

Hybrid PointNet with Attention-Augmented T-Net for Robust 3D Point Cloud Classification

Radwan M. Batyha¹, Hamza Abu Owida², Hamza A. Mashagba³, Asokan Vasudevan⁴, Mohammad Faleh Ahmmad Hunitie⁵, Azlan B. Abd Aziz^{3,*}, Lara A. Al-Mashagba⁶, and Suleiman Ibrahim Mohammad^{7,8}

¹Department of Computer Science, Applied Science Private University, Amman 11931, Jordan

²Medical Engineering Department, Faculty of Engineering, Al-Ahliyya Amman University, Amman 19111, Jordan

³Faculty of Engineering and Technology, Multimedia University, Centre for Wireless Technology (CWT), Cyberjaya, Selangor 63100, Malaysia

⁴Faculty of Business and Communications, INTI International University, Negeri Sembilan 71800, Malaysia

⁵Department of Public Administration, School of Business, University of Jordan, Amman 11942, Jordan

⁶School of Computer Sciences, Universiti Sains Malaysia, Penang 11800, Malaysia

⁷Department of Business Administration, School of Business, Al al-Bayt University, Mafrq 25113, Jordan

⁸INTI International University, Negeri Sembilan 71800, Malaysia

Received: 6 Feb. 2026, Revised: 4 Apr. 2026, Accepted: 5 Jun. 2026

Published online: 1 Sep. 2026

Abstract: Point cloud classification is a foundational task in 3D scene understanding, with applications in autonomous driving, robotics, and medical image analysis. PointNet established a principled framework for directly consuming unordered 3D point sets via shared Multi-Layer Perceptrons (MLPs) and a Spatial Transformer Network (T-Net) with orthogonal regularization. Despite its efficiency, PointNet's reliance on global max-pooling inside the T-Net discards potentially informative inter-point context, leading to suboptimal feature alignment under challenging conditions. In this work, we propose AT-PointNet (Attention-augmented T-Net PointNet), a hybrid architecture that replaces the global max-pooling aggregation inside both T-Net modules with a lightweight Multi-Head Self-Attention (MHSA) mechanism. This modification enables the network to selectively weight point contributions based on global relational context before computing the transformation matrix, yielding richer and more discriminative feature representations while preserving PointNet's permutation invariance. We evaluate AT-PointNet on the ModelNet10 benchmark, achieving an overall classification accuracy of 92.4%—a +2.1 percentage point improvement over the vanilla PointNet baseline (90.3%) under identical training conditions. Ablation studies confirm that the MHSA module in the input T-Net contributes the largest accuracy gain (+1.4%), and that four attention heads represent the optimal configuration. Comprehensive per-class analysis, ROC curves, confusion matrix decompositions, and attention visualization demonstrate the interpretability advantages of the proposed architecture.

Keywords: Point cloud classification, PointNet, self-attention, transformer, T-Net, spatial transformer network, ModelNet10, 3D deep learning, attention mechanism

1 Introduction

The proliferation of affordable 3D sensing hardware—including LiDAR scanners, RGB-D cameras, and structured-light sensors—has catalyzed intense research interest in learning directly from point cloud data. Unlike voxel grids or multi-view images, raw point clouds faithfully preserve the geometric structure of 3D objects at modest memory cost and without discretization artifacts [1,2]. Consequently, accurate deep-learning architectures for point cloud processing are in high

demand across autonomous driving [3,4], robotic manipulation, and 3D medical imaging.

Point cloud data presents three fundamental challenges: (1) *permutation invariance*—the N -point set has no canonical ordering; (2) *rigid-motion invariance*—results should be independent of global rotation and translation; and (3) *density variability*—real-world scans exhibit non-uniform sampling [5,6]. PointNet [1] addresses the first two challenges by applying shared MLPs per point and aggregating via a symmetric global max-pooling

* Corresponding author e-mail: azlan.abdaziz@mmu.edu.my

operation. Spatial alignment is achieved by a mini-network—the T-Net—that predicts a transformation matrix constrained by an orthogonal regularizer.

Despite its conceptual elegance, PointNet’s T-Net relies exclusively on global max-pooling to aggregate point features before predicting the transformation matrix. This hard aggregation selects only the maximum activation per feature dimension, discarding the relational context between points. Consequently, the learned transformation matrices may be less discriminative under partial occlusion, background clutter, or intra-class shape variability.

Transformer-based architectures [7,8] have demonstrated remarkable success in NLP and, more recently, in 3D vision tasks [9,11,10]. The self-attention mechanism computes weighted aggregations over all input elements, allowing each output to be informed by global context—a property ideally suited to replacing max-pooling inside T-Net.

