

PHƯƠNG PHÁP PHÁT HIỆN DỮ LIỆU DỰA TRÊN HỌC SÂU CHO HỆ THỐNG TRUYỀN THÔNG SỢI QUANG CÓ NHIỀU PHI TUYẾN

A DEEP LEARNING-BASED DATA DETECTOR FOR FIBER COMMUNICATION SYSTEM IN THE PRESENCE OF NONLINEAR NOISE

Nguyễn Duy Nhật Viễn^{1*}, Vương Quang Phước^{1,2}

¹Trường Đại học Bách khoa - Đại học Đà Nẵng, Việt Nam

²Trường Đại học Khoa học - Đại học Huế, Việt Nam

*Tác giả liên hệ / Corresponding author: ndnvien@dut.udn.vn

(Nhận bài / Received: 25/12/2024; Sửa bài / Revised: 19/02/2025; Chấp nhận đăng / Accepted: 20/02/2025)

DOI: 10.31130/ud-jst.2025.577

Tóm tắt - Trong bài báo này, nhóm tác giả trình bày phương pháp phát hiện dữ liệu cho hệ thống truyền thông sợi quang có nhiều phi tuyến, và mô hình hóa hệ thống này như một mạng nơ-ron học sâu từ đầu đến cuối, bao gồm bộ phát, kênh truyền và bộ thu. Trên cơ sở đó, nhóm tác giả đề xuất kiến trúc mạng nơ-ron autoencoder cho hệ thống truyền thông sợi quang đường dài để phát hiện dữ liệu. Thông qua các mô phỏng, nhóm tác giả đã so sánh, đánh giá phương pháp đề xuất so với phương pháp phát hiện ML (Maximum likelihood) và autoencoder khác trong trường hợp không có và có nhiễu phi tuyến với tín hiệu 16-QAM. Kết quả mô phỏng cho thấy, tính hiệu quả của phương pháp đề xuất về độ chính xác lẫn độ phức tạp. Nghiên cứu này sẽ là cơ sở để hướng tới việc tối ưu hóa hệ thống truyền thông sợi quang dựa trên học sâu từ đầu đến cuối.

Từ khóa - Hệ thống thông tin sợi quang; bộ tự mã hóa; phát hiện; mạng nơ-ron; học máy; học sâu

1. Giới thiệu

Hệ thống truyền thông sợi quang sử dụng kỹ thuật điều chế biên độ cầu phương đa mức M-QAM (M-ary Quadrature Amplitude Modulation) đã đem lại sự gia tăng đáng kể về tốc độ dữ liệu so với tín hiệu nhị phân truyền thống [1]. Tuy nhiên, các hệ thống M-QAM nhạy cảm với các suy giảm kênh, do việc mã hóa dữ liệu vào pha của tín hiệu quang và chòm sao có mật độ cao. Sự tương tác giữa tín hiệu và nhiễu từ bộ khuếch đại quang thông qua hiệu ứng phi tuyến Kerr dẫn đến nhiễu pha phi tuyến [2]. Các kỹ thuật thông thường để xử lý nhiễu phi tuyến bao gồm thiết kế bộ phát hiện cải tiến và các định dạng điều chế tối ưu [3-5].

Trong [6], các tác giả đã đề xuất thuật toán phát hiện dữ liệu dựa trên tối đa khả năng ML (maximum likelihood) dưới dạng công thức tường minh cho các kênh truyền dẫn quang đường dài có nhiễu phi tuyến. Trong thời gian gần đây, kỹ thuật học sâu được khai thác để giải quyết các bài toán trong thông tin sợi quang [7-9]. Các công trình [10, 11] đề xuất phương pháp thiết kế chòm sao và bộ phát hiện dựa vào học sâu cho các hệ thống truyền dẫn hữu tuyến cũng như vô tuyến.

