

Robust Indoor Localization in VLC Systems Using Machine Learning: A Comparative Study under Environmental Variability

Alhusein Almahjoub^{1*}, Mustafa M. Abuali²

¹ Communications Engineering College of Electronic Technology Bani Walid, Libya

² Department of Computer and Information Technology, College of Electronic Technology, Bani Walid, Libya.

*Corresponding author: al.almahjoub@gmail.com

Received: 08-04-2026	Accepted: 16-07-2026	Published: 22-07-2026
	Copyright: © 2026 by the authors. This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (https://creativecommons.org/licenses/by/4.0/).	

Abstract:

Indoor localization has a growing demand, primarily due to its widespread applications in areas such as robot navigation, indoor navigation, virtual reality, and asset tracking. With recent advancements in solid-state devices, Visible Light Communication (VLC) has emerged as a promising solution for high-speed data communication and accurate indoor localization, harnessing existing LED lighting infrastructure. This research incorporates five machine learning approaches for location estimation using visible light commissioning Received Signal Strength (RSS) as a parameter for location estimation. The study evaluates both static and dynamic scenarios under noise-free and noisy conditions. An indoor environment is simulated with a $5 \times 5 \times 3 \text{ m}^3$ room and a 3×3 LED grid. Five Machine Learning (ML) models, namely Linear Regression (LR), Random Forest, K-nearest neighbours (KNN), Support Vector Regression (SVR), and Multi-Layer Perceptron (MLP), leverage RSS data to generate an estimation of the position of the object in a 2D environment. Evaluation parameters used to determine the performance are Root-Mean-Square Error (RMSE), training/inference time, and spatial error maps. Simulation results show superior performance of the KNN and Random Forest models in a noisy environment. In addition to that, results show that ML offers a robust alternative to analytical methods for Visible Light Positioning (VLP). Future work includes real-world validation and exploration of deep learning techniques.

Keywords: indoor localization, machine learning, received signal strength, root mean square deviation, 6G.

التوطين الداخلي المتين في أنظمة الاتصالات بالضوء المرئي (VLC) باستخدام التعلم الآلي: دراسة مقارنة تحت ظروف التغيرات البيئية

الحسين المحجوب^{1*}، مصطفى مفتاح ابو علي²

¹ قسم هندسة الاتصالات، كلية التقنية الإلكترونية، بني وليد، ليبيا.

² قسم الحاسوب وتقنية المعلومات، كلية التقنية الإلكترونية، بني وليد، ليبيا.

الملخص

يشهد مجال تحديد المواقع داخل المباني (Indoor Localization) طلبًا متزايدًا نتيجة اتساع نطاق تطبيقاته في العديد من المجالات، مثل توجيه الروبوتات، أنظمة الملاحة الداخلية، الواقع الافتراضي، وتتبع الأصول والمعدات. ومع التطورات الحديثة في تقنيات الأجهزة ذات الحالة الصلبة، ظهرت تقنية الاتصالات بالضوء المرئي (Visible Light Communication - VLC) كحل واعد لتحقيق اتصالات عالية السرعة وتحديد دقيق للموقع داخل البيئات المغلقة، وذلك من خلال الاستفادة من البنية التحتية القائمة لمصابيح LED.

تقدم هذه الدراسة تقييمًا مقارنًا لخمس خوارزميات من تقنيات التعلم الآلي (Machine Learning - ML) لتقدير الموقع اعتمادًا على شدة الإشارة المستقبلية (Received Signal Strength - RSS) في أنظمة تحديد المواقع بالضوء المرئي (Visible Light Positioning - VLP). تم تحليل أداء النماذج في سيناريوهات ثابتة ومتحركة، ضمن بيئات خالية من الضوضاء وبيئات تحتوي على ضوضاء.

تم إنشاء نموذج محاكاة لبيئة داخلية بأبعاد $5 \times 5 \times 3$ م تحتوي على شبكة إضاءة مكونة من 3×3 مصابيح LED. وشملت نماذج التعلم الآلي المستخدمة: الانحدار الخطي (Linear Regression - LR)، الغابة العشوائية (Random Forest)، خوارزمية أقرب الجيران (K-Nearest Neighbors - KNN)، انحدار آلة المتجهات الداعمة (Support Vector Regression - SVR)، والإدراك متعدد الطبقات (Multi-Layer Perceptron - MLP)، حيث تم تدريبها باستخدام بيانات RSS لتقدير موقع جسم داخل بيئة ثنائية الأبعاد.