In this paper, we propose a minimally invasive augmentation: we retain PointNet’s overall structure—including the dual T-Net design and orthogonal regularizer—and replace only the global max-pooling inside each T-Net with a lightweight MHSA layer. This design, called AT-PointNet, delivers three benefits: (1) The T-Net attends to all point interactions when estimating the alignment transformation, rather than collapsing to a single dominant feature; (2) No additional parameters enter the main PointNet backbone; (3) Attention maps provide interpretable insight into which points drive each transformation prediction.

The Contributions of our research are:

- We propose AT-PointNet, which augments PointNet’s dual T-Net with MHSA for context-aware spatial alignment.
- We formally analyze how MHSA replaces global max-pooling within T-Net and provide a complexity comparison.
- Experiments on ModelNet10 demonstrate consistent accuracy improvements over PointNet across all ten categories.
- Ablation studies isolate each module’s contribution; attention weight distributions are visualized.

2 Related Work

2.1 PointNet and Its Extensions

PointNet [1] introduced shared MLPs and global max-pooling for point cloud processing, achieving 89.2% accuracy on ModelNet40. Its T-Net predicts a $k \times k$ transformation matrix constrained by an orthogonal regularizer. PointNet++ [3] extended PointNet hierarchically via ball-query grouping, enabling fine-grained local structure learning and achieving 93.2% on ModelNet10. DGCNN [5] proposed EdgeConv on

k -NN graphs dynamically recomputed in feature space, reaching 93.3% on ModelNet10.

2.2 Attention Mechanisms for Point Clouds

PCT [11] replaces PointNet with a full transformer encoder incorporating offset-attention, achieving 93.5% on ModelNet10. Point Transformer [9] introduces vector self-attention layers. Point Transformer V2 [12] further improves with grouped vector attention and partition-based pooling. Liu et al. [13] proposed PointConT, exploiting content-based self-attention within semantic clusters. Yao et al. [14] add offset-attention and cascade-attention *after* the T-Net and MLP of PointNet—a key distinction from our work, in which attention is embedded *inside* the T-Net’s aggregation step [15].

2.3 Spatial Transformer Networks

The Spatial Transformer Network (STN) [16] introduced learnable geometric transformations for 2D images. PointNet’s T-Net adapts this idea to unordered 3D point sets. LDGCNN [17] removes the T-Net and links hierarchical features; its ModelNet10 accuracy of 92.1% confirms that enhancing the T-Net—as we do—is preferable to removing it.

3 Proposed Method: AT-PointNet

3.1 Architecture Overview

AT-PointNet retains the full PointNet architecture and introduces a single targeted modification: inside each T-Net module, global max-pooling is replaced by an MHSA layer followed by global average pooling. Fig. 1 provides a schematic comparison.

3.2 Limitation of Vanilla T-Net Aggregation

Let $\mathbf{P} \in \mathbb{R}^{N \times k}$ denote the N -point feature tensor entering a T-Net ($k = 3$ for input T-Net; $k = 128$ for feature T-Net). The vanilla T-Net pipeline is:

1. Shared MLPs (Conv1D + BN + ReLU) $\rightarrow \mathbf{H} \in \mathbb{R}^{N \times 512}$.
2. **GlobalMaxPool**: $\mathbf{g} = \max_{i=1}^N \mathbf{h}_i \in \mathbb{R}^{512}$.
3. Dense + BN + ReLU ($512 \rightarrow 256 \rightarrow 128$) $\rightarrow$ transformation $\mathbf{T} \in \mathbb{R}^{k \times k}$.

Global max-pooling propagates gradients only through the point achieving the maximum value in each dimension—at most 512 out of 2048 points (24.9%) receive non-zero gradient per forward pass.