Autoencoder (AE) là một loại mạng nơ-ron nhân tạo

Abstract - In this paper, the authors introduce a data detection method for optical fiber communication systems affected by nonlinear noise and model the entire system as an end-to-end deep learning (DL) neural network, incorporating the transmitter, transmission channel, and receiver. Based on this framework, the authors propose an autoencoder-based neural network architecture for long-haul optical communication systems to enhance data detection performance. Through simulations, the authors compare and evaluate the proposed method against Maximum Likelihood (ML) detection and other autoencoder models, considering both linear and nonlinear noise scenarios with a 16-QAM signal. The simulation results demonstrate that, the proposed approach achieves superior accuracy and computational efficiency. This study hence provides a foundation for optimizing optical fiber communication systems using end-to-end deep learning techniques.

Key words - Optical fiber communication; autoencoder; detection; neural networks; machine learning; deep learning

học sâu DL (deep learning) được thiết kế cho việc học không giám sát [12]. AE học cách biểu diễn hiệu quả dữ liệu đầu vào mà không cần dữ liệu được gán nhãn. Khả năng học các biểu diễn dữ liệu hiệu quả và thực hiện các tác vụ như nén, phát hiện bất thường và giảm nhiễu khiến chúng trở thành những ứng dụng rất giá trị trong nhiều lĩnh vực truyền thông [13].

Trên cơ sở đó, bài báo này đề xuất giải thuật phát hiện dữ liệu trên cơ sở AE nhằm giảm thiểu tỷ lệ lỗi ký hiệu SER (symbol error rate).

Các ký hiệu được sử dụng trong bài báo:

- Các chữ in đậm viết thường và viết hoa được dùng để biểu diễn vector và ma trận.

- $\mathbf{C}^{N \times M}$ mô tả ma trận phức kích thước $N \times M$ và $\mathbf{I}_M$ biểu thị ma trận đơn vị $M \times M$.

- $|\mathbf{X}|$, $\mathbf{X}^H$, và $\text{tr}(\mathbf{X})$ là trị tuyệt đối, chuyển đổi Hermitian và phép tính trace của ma trận $\mathbf{X}$.

- $E\{\cdot\}$ và $\|\cdot\|$ là kỳ vọng và phép tính chuẩn (norm) Euclidean.

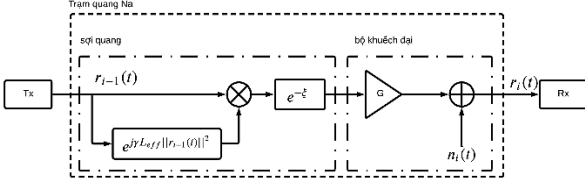
- Ma trận hay vector phức $\mathbf{z}$ ngẫu nhiên Gaussian có kỳ vọng μ và phương sai σ^2 được biểu diễn là $\mathbf{z} \sim \text{CN}(\mu, \sigma^2)$.

¹ The University of Danang - University of Science and Technology, Vietnam (Vien Nguyen-Duy-Nhat)

² Hue University - University of Science, Vietnam (Vuong Quang Phuoc)

2. Mô hình hệ thống

Xét một hệ thống truyền thông cáp quang bao gồm máy phát, máy thu và kênh truyền được mô tả trong Hình 1. Hệ thống truyền dẫn đường dài bao gồm nhiều tầng khuếch đại để bù sự tán xạ và suy hao tín hiệu. Mỗi tầng khuếch đại bao gồm một sợi quang đơn mode SMF (single mode fiber) theo sau sợi bù tán xạ DCF (dispersion compensating fiber) với các tham số phi tuyến Kerr lần lượt là γ_{SMF} và γ_{DCF} . Bộ DCF trong bài báo này được giả sử là bù tán sắc lý tưởng.



Hình 1. Mô hình hệ thống thông tin sợi quang

Công suất tín hiệu suy giảm sau mỗi khoảng đường được xác định theo công thức:

$$e^{-\alpha L} = e^{-(\alpha_{SMF} L_{SMF} + \alpha_{DCF} L_{DCF})}. \quad (1)$$

Trong đó, α_{SMF} và α_{DCF} là các hệ số suy hao; L_{SMF} và L_{DCF} là các độ dài tương ứng của SMF và DMF.