تم تقييم أداء النماذج باستخدام مؤشرات قياس تشمل الجذر التربيعي لمتوسط مربع الخطأ (Root Mean Square Error - RMSE)، وزمن التدريب والتنفيذ، بالإضافة إلى خرائط الخطأ المكاني لتحديد دقة التوطين. أظهرت نتائج المحاكاة أن نموذجي KNN و Random Forest حققا أداءً متميزًا في البيئات التي تحتوي على ضوضاء، كما بينت النتائج أن تقنيات التعلم الآلي توفر منهجية قوية وفعالة مقارنة بالطرق التحليلية التقليدية المستخدمة في أنظمة تحديد المواقع بالضوء المرئي.

تشمل الأعمال المستقبلية التحقق التجريبي من أداء النظام في البيئات الواقعية، بالإضافة إلى دراسة استخدام تقنيات التعلم العميق (Deep Learning) لتحسين دقة وموثوقية أنظمة التوطين الداخلي.

الكلمات المفتاحية: التوطين الداخلي؛ التعلم الآلي؛ شدة الإشارة المستقبلية؛ الجذر التربيعي لمتوسط مربع الخطأ؛ الاتصالات من الجيل السادس (6G).

I. Introduction

Indoor positioning systems are gaining much attention in recent years due to their critical usage in high-precision navigation, asset tracking, real-time tracking, and many other applications. These applications are not limited to the domestic environment but are used in industries, health, and agriculture sectors as well. Global Navigation Satellite System (GNSS) can provide good accuracy in outdoor environments; however, it fails to provide an acceptable location estimation in an indoor environment primarily because of Radio Frequency (RF) attenuation, fading, and blocking [?]. To overcome location accuracy issues in an indoor environment, several technologies were utilized in literature, such as Wireless Local Area Network (WLAN), Bluetooth, Acoustics, ZigBee, and Radio-Frequency Identification (RFID); however, most of these solutions are RF-based and require additional expensive equipment to support positioning [1].

Visible light communication offers an unlicensed and virtually unlimited bandwidth for communication, and this can be used for positioning and sensing as well. With the recent advancements in the field of solid-state devices, light-emitting diodes now have a considerable bandwidth and extremely high switching rate that enables LEDs to be leveraged both for illumination and communication [2]. On the other hand, an off-the-shelf camera or a photodiode (PD) can be integrated to receive the light signals and thus can act as a receiver. With the availability of cameras and PDs in smart gadgets such as smartphones and tablets, and the ubiquitous deployment of LEDs in buildings, infrastructure for VLC and positioning is already available. Non-interfering nature of visible light, along with the inherent security, as the light

cannot pass through an opaque structure, makes it an ideal choice for secure communication [3]. The journey of VLC is just over two decades; however, it has its applications in the field of indoor positioning, asset tracking, smart homes, and IoT. As far as the accuracy and cost of VLC is concerned, in comparison with the competing technologies for localization, such as Wi-Fi, Cellular, GNSS, RFID, Infrared, Bluetooth, and Ultrasonic, VLC has the highest indoor positioning accuracy and lower cost, thus making it an ideal choice for positioning [4]. The first standardization for VLC, IEEE 802.15.7, was led by the IEEE 802.15 working group in 2011, followed by media coverage of VLC by Time magazine and TED Talk. Several notable technology companies from the United States, Europe, Japan, and South Korea are keenly investing in this technology and developing technological solutions.

Indoor positioning systems rely on two fundamental steps i.e.: measurement techniques used to capture signal information, and then, with the help of algorithms, that information is translated into position estimation. The classical techniques used for Visible Light Positioning (VLP) are Received Signal Strength (RSS), Angle of Arrival (AoA), Time Difference of Arrival (TODA), and Image-Sensor-Based Schemes. These techniques are also combined in hybrid systems to enhance the position accuracy [5]. Building on these measurements, frequently employed algorithms for VLP are triangulation, trilateration, fingerprint, machine learning-based methods, and Particle and Kalman filters. However, recently, ML-based methods have been predominantly used in literature [6]. ML algorithms employ the power of statistical models to learn patterns and relationships from labelled training data. However, these algorithms need a considerable amount of data to analyse and anticipate results, which is a time-consuming process. On the contrary, ML-based algorithms have the potential to overcome the limitations of traditional VLP methods, such as nonlinearities of LEDs, impact of noise and multi-path interference, environmental dynamics, and the need for manual calibration, to name a few [7].

The key contributions and the structure of the paper are as follows.