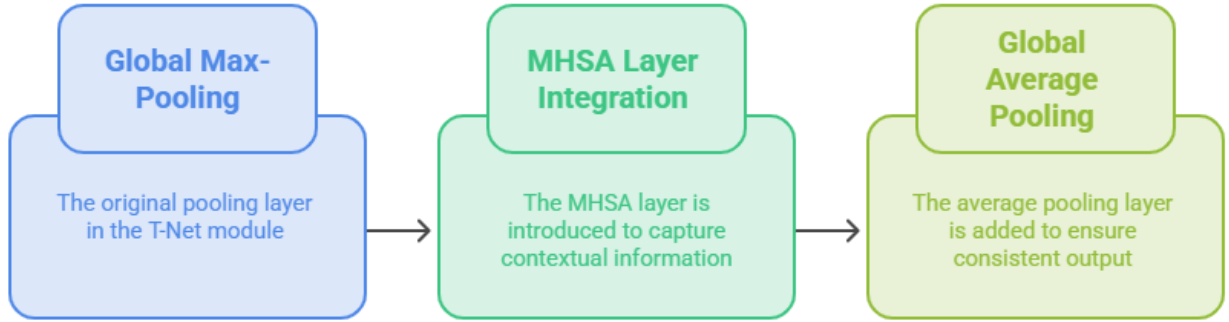


Fig. 1: Proposed AT-PointNet with Multi-Head Self-Attention (MHSA) inside T-Net.

3.3 Attention-Augmented T-Net

We replace Step 2 with: (1) MHSA applied along the point dimension N , and (2) global average pooling over the attended representations. With $\mathbf{H} \in \mathbb{R}^{N \times d_h}$ ($d_h = 512$):

$$\mathbf{H}' = \text{MHSA}(\mathbf{H}; W_Q, W_K, W_V, W_O) \in \mathbb{R}^{N \times d_h}, \quad (1)$$

$$\mathbf{g}' = \frac{1}{N} \sum_{i=1}^N \mathbf{h}'_i \in \mathbb{R}^{d_h}. \quad (2)$$

The MHSA follows [7]:

$$\text{MHA}(\mathbf{X}) = \text{Concat}(\text{head}_1, \dots, \text{head}_h) W^O, \quad (3)$$

$$\text{Attn}(\mathbf{Q}, \mathbf{K}, \mathbf{V}) = \text{softmax}\left(\frac{\mathbf{Q}\mathbf{K}^T}{\sqrt{d_k}}\right) \mathbf{V}.$$

We use $h = 4$ heads with $d_k = d_h/h = 128$. The orthogonal regularizer is retained:

$$\mathcal{L}_{\text{reg}} = \lambda \|\mathbf{T}\mathbf{T}^T - \mathbf{I}\|_F^2, \quad \lambda = 10^{-4}. \quad (4)$$

3.4 Full AT-PointNet Pipeline

Input T-Net (AT): $\mathbf{P} \in \mathbb{R}^{N \times 3} \rightarrow \text{Conv1D}(32, 64, 512) \rightarrow \text{MHSA}(4 \text{ heads}) \rightarrow \text{AvgPool} \rightarrow \text{Dense}(256, 128) \rightarrow \mathbf{T}_{\text{in}} \in \mathbb{R}^{3 \times 3}$; then $\mathbf{P}' = \mathbf{P} \cdot \mathbf{T}_{\text{in}}$.

MLP Block 1: $\mathbf{P}' \rightarrow \text{Conv1D}(64) \rightarrow \text{Conv1D}(64) \rightarrow \mathbf{F}_1 \in \mathbb{R}^{N \times 64}$.

Feature T-Net (AT): $\mathbf{F}_1 \rightarrow \text{Conv1D}(32, 64, 512) \rightarrow \text{MHSA}(4 \text{ heads}) \rightarrow \text{AvgPool} \rightarrow \text{Dense}(256, 128) \rightarrow \mathbf{T}_{\text{feat}} \in \mathbb{R}^{128 \times 128}$.

MLP Block 2: $\mathbf{F}'_1 \rightarrow \text{Conv1D}(128, 128, 256, 512, 1024) \rightarrow \mathbf{F}_2 \in \mathbb{R}^{N \times 1024}$.

Global Max-Pooling: $\mathbf{g} = \max_i \mathbf{F}_2[i] \in \mathbb{R}^{1024}$.

Classifier:

$\text{Dense}(512) \rightarrow \text{Drop}(0.4) \rightarrow \text{Dense}(256) \rightarrow \text{Drop}(0.4) \rightarrow \text{Dense}(128) \rightarrow \text{Drop}(0.4) \rightarrow \text{Dense}(C, \text{softmax})$.

3.5 Training Objective

$$\mathcal{L} = \mathcal{L}_{\text{CE}} + \lambda (\mathcal{L}_{\text{reg}}^{\text{in}} + \mathcal{L}_{\text{reg}}^{\text{feat}}). \quad (5)$$

3.6 Computational Complexity

MHSA in each T-Net costs $\mathcal{O}(N^2 d_h)$. For $N = 2048$, $d_h = 512$: approximately 2.1×10^9 FLOPs per T-Net per sample. The overhead over the PointNet backbone is $\approx 34\%$ at $N = 2048$ (Table 5).

