decoding operations in an end-to-end manner. This joint optimization allows the AE to develop representations (or encoded symbols) that are most suitable for transmission over the channel, leading to potentially better performance than traditional schemes. A message c_i is taken from $\{1, \dots, M\}$, where $M = 2^k$. The vector $\mathbf{1}_c$ used as input is the one-hot encoding of c . We perform one-hot encoding on the input message to ensure that the model is not biased toward any specific value. The input vector is then passed through the transmitter NN and encoded into the vector $\mathbf{x}^n$ of length n , which is used as input to the channel. After $\mathbf{x}^n$ passes through the channel, it is distorted into the noisy signal $\mathbf{y}^n$ as described by (2) and reaches the ground station. There, the receiver NN outputs its reconstruction of the original one-hot encoded vector, denoted as $\hat{\mathbf{1}}_c$. Unlike the standard AE model, we introduce normalization layers in between fully connected (FC) layers to reduce the effect of poor gradient exploding and increase the speed of convergence during training. The encoder utilizes BN layers, which is applied across each batch. We utilize LN at the decoder, treating each sample independently. At the last layer on the decoder NN, we apply a softmax activation function to the decoder output vector $\mathbf{d} \in \mathbb{R}^M$ to determine the most likely value of the original message.

2. Non-differentiable Poisson channel: The Poisson channel model described in (2) is not differentiable, so using the BP algorithm along with gradient descent to tune the AE parameters becomes unfeasible. Within the proposed AE, the CMA-ES algorithm is employed in the BP algorithm to estimate the local gradient of the Poisson channel. To fix the non-differentiability of the Poisson channel, we set the gradient of the channel to a constant, denoted as J , during the BP. This constant J can be considered as a hyper-parameter and thus can be tuned using hyper-tuning optimization algorithms. CMA-ES excels over other ES methods, particularly in non-separable problems with larger search spaces or those requiring numerous function evaluations [9].

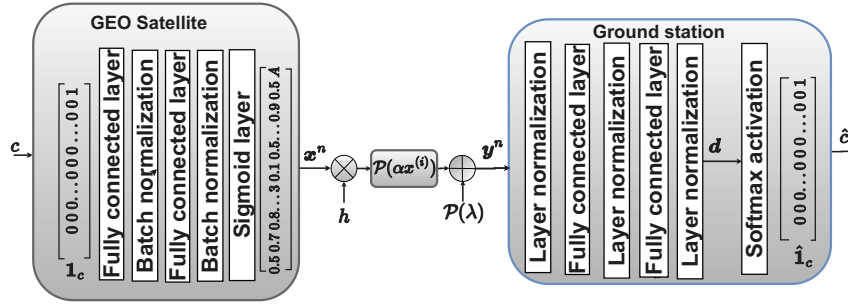


Fig. 2: The proposed AE structure.

3. CMA-ES process: In CMA-ES, candidate solutions (often denoted as w) are produced in a stochastic manner in each iteration, based on the existing candidate solutions from the previous iteration. Subsequently, a selection process is conducted to determine which candidate solutions will serve as parents in the ensuing iteration, with the selection criteria being their value of the objective function, denoted as $f(w)$. As this process continues across successive iterations, the generated candidate solutions progressively improve the objective function [10].

3.1 Candidate solutions: The normal probability distribution responsible for generating new candidate solutions, given the distribution parameters such as mean, variances, and covariances, represents the maximum entropy probability distribution over $\mathbb{R}^p$, where p denotes the number of parameters in the candidate solutions. In other words, this distribution reflects the sample distribution with the least amount of prior information embedded into it. During the k^{th} iteration, the process initiates by sampling $\beta > 1$ candidate solutions $w_i \in \mathbb{R}^p$, where $i = [1, \dots, \beta]$, from a multivariate normal distribution $\mathcal{N}(m_k, \sigma_k^2 \Sigma_k)$, where $m_k \in \mathbb{R}^p$ is the distribution mean and current solution to the optimization problem, $\sigma_k > 0$ is the step-size, and $\Sigma_k \in \mathbb{R}^{p \times p}$ is a symmetric and positive-definite covariance matrix. The candidate solutions w_i are evaluated using the objective function $f: \mathbb{R}^p \rightarrow \mathbb{R}$ and are subsequently sorted as follows [10]:

$$\{w_{f_j} \mid f(w_{f_1}) \leq \dots \leq f(w_{f_\mu}) \leq f(w_{f_{\mu+1}}) \leq \dots\}, \quad (2)$$

where $j = 1 \dots \beta$, and $\mu \leq \beta/2$ is the number of best candidate solutions selected at each iteration.