- A 3×3 LED grid is deployed in a $5 \times 5 \times 3$ meter VLC-enabled room, and an end-to-end simulation framework is developed in conjunction with the mathematical formulation for RSS-based indoor localization. Five machine learning algorithms (MLP, Random Forest, KNN, SVR, and Linear Regression) are evaluated in both noise-free and noisy environments, highlighting their accuracy and robustness.
- A dynamic receiver model is introduced, where the receiver moves along a zigzag trajectory. This setup enables the analysis of model stability and tracking performance in real-time localization scenarios.
- Localization performance is evaluated using Root Mean Square Error (RMSE), spatial error maps, and computational cost analysis. RMSE benchmarks, heatmaps, dynamic error traces, and execution-time graphs collectively provide a comprehensive assessment of each model's behaviour and effectiveness.

The rest of the paper is organized as follows. Section II talks about the mathematical model for the VLP system, and it includes the transmitter, receiver, channel, noise, and receiver mobility. Section III presents the simulation setup and simulation parameters. Section IV presents the results for localization accuracy of ML models with and without noise, position-wise localization error for each model, training and inference time for each model, true and predicted path in a dynamic environment, and localization error over time for a mobile receiver. The paper is concluded in Section V, along with the limitations in our work and a path for future work.

II. MATHEMATICAL MODEL FOR VLP

A. LED-to-Receiver Geometry

Consider an indoor environment of size $L_x \times L_y \times L_z$ meters, where a set of N LED's are mounted on the ceiling at known positions $t_i = (x_i, y_i, z_i)$ for $i = 1, 2, \dots, N$. The receiver is located at position $r = (x_r, y_r, z_r)$.

The Euclidean distance between LED i and the receiver r is given by:

$$d_i = \|t_i - r\| = \sqrt{(x_r - x_i)^2 + (y_r - y_i)^2 + (z_r - z_i)^2} \quad (1)$$

Given that the LEDs are mounted on the ceiling and oriented downward, and the receiver is facing upward, the angle of irradiance ϕ_i and angle of incidence θ_i , are equal and defined as:

$$\cos(\phi_i) = \cos(\theta_i) = \frac{z_i - z_r}{d_i} \quad (2)$$

B. Lambertian VLC Channel Model

The channel gain H_i between LED i and the receiver is modeled by a Lambertian radiation pattern [8]:

$$H_i = \begin{cases} \frac{(m+1)A}{2\pi d_i^2} \cos^m(\phi_i) \cos(\theta_i) T_s g(\theta_i), & \theta_i \leq \text{FOV} \\ 0, & \text{otherwise} \end{cases} \quad (3)$$

Where:

- $m = -\ln(2)/\ln(\cos(\Phi_{1/2}))$ is the Lambertian order.
- A is the photo-diode active area.
- T_s is the gain of the optical filter.

$g(\theta)$ is the concentrator gain defined by:

$$g(\theta) = \begin{cases} \frac{n^2}{\sin^2(\Theta_c)}, & \theta \leq \Theta_c \\ 0, & \text{otherwise} \end{cases} \quad (4)$$

- n is the refractive index of the lens at the receiver.
- Θ_c is the concentrator's field of view (FOV).

C. Received Signal Strength

The received power $P_{r,i}$ from LED i is computed as [9]:

$$P_{r,i} = P_t H_i \quad (5)$$

where P_t is the transmitted optical power from each LED.

D. Machine Learning Input/Output

The machine learning model uses RSS values as input features and corresponding coordinates as output labels [10]:

$$\mathbf{x} = [P_{r,1}, P_{r,2}, \dots, P_{r,N}] \in \mathbb{R}^N, \quad \mathbf{y} = [x, y] \in \mathbb{R}^2 \quad (6)$$

The ML function approximates:

$$\hat{\mathbf{y}} = f_\theta(\mathbf{x}) \quad (7)$$

where f_θ is a supervised model parameterized by θ .

E. Noise Modeling

To simulate real-world imperfections, Gaussian noise is added to the RSS measurements [11]:

$$\tilde{P}_{r,i} = P_{r,i} + \epsilon_i, \quad \epsilon_i \sim \mathcal{N}(0, \sigma^2) \quad (8)$$

F. Dynamic Receiver Model

In dynamic scenarios, the receiver is assumed to move along a trajectory in the 2D plane. Its true position at time t is denoted as:

$$\mathbf{r}(t) = [x(t), y(t), z_r] \quad (9)$$

where z_r is fixed (e.g., receiver height from floor), and $t \in \{1, 2, \dots, T\}$ denotes the time step index.