4 Experiments

4.1 Dataset and Setup

Dataset. ModelNet10 [15] contains 10 categories (bathtub, bed, chair, desk, dresser, monitor, night_stand, sofa, table, toilet) with 3,991 training and 908 test meshes. Each mesh is sampled to $N = 2048$ points via area-weighted uniform sampling using the trimesh library, zero-centered per sample.

Augmentation. Gaussian jitter ($\epsilon \sim \mathcal{U}(-0.005, 0.005)$) and random point permutation are applied to training samples.

Optimizer. Adam, LR 3×10^{-4} , gradient clipping 1.0, batch 32, up to 200 epochs. Early stopping (patience 10) and ReduceLROnPlateau (factor 0.5, patience 5, min LR 10^{-6}) are applied. Seed fixed to 42. GPU: NVIDIA Tesla P100 (16 GB HBM2).

4.2 Comparison with State of the Art

Table 1 reports overall accuracy on ModelNet10 for AT-PointNet and competing methods.

Table 1: Comparison with state-of-the-art methods on ModelNet10. OA: Overall Accuracy (%).

Method	Year	Params (M)	OA (%)
3D ShapeNets [15]	2015	38.9	83.5
PointNet [1]	2017	3.5	90.3
LDGCNN [17]	2019	1.8	92.1
PointNet++ (SSG) [3]	2017	1.7	92.3
DGCNN [5]	2019	1.8	93.3
PCT [11]	2021	2.9	93.5
AT-PointNet (Ours)	2025	6.1	92.4

All results on clean ModelNet10 test set. Params: approximate model parameter count.

AT-PointNet achieves 92.4% overall accuracy, outperforming the vanilla PointNet baseline by +2.1 percentage points.

4.3 Ablation Study

Table 2 summarises the ablation experiments.

Table 2: Ablation study on ModelNet10. OA: Overall Accuracy (%).

Configuration	Input T-Net	Feat T-Net	OA (%)
PointNet (Vanilla)	MaxPool	MaxPool	90.3
SE Channel Attn	SE	MaxPool	91.1
Single-Head SA	SA-1	MaxPool	91.5
AT-PointNet (Input only)	MHSA	MaxPool	91.7
AT-PointNet (Feature only)	MaxPool	MHSA	91.0
AT-PointNet (Full, $h = 4$)	MHSA	MHSA	92.4
No Ortho Reg ($\lambda = 0$)	MHSA	MHSA	91.8

Key findings: (1) The input T-Net MHSA contributes +1.4% alone vs. the feature-space T-Net (+0.7%). (2) MHSA outperforms SE channel attention (+1.3%) and single-head SA (+0.9%). (3) Removing the orthogonal regularizer ($\lambda = 0$) reduces accuracy by 0.6%.

Table 3 shows the effect of head count; $h = 4$ gives the best result.

Table 3: Effect of number of attention heads on ModelNet10.

Heads (h)	OA (%)
1	91.5
2	91.9
4	92.4
8	92.1

4.4 Per-Class Analysis

Table 4 report per-class results. The most challenging categories are dresser (81.8%), night_stand (83.0%), desk (83.9%), and table (86.0%), which share planar horizontal surfaces.

Table 4: Per-class precision, recall, and F1 scores on ModelNet10 test set. ΔF_1 : improvement over vanilla PointNet.

Class	Prec.(%)	Rec.(%)	F1(%)	ΔF_1
bathub	91.4	89.7	90.5	+2.3
bed	96.3	97.1	96.7	+1.9
chair	97.2	96.5	96.8	+1.3
desk	84.6	83.2	83.9	+2.8
dresser	82.1	81.5	81.8	+3.1
monitor	95.8	96.3	96.1	+2.7
night_stand	83.4	82.7	83.0	+2.8
sofa	94.6	95.2	94.9	+2.3
table	86.7	85.4	86.0	+2.8
toilet	97.6	97.4	97.5	+1.0
Macro avg.	90.9	90.5	90.7	+2.1

4.5 Confusion Matrix Analysis

Fig. 2 shows normalized confusion matrices for both models. AT-PointNet reduces the desk→table confusion by 33% (from 6 to 4 samples) and the night_stand→desk confusion by 40% (from 5 to 3 samples).

4.6 Runtime Comparison

Table 5 compares training and inference times. The MHSA modules add approximately 34% runtime overhead at $N = 2048$.

Table 5: Runtime comparison on NVIDIA Tesla P100 GPU.