3.2 Parameters update: Now, we discuss how the mean, step size, and covariance matrix, respectively, are calculated at each iteration. These are the three main distribution parameters that will be updated in each iteration. First, the mean is updated as follows:

$$m_{k+1} = \sum_{r=1}^{\mu} \theta_r w_{f_r}. \quad (3)$$

In this context, the positive recombination weights $\theta_1 \geq \theta_2 \geq \dots \geq \theta_\mu > 0$ are selected such that their sum equals one. Conventionally, these weights are determined to satisfy the condition $1/\sum_{r=1}^{\mu} \theta_r^2 \approx \beta/4$ [9]. Next, The step-size σ_k is updated using cumulative step-size adaptation, sometimes also denoted as path length control [9], i.e.,

$$\sigma_{k+1} = \sigma_k \exp \left(\frac{c_\sigma}{d_\sigma} \left(\frac{\|p_\sigma\|}{E[\|\mathcal{N}(0, I)\|]} - 1 \right) \right), \quad (4)$$

where $c_\sigma^{-1} \approx p/3$ is the backward time horizon for the evolution path p_σ and larger than one, d_σ is the damping parameter, and E denotes the expectation. Finally, the covariance matrix is updated as follows [10]:

$$\Sigma_{k+1} = \left(1 - p^2 - \frac{\mu_w}{p^2} + c_s \right) \Sigma_k + c_1 p_c p_c^T + c_\mu \sum_{r=1}^{\mu} \theta_i \frac{w_{r:\beta} - m_k}{\sigma_k} \left(\frac{w_{r:\beta} - m_k}{\sigma_k} \right)^T, \quad (5)$$

where $\mu_\theta = 1/(\sum_{i=1}^{\mu} \theta_i^2)$ is the variance effective selection mass, and $c_s = \left(1 - \mathbf{1}_{[0, 1.5\sqrt{p}]}(\|p_\sigma\|) \right) c_1 c_c (2 - c_c)$. The backward time horizon for the evolution path p_c is indicated as $c_c^{-1} \approx n/4$, $c_1 \approx 2/n^2$ represents the learning rate for the rank-one update of the covariance matrix, and $c_\mu \approx \mu_w/n^2$ is the learning rate for the rank- μ update of the covariance matrix, ensuring it does not exceed $1 - c_1$. The indicator function $\mathbf{1}_{[0, 1.5\sqrt{p}]}(\|p_\sigma\|)$, yields a value of one if and only if $\|p_\sigma\| \leq 1.5\sqrt{p}$. Algorithm 1 summarizes the CMA-ES algorithm utilized for the hyper-tuning optimization problem.

4. Channel gradient estimation: In the proposed AE, the training process is repeated multiple times. At each time, candidate solutions for the channel gradient J are generated at the training stage of the AE. Next, each candidate solution is evaluated based on the value of the objective function f . The candidate solutions are sorted, and the best candidate solutions are chosen. Since f no longer needs to be differentiable, we use the BER produced by the model as an objective function. Then, the distribution parameters are updated for an adequate number of iterations or until convergence. At the last iteration, the optimal channel gradient value is set as the final mean of the multivariate normal distribution.

IV. NUMERICAL RESULTS

In this section, we evaluate the performance of our proposed AE with normalization layers and gradient fix against several channel coding schemes for code rate $R = \frac{1}{2}$ across the Poisson channel. The proposed AE presented in Fig. 2 is trained over 25 epochs with 8 million training samples. In the proposed AE, the CMA-ES algorithm is utilized in the BP process to estimate the local gradient of the Poisson channel. The CMA-ES algorithm is employed not only for optimizing the channel gradient of the Poisson channel but also for optimizing the learning rate, peak intensity, and batch size. Without gradient problem resolution, the default action

Algorithm 1 CMA-ES

Require: number of parameters p , number of iterations v
number of candidate solutions β , number of solutions
selected at each iteration μ , objective function f , and
recombination weights $\theta_1, \theta_2, \dots, \theta_\mu$.