1) *Trajectory*: Let the trajectory be parameterized by a known path function or sequence, such as:

$$\mathbf{r}(t) = [x_0 + v_x t, y_0 + A \sin(\omega t), z_r] \quad (10)$$

where:

- x_0, y_0 : initial position

- v_x : constant horizontal speed in x-direction
- A : amplitude of lateral oscillation
- ω : frequency of the sine wave (zigzag path)

III. SIMULATION SETUP AND PARAMETERS

Algorithm 1 VLC-Based Indoor Localization with Machine Learning

```

1: Input: Room size, LED positions, receiver height, system parameters
2: Initialize: ML models: MLP, Random Forest, KNN, SVR, Linear Regression
3: for each scenario in {No Noise, With Noise} do
4:   Generate dataset:  $X, y \leftarrow \text{GENERATE-DATASET}(\text{samples} = 2000, \sigma)$ 
5:   Normalize RSS features:  $X \leftarrow \text{STANDARDSCALER}(X)$ 
6:   Split dataset into train/test sets
7:   for each model in ML models do
8:     Train model on training set
9:     Predict  $\hat{y}$  on test set
10:    Compute RMSE and store in Results
11:   end for
12: end for
13: Plot bar chart: RMSE comparison (with/without noise)
14: for each model in ML models do
15:   Predict on test set
16:   Compute error at each position
17:   Generate 2D heatmap of localization error
18:   Measure and store training and inference time
19: end for
20: Plot training and inference time bar charts
21: Dynamic Receiver Simulation:
22: Define path  $\mathbf{r}(t)$  as zigzag trajectory
23: for each point  $t$  in trajectory do
24:   Simulate RSS with noise:  $\mathbf{x}(t)$ 
25:   Predict location:  $\hat{\mathbf{y}}(t) = f(\mathbf{x}(t))$ 
26:   Compute error  $e(t) = \|\hat{\mathbf{y}}(t) - \mathbf{r}(t)\|$ 
27: end for
28: Compute dynamic RMSE over all time steps
29: Plot true vs predicted path and error over time
30: Output: Graphs of RMSE, Error Heatmaps, Training and Inference Time maps, dynamic performance

```

TABLE I. SIMULATION PARAMETERS

Parameter	Value
Room size ($L_x \times L_y \times L_z$)	5 m $\times$ 5 m $\times$ 3 m
Receiver height z_r	0.85 m
Number of LEDs	9 (arranged in 3 $\times$ 3 grid)
LED positions	Uniform grid on ceiling
Photodetector (PD) area A	10^{-4} m ²
Lambertian order m	1
Transmitted power P_t	1 W
Field of View (FOV)	90°
Refractive index n	1.5
Optical filter gain T_s	1
Concentrator gain $g(\theta)$	$n^2/\sin^2(\Theta_c)$
Noise standard deviation σ	10^{-6}
Training samples	2000
Test samples	400 (20% split)
Machine Learning models	MLP, RF, KNN, SVR, Linear Regression
Dynamic path	Zigzag scan across room

IV. RESULTS

In this section, we present the results of the localization accuracy of ML models in the absence and presence of noise, position-wise localization error for each ML model, training and inference time for ML models, and the localization results regarding the dynamic path.

A. Performance Metric: RMSE and Localization Error

Localization accuracy is evaluated using Root Mean Square Error (RMSE) [12]:

$$\text{RMSE} = \sqrt{\frac{1}{M} \sum_{j=1}^M [(x_j - \hat{x}_j)^2 + (y_j - \hat{y}_j)^2]} \quad (11)$$

where M is the number of test samples, and $(\hat{x}_j, \hat{y}_j)$ is the predicted location for the j -th test point.

Figure 1 compares the localization accuracy of five machine learning models under the aforementioned conditions. From the bar chart, it is evident that there is some degradation in all the models where Gaussian noise is added to RSS values, with Linear Regression showing the steepest increase in RMSE due to its inability to model non-linear relationships. As far as KNN and Random Forest are concerned, they maintain high accuracy with relatively small increases in RMSE, demonstrating better robustness to input uncertainty. SVR deals with the noise well, but with a slightly higher baseline error. These results suggest that tree-based and instance-based models are better suited for noisy VLC environments.

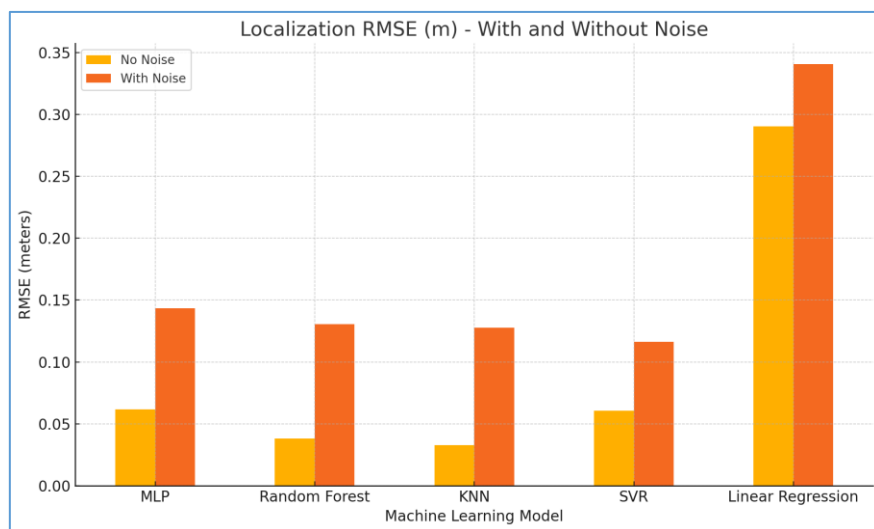


Fig. 1. Localization accuracy of ML models with and without noise.