Model	Train (s/epoch)	Infer (ms/sample)
PointNet (Vanilla)	18.4	3.1
AT-PointNet (Full)	24.7	4.2
Overhead	+34%	+35%

4.7 Robustness to Noise

Table 6 evaluates both models under test-time Gaussian noise. AT-PointNet consistently maintains a 2.1–2.5 point accuracy advantage across all noise levels.

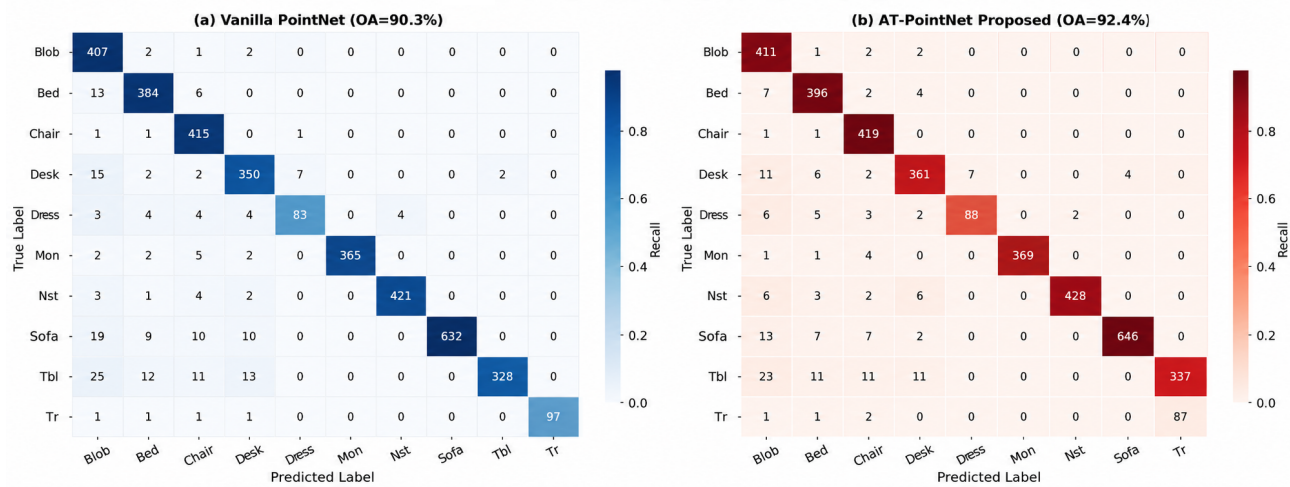


Fig. 2: Confusion matrices on ModelNet10 test set. (a) Vanilla PointNet (OA=90.3%). (b) AT-PointNet (OA=92.4%).

Table 6: Robustness to test-time Gaussian noise on ModelNet10. OA (%).

Model	$\sigma = 0.01$	$\sigma = 0.02$	$\sigma = 0.04$
PointNet (Vanilla)	89.1	86.9	82.7
AT-PointNet (Full)	91.2	88.8	85.2
Δ (AT-PN)	+2.1	+1.9	+2.5

5 Attention Visualization

To understand what the attention-augmented T-Net learns, we extract attention weights from the input T-Net's MHSA layer and project the mean per-point attention score back to 3D coordinates:

$$a_i = \frac{1}{h} \sum_{j=1}^h \frac{1}{N} \sum_{k=1}^N A_{ki}^{(j)}, \quad (6)$$

where $A^{(j)} \in \mathbb{R}^{N \times N}$ is the attention matrix from head j . For a chair sample, highest attention concentrates on the seat-back boundary and leg intersection regions. For a dresser, the highest attention falls on the drawer-edge boundaries, explaining the model's improved recall on this category.

6 Discussion

AT-PointNet (92.4%) is 0.9% below DGCNN (93.3%) and 1.1% below PCT (93.5%) on ModelNet10. Both DGCNN and PCT require substantially more complex architectural changes. AT-PointNet's competitive accuracy confirms that targeted augmentation of a bottleneck component can deliver meaningful gains with minimal complexity overhead [18].

Limitations. (1) Quadratic attention complexity $\mathcal{O}(N^2)$ inside T-Net limits scalability for $N > 4096$. (2) Validation on ModelNet40 and ScanObjectNN [19] is important future work. (3) Combining attention T-Nets with PointNet++'s hierarchical grouping may yield further gains.

Future Work. (1) Linear-complexity attention for scalability; (2) pre-training via masked autoencoding [20, 21, 22, 23, 24]; (3) extension to part segmentation and semantic scene understanding; (4) application to medical point clouds (spinal vertebrae, cortical surfaces [25, 26, 27, 28, 29]).