Các bộ khuếch đại khôi phục công suất tín hiệu về mức phát bởi độ lợi khuếch đại $G = e^{\xi}$. Ngoài ra, trong các bộ khuếch đại phát sinh nhiễu phát xạ tự phát khuếch đại ASE (amplified spontaneous emission) được mô hình là một quá trình nhiễu Gaussian trắng cộng. Tín hiệu băng cơ sở đưa vào tuyến quang bởi bộ phát TX được xác định bởi:

$$s(t) = \sqrt{P_{in}} \sum_{n=-\infty}^{+\infty} a_n p(t - nT). \quad (2)$$

Với $a_n = [a_n^{(x)}, a_n^{(y)}]^T \in \Omega^2$ là véc-tơ dữ liệu thứ n từ chòm sao chuẩn hóa Ω có $\mathbb{E}[a_n, a_n^H] = I$, T là chu kỳ tín hiệu, P_{in} là công suất khởi động, và $p(t)$ là xung với đỉnh bằng 1 tại thời điểm $t = 0$. Tín hiệu sau tầng khuếch đại thứ i nhận được từ phương trình Schrödinger bỏ qua tán xạ như sau:

$$r_{k+1}(t) = r_k(t) e^{\frac{j\gamma L \|r_k(t)\|^2}{K}} + n_k(t), 0 \leq k \leq K, \quad (3)$$

với $r_0(t) = s(t)$, $\gamma = \gamma_{SMF}$ là tham số phi tuyến của sợi quang, L là chiều dài tương đối của sợi quang ($L = L_{SMF}$) có hệ số suy giảm $\alpha = \alpha_{SMF}$, và $n_i(t)$ là nhiễu ASE có kỳ vọng bằng 0 và phương sai $\frac{N_{0,ASE}}{K}$, với mật độ phổ công suất $N_{0,ASE} = h\nu n_{sp}(G - 1)$, trong đó h là hằng số Planck, ν là tần số quang, n_{sp} là hệ số phát xạ tự phát và G là độ lợi bộ khuếch đại.

Để tính toán đơn giản, bỏ qua ảnh hưởng tán xạ mode phân cực (Polarization Mode Dispersion) và nhiễu pha của bộ dao động nội và giả sử đạt được đồng bộ thời gian và tần số. Nghĩa là hiệu suất đạt được có thể được xem là giới hạn dưới (lower bound) của tỷ lệ lỗi ký hiệu SER (Symbol Error Rate) của một hệ thống thực tế.

Tín hiệu sau khi đi qua K bộ khuếch đại được chuyển về miền điện, lọc với băng thông B và lấy mẫu với tốc độ lấy mẫu $\frac{1}{T}$. Giả sử bộ lọc là lý tưởng, trong miền rời rạc, ta có:

$$r_{k+1}(nT) = r_k(nT) e^{\frac{j\gamma L \|r_k(nT)\|^2}{K}} + n_k(nT), \quad (4)$$

Bằng phương pháp truy hồi, ta tính được

$$r_{k+1}(nT) = s_k(nT) e^{\frac{j\gamma L}{K} \sum_{k=1}^K \|r_k(t)\|^2} + w_k(nT), \quad (5)$$

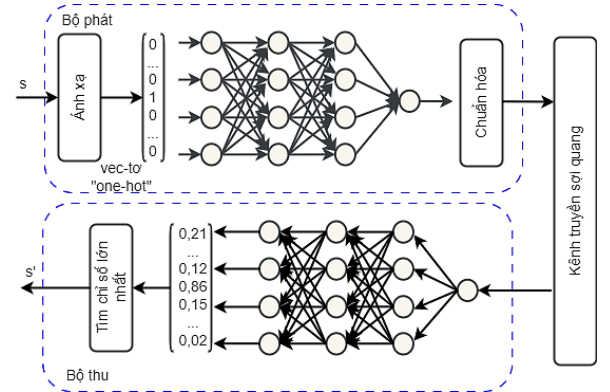
với $w_k \sim CN(0, \sigma_{ASE}^2 \mathbf{I})$, $\sigma_{ASE}^2 = BN_{0,ASE}$.