Figure 2 presents the training time (in seconds) required for each model to learn from the noisy dataset. Random Forest and KNN show a substantial decrease in training time when compared to MLP and SVR. MLP needs the most time to converge due to its recursive optimization and complicated structure. SVR, while simpler, is computationally expensive in training because it involves solving a quadratic problem. Linear Regression, being analytically solvable, is the fastest to train. These results are critical for real-time or resource-constrained deployments.

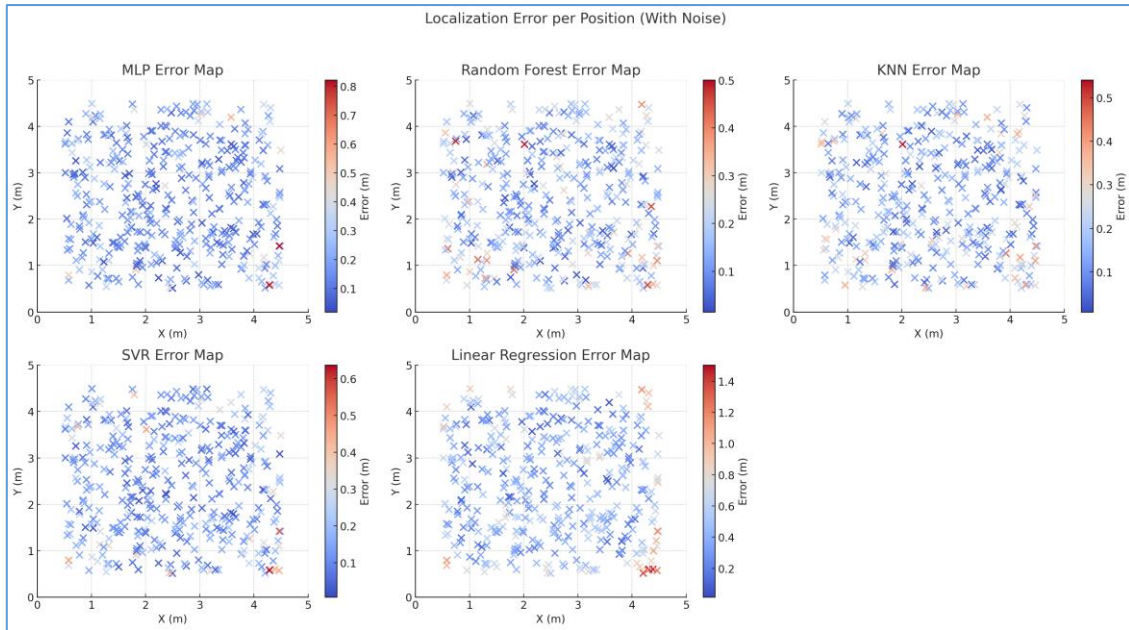


Fig. 2. Position-wise localization error for a each ML model.

B. Training & Model Inference

At each time step, the trained ML model gives an estimated location:

$$\hat{\mathbf{y}}(t) = f_{\theta}(\mathbf{x}(t)) \quad (12)$$

Figure 3 shows the computational cost of training each machine learning model on noisy RSS data. As expected, Linear Regression has the shortest training time due to its closed-form solution. KNN also trains swiftly, as it only stores the training samples without any optimization process. Random Forest experiences a minor delay, as it must build and optimize multiple decision trees, yet preserves a high level of efficiency. In contrast, MLP and SVR require a much longer runtime to train. MLP involves iterative weight updates via backpropagation, while SVR requires solving a quadratic optimization problem for each output dimension. These results demonstrate the trade-off between training time and model complexity, which is important when models are re-trained often or deployed in an edge environment.