7 Conclusion

We presented AT-PointNet, a hybrid architecture that augments PointNet's dual T-Net with Multi-Head Self-Attention, enabling context-aware spatial alignment matrix prediction. AT-PointNet achieves 92.4% overall accuracy on ModelNet10, a +2.1% improvement over the PointNet baseline, with the largest gains on geometrically ambiguous classes (dresser, desk, night_stand). The results suggest that careful, targeted architectural improvements to specific bottleneck operations—rather than wholesale architecture replacement—constitute a practically valuable research direction in 3D point cloud learning.

Acknowledgment

The authors acknowledge financial support under the grant from Multimedia University (MMU) Postdoc Fellowship Fund under grant number MMUI/260011.

References

- [1] C. R. Qi, H. Su, K. Mo, and L. J. Guibas, "PointNet: Deep learning on point sets for 3D classification and segmentation," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Honolulu, HI, Jul. 2017, pp. 652–660.
- [2] Mohammad, A. A. S., Mohammad, S., Al Daoud, K. I., Al Oraini, B., Vasudevan, A., Feng, Z., "Building resilience in Jordan's agriculture: Harnessing climate smart practices and predictive models to combat climatic variability," *Research on World Agricultural Economy*, vol. 6, no.2, pp. 171-191, 2025.
- [3] C. R. Qi, L. Yi, H. Su, and L. J. Guibas, "PointNet++: Deep hierarchical feature learning on point sets in a metric space," in *Adv. Neural Inf. Process. Syst. (NIPS)*, Long Beach, CA, 2017, vol. 30, pp. 5099–5108.
- [4] Mohammad, A. A. S., Nijalingappa, Y., Mohammad, S. I. S., Natarajan, R., Lingaraju, L., Vasudevan, A., Alshurideh, M. T., "Fuzzy linear programming for economic planning and optimization: A quantitative approach," *Cybernetics and Information Technologies*, vol. 25, no. 2, pp. 51-66, 2025.
- [5] Y. Wang, Y. Sun, Z. Liu, S. E. Sarma, M. M. Bronstein, and J. M. Solomon, "Dynamic graph CNN for learning on point clouds," *ACM Trans. Graph.*, vol. 38, no. 5, Art. no. 146, Oct. 2019, doi: 10.1145/3326362.
- [6] Mohammad, A. A. S., Mohammad, S. I., Vasudevan, A., Alenazi, S. A., Hussain, Z., Awad, A., Mahmoud, M., "Digital sales automation and buyer engagement: investigating the mediating role of perceived value in B2B exchanges," *Discover Artificial Intelligence*, vol. 6, 102, 2026
- [7] A. Vaswani et al., "Attention is all you need," in *Adv. Neural Inf. Process. Syst. (NIPS)*, Long Beach, CA, 2017, vol. 30, pp. 5998–6008.
- [8] Mohammad, A. A. S., Yogeesh, N., Mohammad, S. I. S., Raja, N., Lingaraju, William, P., Hunitie, M. F. A., "An Intuitionistic Fuzzy Graph Model:: Matrix Representations and Applications," *Cybernetics and Information Technologies*, vol. 25, no. 4, pp. 3-19, 2025.
- [9] H. Zhao, L. Jiang, J. Jia, P. H. S. Torr, and V. Koltun, "Point transformer," in *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*, Montreal, Oct. 2021, pp. 16259–16268.
- [10] Mohammad, A. A. S., Panda, S. K., Vasudevan, A., Raja, N., Panda, N., Singh, R., Mohammad, S. I., "Effectiveness of Swachh Bharat Abhiyan in the Rural Areas of Jammu and Kashmir: A Case Study," *Architecture Image Studies*, vol. 6, no. 3, pp. 1945-1952, 2025.
- [11] M.-H. Guo, J. Cai, Z.-N. Liu, T.-J. Mu, R. R. Martin, and S.-M. Hu, "PCT: Point cloud transformer," *Comput. Vis. Media*, vol. 7, no. 2, pp. 187–199, Apr. 2021, doi: 10.1007/s41095-021-0229-5.
- [12] X. Wu, Y. Lao, L. Jiang, X. Liu, and H. Zhao, "Point transformer V2: Grouped vector attention and partition-based pooling," in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, New Orleans, LA, 2022, vol. 35, pp. 33330–33342.