Dữ liệu có thể phát hiện bởi bộ phát hiện ML (Maximum likelihood) [6]:

$$\hat{a}_k = \arg \min_{a_k \in \Omega^2} \left\| r_k - s_k e^{\frac{j\gamma L}{K} \|s_k\|^2} \right\|^2. \quad (6)$$

3. Ứng dụng học sâu để phát hiện dữ liệu

Trong phần này, đề xuất phương pháp học sâu để phát hiện dữ liệu cho hệ thống thông tin sợi quang. Một hệ thống truyền thông nói chung, cũng như hệ thống thông tin sợi quang, có thể được xét như một bộ autoencoder [14]. Một autoencoder được mô tả như một cấu trúc mạng trí tuệ học sâu mà được đào tạo để khôi phục thông tin ngõ vào ở ngõ ra như được mô tả trong Hình 2.



Hình 2. Kiến trúc mạng nơ-ron mô tả hệ thống thông tin sợi quang

Nói cách khác, AE bao gồm 2 phần chính: phần encoder ánh xạ các ngõ vào s thành một dạng biểu diễn (mã) có số chiều thấp hơn; và phần decoder cố gắng khôi phục ngõ vào từ mã ấy.

Bên phát muốn truyền một thông điệp $m \in M = \{1, 2, \dots, M\}$ đến bên thu. Tốc độ thông tin là $R = \frac{k}{n}$, trong đó, $k = \log_2(M)$ và n là số ký tự truyền trong một thông điệp.

Quá trình mã hóa được thực hiện như sau: thông điệp m được mã hóa thành một véc-tơ "one-hot" – ký hiệu là $\mathbf{1}_m \in \mathbb{R}^M$, trong đó chỉ có thành phần thứ m có giá trị 1, còn các thành phần khác có giá trị 0. Việc mã hóa "one-hot" như vậy được sử dụng khá phổ biến trong các thuật toán học máy và là cơ sở để tối thiểu hóa tỷ lệ lỗi ký hiệu. Véc-tơ one-hot được đưa vào lớp ẩn đầu tiên của mạng nơ-ron, các lớp ẩn này gồm các nơ-ron, ma trận trọng số, véc-tơ độ lệch (bias), các hàm kích hoạt, các kết nối và các kỹ thuật chuẩn hóa.

Mạng nơ-ron bao gồm nhiều lớp ẩn, lớp ẩn sau nhận dữ liệu là ngõ ra của lớp ẩn trước và tạo ra ngõ ra $\mathbf{z}_{out} = f(\mathbf{w}^T \mathbf{z}_{in} + \mathbf{b})$, trong đó, $\mathbf{w}$ là ma trận trọng số, $\mathbf{b}$ là véc-tơ độ lệch và $f(\cdot)$ là hàm kích hoạt.

Ngõ ra được chuẩn hóa và truyền lên kênh truyền đến bộ thu. Tín hiệu nhận được tại bộ thu là y lại đưa đến mạng

ơ-nơ thu. Mạng ơ-nơ thu cũng bao gồm các thành phần cơ bản như mạng ơ-nơ phát, tuy nhiên với vai trò ngược lại. Quá trình mã hóa được tiến hành tại ơ-nơ phát thì giải mã sẽ được thực hiện tại ơ-nơ thu. Nếu gọi ngõ ra của mạng ơ-nơ thu là $f_y(s') \in [0, 1], s' \in M$, và giả sử áp dụng hàm *sigmoid* ở lớp ẩn cuối cùng rồi chuẩn hóa sao cho tổng bằng 1, ta thu được $\hat{s} = \arg \max_{s'} f_y(s')$.

Hàm mất mát sử dụng trong mạng AE là cross-entropy, xác định sự khác biệt giữa phân phối xác suất dự đoán và phân phối thực tế, được định nghĩa như sau:

$$\mathcal{L} = \frac{1}{N} \sum_{n=1}^N -u_s^{(i)} \log f_y(s')^{(i)}. \quad (7)$$

Với N là kích thước nhóm dữ liệu (batch size), $u_s^{(i)}$ là nhãn của ký hiệu thứ i .

4. Kết quả và thảo luận

Trong phần này, các kết quả mô phỏng được trình bày để đánh giá phương pháp phát hiện tín hiệu dựa trên học sâu đề xuất. Giả sử chiều dài tuyến cáp quang là $L = 5000$ km, hệ số phi tuyến bằng 0 (không có nhiễu phi tuyến) và $\gamma = 1,27$ (có xét nhiễu phi tuyến), công suất nhiễu $P_N = -21,3$ dBm, số bộ khuếch đại $K = 50$, tín hiệu điều chế 16-QAM ($M = 16$).