Ensure: minimize $f(\mathbf{w}_i) \forall i \in \{1, 2, \dots, \mu\}$

```

1:  $\mathbf{m}_0 \leftarrow$  initialize mean vector.
2:  $\sigma_0 \leftarrow$  initialize step size.
3:  $\Sigma_0 \leftarrow$  initialize covariance matrix.
4: for  $k \leftarrow 1$  to  $v$  do
5:   for  $i \leftarrow 1$  to  $\mu$  do
6:      $\mathbf{w}_i \leftarrow$ 
       sample_multivariate_normal( $\mathbf{m}_{k-1}, \sigma_{k-1}^2 \Sigma_{k-1}$ )
7:      $f_i = f(\mathbf{w}_i)$ 
8:   end for
9:    $\mathbf{w}_{1..\mu} \leftarrow \mathbf{w}_{s(1)..s(\mu)}$  where  $s(i) = \text{argsort}(f_{1..\mu}, i)$ 
10:   $\mathbf{m}_k \leftarrow$  update_m( $\mathbf{w}_1, \dots, \mathbf{w}_\mu$ )
11:   $\mathbf{p}_{\sigma k} \leftarrow$ 
       update_ps( $\mathbf{p}_{\sigma k-1}, \sigma_{k-1}^{-1} \Sigma_{k-1}^{-1/2} (\mathbf{m}_k - \mathbf{m}_{k-1})$ )
12:   $\mathbf{p}_c \leftarrow$  update_pc( $\mathbf{p}_c, \sigma^{-1} (\mathbf{m}_k - \mathbf{m}_{k-1}), \|\mathbf{p}_{\sigma k-1}\|$ )
13:   $\Sigma_k \leftarrow$ 
       update_cov( $\Sigma_{k-1}, \mathbf{p}_{\sigma k}, (\mathbf{w}_1 - \mathbf{m}_{k-1}) / \sigma_{k-1}, \dots,$ 
        $(\mathbf{w}_\mu - \mathbf{m}_{k-1}) / \sigma_{k-1}$ )
14:   $\sigma_k \leftarrow$  update_step( $\sigma_{k-1}, \|\mathbf{p}_{\sigma k}\|$ )
15: end for

```

is that the gradients are set to zero. Consequently, in the BP algorithm, the encoder parameters remain unchanged. Em-

ploying gradient approximation, such as utilizing the CMA-ES algorithm, resolves this issue. In the testing stage, we utilize 5 million samples and the optimal learnable parameters proposed from the training stage. In Fig. 3a, the proposed AE

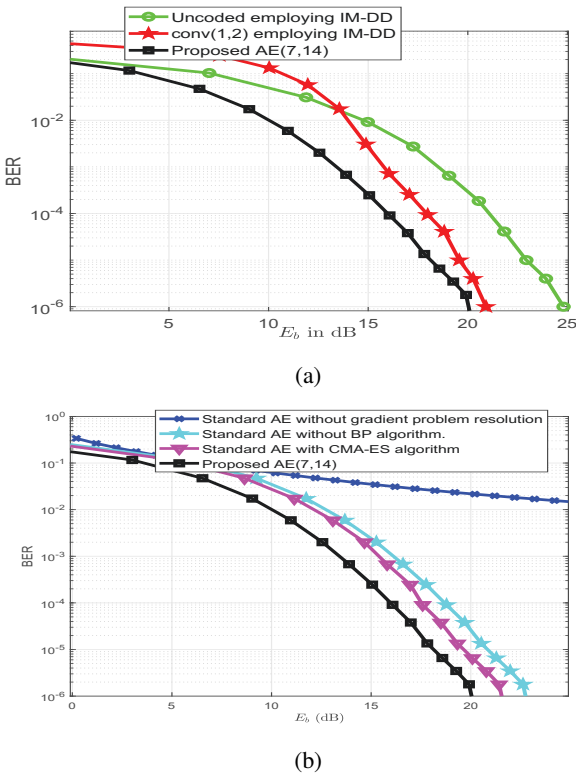


Fig. 3: BER versus E_b for $\lambda = 1$ of the proposed AE with (a) model-based frameworks (b) learning frameworks.