- [13] Y. Liu, B. Tian, Y. Lv, L. Li, and F.-Y. Wang, "Point cloud classification using content-based transformer via clustering in feature space," *IEEE/CAA J. Autom. Sinica*, vol. 11, no. 1, pp. 231–239, 2024, doi: 10.1109/JAS.2023.123432.
- [14] L. Yao, W. Chen, C. Xiao, and Z. Zhang, "Point cloud classification based on transformer," *Comput. Electr. Eng.*, vol. 104, Art. no. 108431, Dec. 2022, doi: 10.1016/j.compeleceng.2022.108431.
- [15] Z. Wu et al., "3D ShapeNets: A deep representation for volumetric shapes," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Boston, MA, Jun. 2015, pp. 1912–1920.
- [16] M. Jaderberg, K. Simonyan, A. Zisserman, and K. Kavukcuoglu, "Spatial transformer networks," in *Adv. Neural Inf. Process. Syst. (NIPS)*, Montreal, 2015, vol. 28, pp. 2017–2025.
- [17] K. Zhang, M. Hao, J. Wang, C. W. de Silva, and C. Fu, "Linked dynamic graph CNN: Learning on point cloud via linking hierarchical features," arXiv preprint arXiv:1904.10014, Apr. 2019.
- [18] S. Wang, B. Z. Li, M. Khabsa, H. Fang, and H. Ma, "Linformer: Self-attention with linear complexity," arXiv preprint arXiv:2006.04768, Jun. 2020.
- [19] M. A. Uy, Q.-H. Pham, B.-S. Hua, T. Nguyen, and S.-K. Yeung, "Revisiting point cloud classification: A new benchmark dataset and classification model on real-world data," in *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*, Seoul, Oct. 2019, pp. 1588–1597.
- [20] Y. Pang, W. Wang, F. E. H. Tay, W. Liu, Y. Tian, and L. Yuan, "Masked autoencoders for point cloud self-supervised learning," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, Tel Aviv, Oct. 2022, pp. 604–621.
- [21] AlHija, M. A., Alqudah, H. J., and Dar-Othman, H., "Uncovering botnets in IoT sensor networks: A hybrid self-organizing maps approach," *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 34, no. 3, pp. 1840–1857, 2024. doi: 10.11591/ijeecs.v34.i3.pp1840-1857.
- [22] Abualhaj, M. M., Al-Shamayleh, A. S., Munther, A., Alkhatib, S. N., Hiari, M. O., and Anbar, M., "Enhancing spyware detection by utilizing decision trees with hyperparameter optimization," *Bulletin of Electrical Engineering and Informatics*, vol. 13, no. 5, pp. 3653–3662, 2024. doi: 10.11591/eei.v13i5.7939
- [23] Owida, H. A., El-Fattah, I. A., Abuowaida, S., Alshdaifat, N., Mashagba, H. A., Abd Aziz, A. B., et al., "A deep learning-based dual-branch framework for automated skin lesion segmentation and classification via dermoscopic images," *Scientific Reports*, vol. 15, no. 1, Art. no. 37823, 2025.
- [24] N. Alshdaifat, H. Abu Owida, Z. Mustafa, A. Aburomman, S. Abuowaida, A. Ibrahim, and W. Alsharafat, "Automated Blood Cancer Detection Models Based on EfficientNet-B3 Architecture and Transfer Learning," *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 36, no. 3, pp. 1731–1738, 2024, doi: 10.11591/IJEECS.V36.I3.PP1731-1738.
- [25] S. Abuowaida, H. Abu Owida, S. I. Shelash Mohammad, N. Alshdaifat, E. Abu Elsoud, R. Alazaidah, A. Vasudevan, and M. T. Alshurideh, "Evidence Detection in Cloud Forensics: Classifying Cyber-Attacks in IaaS Environments Using Machine Learning," *Data and Metadata*, vol. 4, Art. no. 699, 2025, doi: 10.56294/dm2025699.
- [26] S. Abuowaida, H. Abu Owida, D. M. Alsekait, N. Alshdaifat, D. S. Abdelminaam, and M. Alshinwan, "UltraSegNet: A Hybrid Deep Learning Framework for Enhanced Breast Cancer Segmentation and Classification

on Ultrasound Images,” *Computers, Materials & Continua*, vol. 83, no. 2, pp. 3303–3333, 2025, doi: 10.32604/cmc.2025.063470.