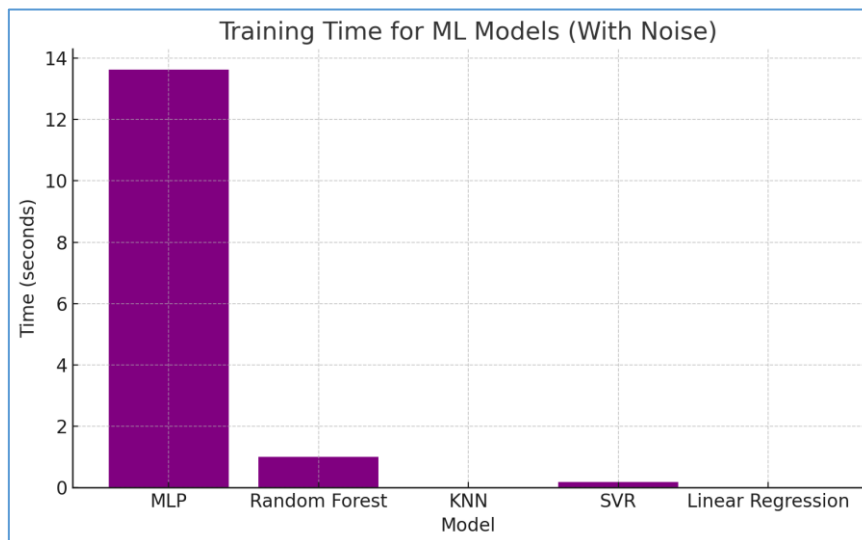


Fig. 3. Training Time for ML models in the presence of noise

The inference time plot elucidates the time each model requires to predict receiver positions on the test data. Random Forest and Linear Regression demonstrate fast prediction capabilities, making them highly suitable for real-time localization applications. In contrast, KNN and SVR are the slowest, as KNN must compute distances to all training samples during inference, and SVR involves kernel-based operations. MLP shows moderate inference time, balancing complexity and speed. These results highlight that inference efficiency is just as important as accuracy in practical VLC positioning systems, particularly in mobility scenarios.

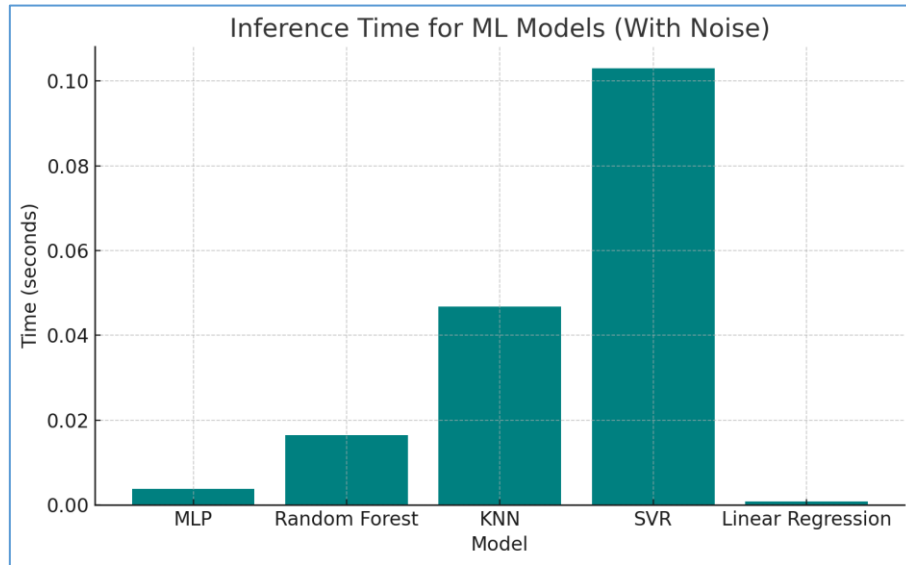


Fig. 4. Inference Time for ML models in the presence of noise.

C. Error Over Time and Time-Varying RSS

Localization error is evaluated at each frame:

$$e(t) = \|\mathbf{r}(t) - \hat{\mathbf{y}}(t)\|_2 \quad (13)$$

The overall tracking performance across the path is summarized using:

$$\text{RMSE}_{\text{dynamic}} = \sqrt{\frac{1}{T} \sum_{t=1}^T e(t)^2} \quad (14)$$

At each step t , the received signal vector is:

$$\mathbf{x}(t) = [P_{r,1}(t), P_{r,2}(t), \dots, P_{r,N}(t)] \quad (15)$$

Each $P_{r,i}(t)$ is computed based on the time-varying position $r(t)$:

$$P_{r,i}(t) = \frac{(m+1)A}{2\pi d_i(t)^2} \cos^m(\phi_i(t)) \cos(\theta_i(t)) T_s g(\theta_i(t)) \quad (16)$$

Next, we compare the ground truth path of the receiver with the trajectory estimated by the Random Forest model as it moves dynamically in a zigzag fashion across the room. The result is given in Figure 5. The blue line is for the true path, and the red line is for the predicted path using the ML model. From Figure 5, it can be seen that the ML model accurately tracks the general shape and direction of movement, preserving precise alignment with the true path over the majority of the trajectory length. As far as the deviations are concerned, they are minimal. The performance is superior, especially in the central and mid-range positions. Minor discrepancies at the path edges may be overcome by training the model more on sharp turns.