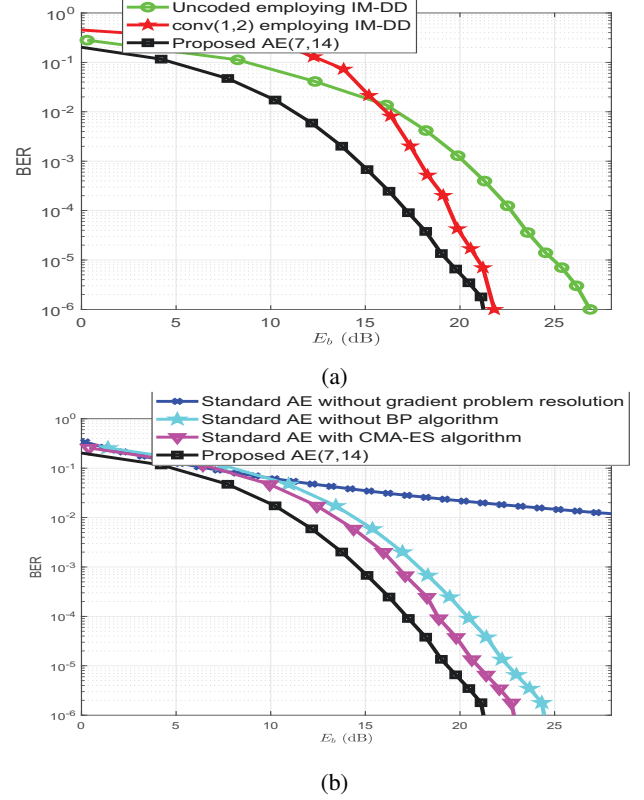


Fig. 4: BER versus E_b for $\lambda = 2$ of the proposed AE with (a) model-based frameworks (b) learning frameworks.

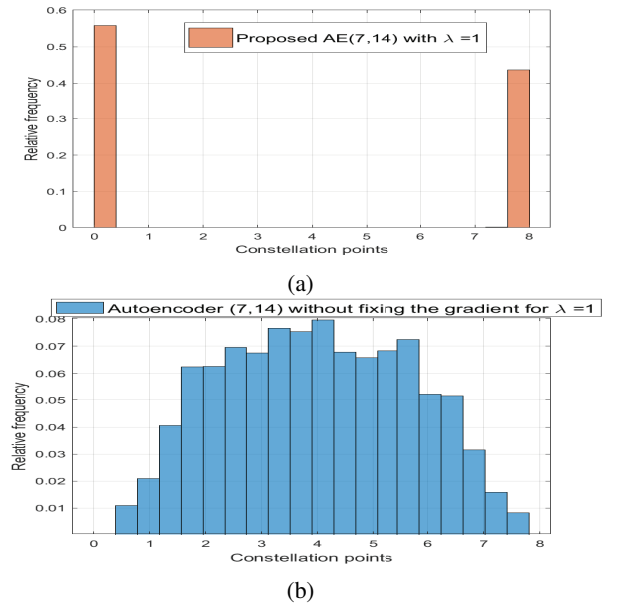


Fig. 5: The constellation points versus the relative frequency: (a) Proposed AE with CMA-ES algorithm, (b) AE without fixing the channel gradient.

is compared to the uncoded IM/DD and convolutional codes examined at $\lambda = 1$. These results show that the proposed AE achieves better performance than model-based approaches across the Poisson channel. In Fig. 3a, at BER 10^{-5} , the proposed AE has an improvement of about 1.5 dB over convolutional codes. We note that the AE's BER performance is 1.2 dB greater than that of the convolution codes. In Fig. 3b, the proposed AE is compared to the standard AE(7, 14) in [4], [8] with and without gradient problem resolution examined at $\lambda = 1$. Both references utilize the standard AE structure initiated by [3]. In Fig. 3b, at BER 10^{-5} , the proposed AE is better by about 1.5 dB over the standard AE in [4] with the CMA-ES algorithm using $\lambda = 1$. Also, at BER 10^{-6} , the proposed AE is better than standard AE without BP algorithm by 2.3 dB. Avoiding the BP algorithm is employed using PSO. Obviously, the standard AE without channel gradient problem resolution of the Poisson channel has the worst BER performance as the encoder module weights are not optimized. The hyper-tuning parameters, optimized by the CMA-ES algorithm for $\lambda = 1$, are as follows: $A = 7.3$, learning rate = 0.00045, batch size = 32, and Poisson channel gradient = 1.23.

In Fig. 4a, the proposed AE is compared to the uncoded IM/DD and convolutional codes examined at $\lambda = 2$. The proposed AE is better by at least 1 dB over all other encoding schemes tested at $\lambda = 2$. At BER 10^{-6} , the proposed AE outperforms the uncoded modulations by 5.3 dB. In Fig. 4b the proposed AE is compared to the standard AE(7, 14) with and without gradient problem resolution examined at $\lambda = 2$. The proposed AE has an improvement of approximately 1.3 dB over the standard AE with CMA-ES algorithm at BER 10^{-5} . In addition, the standard AE has worse performance by 16 dB compared with the proposed AE. For $\lambda = 2$, the optimized hyper-tuning parameters are $A = 7.8$, learning rate = 0.00012, batch size = 32, and Poisson channel gradient = 1.8. To evaluate training loss, we utilize the cross entropy loss and for updating the weights we utilize Adam optimizer.