Bảng 1. Các tham số thiết lập mạng ơ-nơ Auto Encoder so với [13]

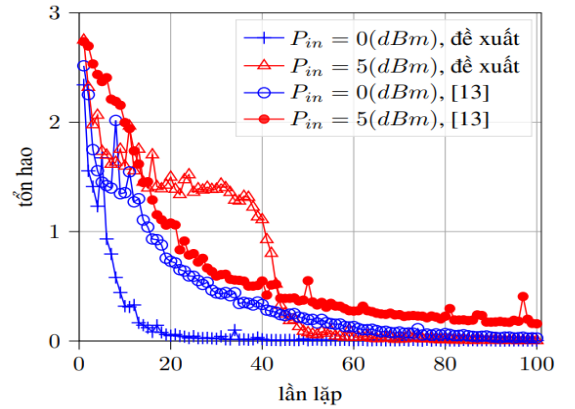
Tham số	Phương pháp [13]	Phương pháp đề xuất
Số ơ-nơ bộ phát (lớp ngõ vào, lớp ẩn, lớp ngõ ra)	$(M, (2 - 6) \times M, 7 \times 2)$	$(M, 2)$
Số ơ-nơ bộ thu (lớp ngõ vào, lớp ẩn, lớp ngõ ra)	$(2, (2 - 7) \times M, 8 \times M)$	$(2, 4M, M, 4M)$
Hàm kích hoạt lớp ẩn	$\tanh(\)$	$\text{LeakyELU}(\)$
Hàm kích hoạt lớp ngõ ra	$\text{sigmoid}(\)$	$\text{sigmoid}(\)$

Bảng 1 trình bày các tham số thiết lập mạng ơ-nơ nhân tạo AE của bài báo [13] và phương pháp đề xuất. Hàm kích hoạt ngõ ra đều là *sigmoid*() do ngõ ra là véc-tơ “one-hot”, trong khi đó, với các lớp thì phương pháp đề xuất sử dụng hàm kích hoạt *LeakyELU*() thay vì *tanh*(). Hàm kích hoạt *LeakyELU*() hạn chế điểm chết và ổn định hơn trong quá trình huấn luyện. Ngoài ra, trong Bảng 1, ta thấy rằng, thuật toán đề xuất có độ phức tạp (số ơ-nơ trong các lớp) thấp hơn so với [13], trong khi kết quả của phương pháp đề xuất cho kết quả tốt hơn vì trong [13], kiến trúc mạng có số lượng neuron khá đồng đều và dư thừa, trong khi đó, phương pháp đề xuất lựa chọn cấu trúc vừa đủ, đồng thời tạo được nút thắt trong lớp ẩn, khai thác được ưu điểm của autoencoder.

Tập dữ liệu được tạo ngẫu nhiên theo các công thức (2) và (3) với kích thước 100000 mẫu, do tập dữ liệu là khá lớn, nên không chia dữ liệu test mà chỉ chia tỷ lệ train/validation là 90%/10%.

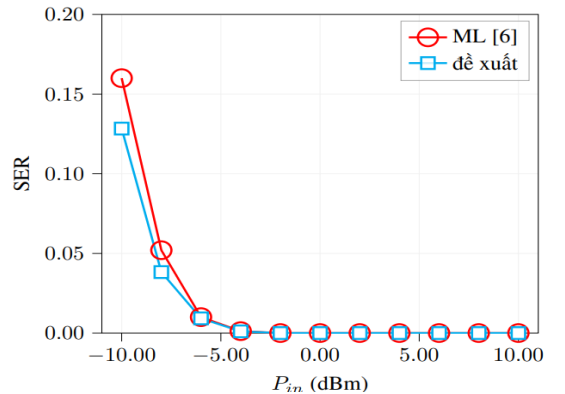
Hình 3 biểu diễn kết quả hàm mất mát trong trường hợp có nhiễu phi tuyến có $\gamma = 1,27$ với công suất phát lần lượt là $P_{in} = 0, 5$ (dBm). Theo hình vẽ, ta thấy, trường hợp công suất vào nhỏ (0 dBm) có tốc độ hội tụ nhanh hơn công

suất vào lớn. Điều này là bởi công suất phát tham gia vào hàm $\exp(\cdot)$ trong biểu thức tín hiệu thu được. Ngoài ra, so với [13], phương pháp đề xuất có giá trị mất mát bé hơn so với [13] do lớp mã hóa (encoding layer) và lớp giải mã (decoding layer) có kích thước lớn hơn, đồng thời kích thước của lớp mã hóa tạo nút thắt cần thiết để nén thông tin tốt hơn.