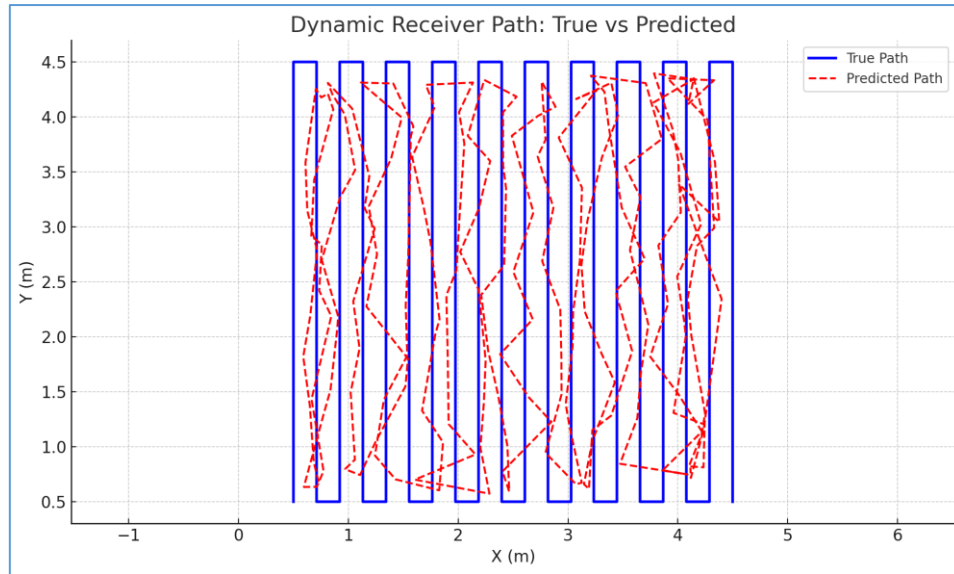


Fig. 5. True and predicted path in dynamic environment.

Finally, we demonstrate the localization error over time in Figure 6 as the receiver moves along its dynamic path. Generally, the error remains steady and bounded, mostly under 0.2 meters, demonstrating the model's robustness to continuous movement and noisy RSS readings. The performance at the corners is affected primarily due to poorer LED coverage and more pronounced angle effects. In a nutshell, the performance of the ML model is persistent and demonstrates its suitability for real-time tracking applications in indoor VLC environments.

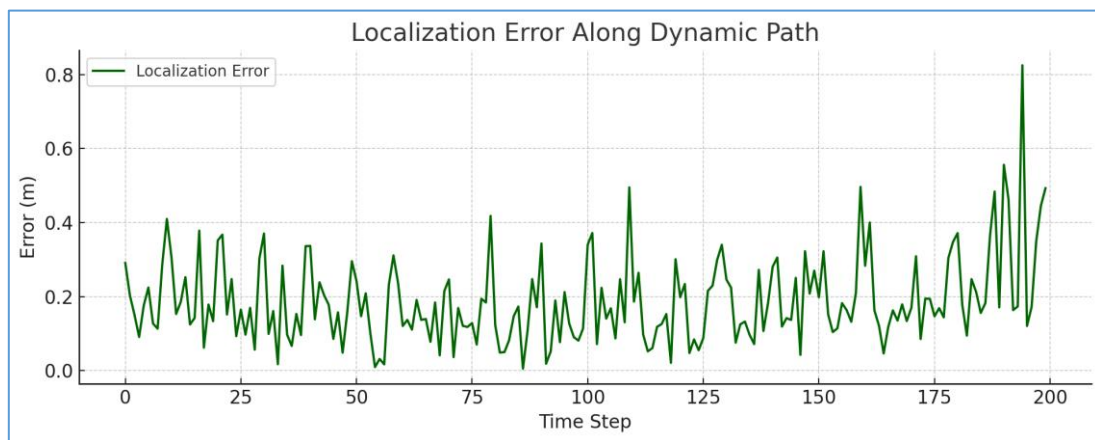


Fig. 6. Localization error over time for a mobile receiver.