Through these baselines, the proposed AE is shown to be more effective across the Poisson channel than both the model and learning-based frameworks. The better performance of the proposed AE over the standard AE proves that introducing normalization layers into the AE model greatly improves the BER performance, with the advantage of providing consistent distribution to the AE model. Moreover, the proposed AE incorporates the CMA-ES algorithm. demonstrates enhancements over the standard AE, when compared with the standard AE based on PSO. CMA-ES dynamically adjusts its step sizes and covariance matrix throughout the optimization process, enabling it to improve the objective function. This adaptive capability allows CMA-ES to efficiently explore complex search spaces without needing lots of parameter tuning. In contrast, PSO often demands precise parameter tuning to attain satisfactory performance across diverse problem sets [9], [10]. The testing time of the Proposed AE and the standard AE is 17766s, 18789s.

As shown in Fig. 5a, upon calculating the channel gradient for the Poisson channel, the histogram illustrates that the modulation utilized by the proposed AE resembles on-off keying

modulation at its peak intensity $A = 8$. Conversely, when employing the AE without estimating the channel gradient as illustrated in Fig. 5b, the histogram displays a randomly distributed pattern of the modulated signal from the AE. This randomness substantiates the notable degradation in BER performance of the AE when the gradient of the Poisson channel is not estimated. The incorporation of normalization in the proposed AE layers speeds up training and allows for a reduction in the number of neurons per layer in the conventional AE. The testing time for both the Proposed AE and the standard AE is 17,766s and 18,789s, respectively. This indicates a decrease in testing time of around 5.45% between the Proposed AE and the standard AE.

V. CONCLUSION

This letter explores the utilization of an AE model within a point-to-point SOC scenario, considering the impact of a practical channel model referred to as the Poisson channel. While the Poisson channel model effectively characterizes the SOC, its non-differentiable attributes pose challenges for DL models. A novel non-gradient-based optimization framework has been employed to estimate the channel gradient, effectively tackling the non-differentiable nature of Poisson channels. Moreover, our AE integrates normalization layers in both the encoding and decoding modules. The numerical results demonstrate that the proposed AE outperforms both state-of-the-art learning frameworks and model-based schemes in BER performance within the Poisson channel. Our model presents a promising solution to encourage researchers to incorporate the Poisson channel into their deep learning models without relying on approximations or transformations.

REFERENCES

- [1] A. Chaaban, Z. Rezki, and M.-S. Alouini, "On the capacity of intensity-modulation direct-detection Gaussian optical wireless communication channels: A tutorial," *IEEE Commun. Surveys Tuts.*, vol. 24, pp. 455–491, Firstquarter 2022.
- [2] H. Kaushal and G. Kaddoum, "Optical communication in space challenges and mitigation techniques," *IEEE Commun. Surveys Tuts.*, vol. 19, pp. 57–96, Firstquarter 2017.
- [3] T. O'shea and J. Hoydis, "An introduction to deep learning for the physical layer," *IEEE Trans. Cogn. Commun. Netw.*, vol. 3, pp. 563–575, December 2017.
- [4] Z.-R. Zhu, J. Zhang, R.-H. Chen, and H.-Y. Yu, "Autoencoder-based transceiver design for owc systems in Log-normal fading channel," *IEEE Photon. J.*, vol. 11, pp. 1–12, Aug. 2019.
- [5] A. Elfikky, M. Soltani, and Z. Rezki, "Learning-based autoencoder for multiple access and interference channels in space optical communications," *IEEE Comm. Lett.*, vol. 27, no. 10, pp. 2662–2666, 2023.
- [6] A. E.-R. A. El-Fikky and Z. Rezki, "On the performance of autoencoder-based space optical communications," in *Proc. IEEE GLOBECOM*, Rio de Janeiro, Brazil, pp. 1466–1471, Dec. 2022.
- [7] A. Elfikky and Z. Rezki, "Symbol detection and channel estimation for space optical communications using neural network and autoencoder," *IEEE Transactions on Machine Learning in Communications and Networking*, vol. 2, pp. 110–128, 2024.
- [8] L.-H. Si-Ma, Z.-R. Zhu, and H.-Y. Yu, "Model-aware end-to-end learning for siso optical wireless communication over poisson channel," *IEEE Photon. J.*, vol. 12, no. 6, pp. 1–15, 2020.
- [9] I. Loshchilov, "CMA-ES with restarts for solving cec 2013 benchmark problems," in *2013 IEEE Congress on Evolutionary Computation*, pp. 369–376, Ieee, 2013.
- [10] Y. Wu, X. Peng, H. Wang, Y. Jin, and D. Xu, "Cooperative coevolutionary CMA-ES with landscape-aware grouping in noisy environments," *IEEE Trans. Evol. Comp.*, vol. 27, no. 3, pp. 686–700, 2023.