Hình 3. Giá trị hàm mất mát tương ứng với trường hợp $\gamma = 1,27$

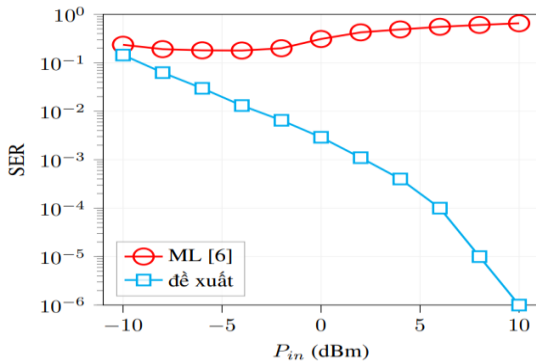
Hình 4 so sánh phương pháp đề xuất so với kỹ thuật phát hiện tín hiệu bằng ML khi không có nhiễu phi tuyến. Ta thấy rằng, kỹ thuật phát hiện ML cho kết quả tương đồng với phương pháp đề xuất. Do công suất phát không tham gia vào hàm $\exp(\cdot)$ trong biểu thức tín hiệu thu được ($\exp(0)$), nên ta dễ dàng khôi phục ký tự từ tín hiệu thu được thông qua kỹ thuật ML.



Hình 4. Tỷ số lỗi ký hiệu là hàm của công suất phát trong trường hợp không có nhiễu phi tuyến $\gamma = 0$

Trong kỹ thuật phát hiện dữ liệu, mô hình học sâu khó có thể đáp ứng được các kiểu điều chế khác nhau, vì để hỗ trợ được cho tất cả kiểu điều chế, yêu cầu phải thêm tham số điều chế như là 1 dữ liệu vào của mô hình hoặc thiết kế mô hình tổng quát (không tối ưu). Tuy nhiên, một mô hình đáp ứng được điều chế đa mức có số mức cao thì sẽ đáp ứng được cho kiểu điều chế có số mức thấp hơn. Trong thông tin sợi quang, điều chế 16-QAM được xem là kiểu điều chế có chòm sao phức tạp hơn các kiểu điều chế thông dụng như OOK, BPSK, QPSK. Hình 5 so sánh phương pháp đề xuất so với kỹ thuật phát hiện tín hiệu bằng ML [6] với tín hiệu được điều chế 16-QAM. Bằng quan sát, ta thấy rằng, kỹ thuật phát hiện ML cho kết quả tương đồng với phương pháp đề xuất khi công suất phát thấp, và phát hiện

ML có kết quả khá xấu khi công suất phát tăng (do công suất phát góp phần vào nhiễu phi tuyến trong tín hiệu thu). Trong khi đó, phương pháp đề xuất có kết quả càng tốt khi công suất phát càng cao. Nghĩa là mạng nơ-ron có thể tìm được chòm sao thích hợp trong sự hiện diện của nhiễu pha, hiểu được tín hiệu phát và học được cách xấp xỉ phân bố kênh gần giống như ML.



Hình 5. Tỷ số lỗi ký hiệu là hàm của công suất phát trong trường hợp có nhiễu phi tuyến ($\gamma = 1,27$)

5. Kết luận

Trong bài báo này, nhóm tác giả trình bày một phương pháp phát hiện dữ liệu trên cơ sở học sâu. Nhóm tác giả đã đánh giá bộ phát hiện dữ liệu bằng tỷ lệ lỗi ký hiệu theo hàm của công suất đầu vào và so sánh với các bộ phát hiện ML. Kết quả mô phỏng đã chứng minh rằng, phương pháp đề xuất có thể học từ các dữ liệu chòm sao mã hóa với nhiễu phi tuyến, tốt hơn phương pháp ML mà không cần kiến thức kênh tường minh.