TABLE II. MACHINE LEARNING MODEL CONFIGURATIONS

Model	Configuration / Hyperparameters
MLP Regressor	2 hidden layers (64, 64), activation: ReLU, max iter: 1000
Random Forest	100 trees, max depth: auto, criterion: MSE
KNN Regressor	$k = 5$, distance metric: Euclidean
Support Vector Regression	RBF kernel, default C and ϵ , wrapped in MultiOutputRegressor
Linear Regression	Ordinary Least Squares, wrapped in MultiOutputRegressor

V. CONCLUSION

Due to the high demand for location-based services, indoor localization has gained a lot of attention both in industry and academia, given its potential market share in logistics, smart buildings, asset management, and entertainment. Machine learning has influenced all walks of life, and indoor localization is no exception. This work demonstrates the use of five machine learning models, namely linear regression, random forest, KNN, SVR, and MLP, to estimate the position of static and mobile receivers. A simulation framework was developed to model a realistic indoor room illuminated by a 3×3 LED grid, and datasets were generated under both noise-free and noisy conditions. ML model performance was evaluated based on RMSE, computational efficiency, and spatial error distribution. Simulation results show that both KNN and Random Forest achieved the best localization accuracy in static and mobile environments. SVR results show robustness to noise; however, the training and inference times are higher. LR fails to perform well due to its inefficiency in handling nonlinear RSS to position relationships. The dynamic simulation allows us to validate the performance of ML models useful to real-time tracking applications.

The limitation of this work includes the assumption of ideal Lambertian propagation with no multi-path or diffuse reflections. The simulation only considers 2D planar receiver estimation. The data generated to train and test the model was simulation-based. Future work includes hardware implementation to achieve real-world data for training and testing purposes. Additionally, we plan to incorporate multi-path effects and extend the framework to a 3D environment.

REFERENCES

- [1] J. Singh, N. Tyagi, S. Singh, F. Ali, and D. Kwak, "A systematic review of contemporary indoor positioning systems: Taxonomy, techniques, and algorithms," *IEEE Internet of Things Journal*, 2024.
- [2] R. Wang, G. Niu, Q. Cao, C. S. Chen, and S.-W. Ho, "A survey of visible-light-communication-based indoor positioning systems," *Sensors*, vol. 24, no. 16, p. 5197, 2024.
- [3] Z. Zhu, Y. Yang, M. Chen, C. Guo, J. Cheng, and S. Cui, "A survey on indoor visible light positioning systems: Fundamentals, applications, and challenges," *IEEE Communications Surveys & Tutorials*, 2024.
- [4] S. Bastiaens, M. Alijani, W. Joseph, and D. Plets, "Visible light positioning as a next-generation indoor positioning technology: A tutorial," *IEEE Communications Surveys & Tutorials*, 2024.
- [5] M. Saadi, T. Ahmad, M. Kamran Saleem, and L. Wuttisittikulkij, "Visible light communication—an architectural perspective on the applications and data rate improvement strategies," *Transactions on emerging telecommunications technologies*, vol. 30, no. 2, p. e3436, 2019.
- [6] Q. Chen, F. Wang, B. Deng, L. Qin, and X. Hu, "Indoor visible light positioning system based on memristive convolutional neural network," *Optics Communications*, vol. 577, p. 131340, 2025.
- [7] M. Saadi, A. Bajpai, D. Z. Rodriguez, and L. Wuttisittikulkij, "Investigating the role of channel state information for mimo based visible light communication system," in *2022 37th International Technical Conference on Circuits/Systems, Computers and Communications (ITC-CSCC)*. IEEE, 2022, pp. 1–4.
- [8] B. A. Vijayalakshmi, P. Gandhimathi, and M. Nesasudha, "Vlc channel characteristics and data transmission model in indoor environment for future communication: an overview," *Journal of Optics*, vol. 53, no. 2, pp. 933–939, 2024.

- [9] K. Zhang, Z. Zhang, and B. Zhu, “Beacon led coordinates estimator for easy deployment of visible light positioning systems,” *IEEE Transactions on Wireless Communications*, vol. 21, no. 12, pp. 10208–10223, 2022.
- [10] H. Q. Tran and C. Ha, “Machine learning in indoor visible light positioning systems: A review,” *Neurocomputing*, vol. 491, pp. 117–131, 2022.
- [11] N. Liu, L.-H. Hong, P. Feng, H.-N. Yang, and J.-Y. Wang, “Performance analysis for visible light communications with gaussian-plus-laplacian additive noise,” *Telecommunication Systems*, vol. 85, no. 3, pp. 373–387, 2024.
- [12] T. O. Hodson, “Root mean square error (rmse) or mean absolute error (mae): When to use them or not,” *Geoscientific Model Development Discussions*, vol. 2022, pp. 1–10, 2022.

Compliance with ethical standards**Disclosure of conflict of interest**

The authors declare that they have no conflict of interest.

Disclaimer/Publisher’s Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of ALBAHIT and/or the editor(s). ALBAHIT and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content