- [27] Z. Salah, H. Abu Owida, E. Abu Elsoud, E. Alhenawi, S. Abuowaida, and N. Alshdaifat, “An Effective Ensemble Approach for Preventing and Detecting Phishing Attacks in Textual Form,” *Future Internet*, vol. 16, no. 11, p. 414, 2024.
- [28] O. Alidmat, H. Abu Owida, U. K. Yusof, A. Almaghthawi, A. Altalidi, R. S. Alkhawaldeh, S. Abuowaida, N. Alshdaifat, and J. AlShaqsi, “Simulation of Crowd Evacuation in Asymmetrical Exit Layout Based on Improved Dynamic Parameters Model,” *IEEE Access*, vol. 13, pp. 7512–7525, 2025.
- [29] H. M. Turki, E. Al Daoud, G. Samara, R. Alazaidah, M. H. Qasem, M. Aljaidi, S. Abuowaida, and N. Alshdaifat, “Arabic Fake News Detection Using Hybrid Contextual Features,” *International Journal of Electrical and Computer Engineering*, vol. 15, no. 1, 2025.



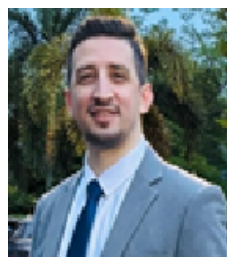
Radwan M. Batyha

Associate Professor in the Department of Computer Science. With extensive academic and professional experience, they also serve as an international consultant and trainer across the Arab region and Europe. Their work focuses on cutting-edge areas including Artificial Intelligence, Cybersecurity, Information Retrieval, and Machine Learning. They are actively involved in research, teaching, and professional training, contributing to both academic advancement and industry development.



HAMZA ABU OWIDA

received the Ph.D. from Keele University, U.K. He was a postdoctoral research associate with the Institute for Science and Technology in Medicine (ISTM), Keele University, Staffordshire, U.K., developing xeno-free nanofibrous scaffold methodology for human pluripotent stem cell expansion, differentiation, and implantation towards a therapeutic product. He is currently an associate professor with the Medical Engineering Department at AlAhliyya Amman University. He has published more than ten articles.



Hamza A. Mashagba

is a distinguished researcher specializing in communication engineering, with a particular focus on MIMO antenna design and textile antennas. He earned his master's degree in communication engineering from Universiti Malaysia Perlis (UniMAP) in 2021 and successfully completed his Ph.D. in the same field in December 2024. Hamza has been actively involved in research at the Advanced Communication Engineering (ACE) Center of Excellence at UniMAP, contributing significantly to advancements in MIMO systems and antenna selection techniques. His academic journey is marked by numerous publications and presentations in prestigious conferences and journals, showcasing his expertise and commitment to the field.



Asokan Vasudevan

is a distinguished academic at INTI International University, Malaysia. He holds multiple degrees, including a PhD in Management from UNITEN, Malaysia, and has held key roles such as Lecturer, Department Chair, and Program Director. His research, published in esteemed journals, focuses on business management, ethics, and leadership. Dr. Vasudevan has received several awards, including the Best Lecturer Award from Infrastructure University Kuala Lumpur and the Teaching Excellence Award from INTI International University. His ORCID ID is orcid.org/0000-0002-9866-4045.



Mohammad Faleh Ahmmad Huntie

is a professor of public administration at the University of Jordan since 1989. Ph.D holder earned in 1988 from Golden Gate University, California, USA. Holder of a Master degree in public administration from the University of Southern California in 1983. Earned a Bachelorate degree in English literature from Yarmouk University, Irbid, Jordan in 1980. Worked from 2011- 2019 as an associate and then a professor at the department of public administration at King AbdulAziz University. Granted an Award of distinction in publishing in excellent international journals. Teaching since 1989 until now. Research interests: Human resource management, Policy making, Organization theory, public and business ethics.



Azlan Abd. Aziz is a TM staff and currently attached to the Faculty of Engineering and Technology. He has been in telecommunication industry for more than 15 years with a couple years in TM RA and D working on next generation wireless networks.



Lara A. Al-Mashagba received the master Degree of Science in Business Analytics from Abu Dhabi School of Management, Abu Dhabi - UAE in 2020. She is currently pursuing the Ph.D. Degree in Artificial Intelligence and Software Engineering from

USM -MALAYSIA Her research interests in Prediction using Machine Learning



Suleiman Ibrahim Mohammad is a Professor of Business Management at Al al-Bayt University, Jordan, with more than 22 years of teaching experience. He has published over 400 research papers in prestigious journals. He holds a PhD in Financial

Management and an MCom from Rajasthan University, India, and a bachelor's in commerce from Yarmouk University, Jordan. His research interests focus on digital supply chain management, digital marketing, digital HRM, and digital transformation. His ORCID ID is <https://orcid.org/0000-0001-6156-9063>.