Lời cảm ơn: Nghiên cứu này được tài trợ bởi Bộ Giáo dục và Đào tạo trong đề tài mã số B2024.DNA.19.

TÀI LIỆU THAM KHẢO

- [1] K. P. Ho, *Phase-Modulated Optical Communication Systems*, 1st edition, New York, USA: Springer New York, 2005.
- [2] A. Lau and J. M. Kahn, "Signal Design and Detection in Presence of Nonlinear Phase Noise", *Journal of Lightwave Technology*, vol. 25, pp. 3008-3016, 2007. <https://doi.org/10.1109/JLT.2007.905217>
- [3] F. Khan, Y. Zhou, A. Lau, and C. Lu, "Modulation format identification in heterogeneous fiber-optic networks using artificial neural networks", *Opt. Express*, vol. 20, no. 11, pp. 12422-12431, 2012. <https://doi.org/10.1364/OE.20.012422>
- [4] K. P. Ho and J. M. Kahn, "Electronic compensation technique to mitigate nonlinear phase noise", *Journal of Lightwave Technology*, vol. 22, no. 3, pp. 779-783, 2004. <https://doi.org/10.1109/JLT.2004.825792>
- [5] L. F. Mollenauer, "Distributed amplification for light wave transmission system", Oct. 221991, U.S. Patent 5,058,974.
- [6] A. S. Tan *et al.*, "An ML-based detector for optical communication in the presence of nonlinear phase noise", in *2011 IEEE International Conference on Communications (ICC)*. IEEE, 2011. <https://doi.org/10.1109/icc.2011.5962741>
- [7] N. T. Hung *et al.*, "Digital back-propagation optimization for high-baudrate single-channel optical fiber transmissions", *Optical Communications*, vol. 491, pp. 411-416, 2021. <https://doi.org/10.1016/j.optcom.2021.126913>
- [8] C. Zuo *et al.*, "Deep learning in optical metrology: a review", *Light: Science & Applications*, vol. 11, no.1, pp. 1-54, 2022. <https://doi.org/10.1038/s41377-022-00757-0>
- [9] B. Xue, S. Wei, X. Yang, Y. Ma, T. Xi, and X. Shao, "Simplified design method for optical imaging systems based on deep learning". *Applied Optics*, vol. 63, no. 28, pp. 7433-7441. 2024. <https://doi.org/10.1364/AO.530390>
- [10] T. O'Shea and J. Hoydis, "An Introduction to Deep Learning for the Physical Layer", *IEEE Transactions on Cognitive Communications and Networking*, vol. 3, pp. 563-575, 2017. <https://doi.org/10.1109/TCCN.2017.2758370>
- [11] D. Zibar, M. Piels, R. Jones, and C. G. Schäffer, "Machine learning techniques in optical communication", *Journal of Lightwave Technology*, vol. 34, pp. 1442 - 1452, 2016. <https://doi.org/10.1109/JLT.2015.2508502>
- [12] C. Catanese, A. Triki, E. Pincemin, and Y. Jaouën, "A Survey of Neural Network Applications in Fiber Nonlinearity Mitigation", in *2019 21st International Conference on Transparent Optical Networks (ICTON)*, Angers, France, 2019, pp. 1-4. <https://doi.org/10.1109/ICTON.2019.8840355>
- [13] S. Li, C. Häger, N. Garcia and H. Wymeersch, "Achievable information rates for nonlinear fiber communication via end-to-end autoencoder learning", in *2018 European Conference on Optical Communication (ECOC)*. IEEE, 2018. <https://doi.org/10.1109/ECOC.2018.8535456>
- [14] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. MIT Press, 2016, <http://www.deeplearningbook.org>
- [15] C. Pham-Quoc *et al.*, "Robust 3D beamforming for secure UAV communications by DAE", *Mobile Networks and Applications*, vol. 28, no.3, pp. 1-9. 2023, 1197-1205. <https://doi.org/10.1007/s11036-023-02130-w>.