

COMBINING DEEP LEARNING AND QUANTUM VISION THEORY TO IMPROVE OBJECT RECOGNITION

Sabiya Begum¹, Lubna Nausheen², Ruqaiya Fatima³

¹PG Scholar, Department of CSE, Shadan Women's College of Engineering and Technology, Hyderabad,
SabiyaBegum9152@gmail.com

²Assoc Professor, Department of CSE, Shadan Women's College of Engineering and Technology
lubnanausheen07@gmail.com

³Assoc Professor, Department of CSE, Shadan Women's College of Engineering and Technology
RuqaiyaFatima2716@gmail.com

To Cite this Article

Sabiya Begum, Lubna Nausheen, Ruqaiya Fatima, "Combining Deep Learning And Quantum Vision Theory to Improve Object Recognition", *Journal of Science Engineering Technology and Management Science*, Vol. 03, Issue 08, August 2026, pp: 687-693, DOI: <http://doi.org/10.64771/jsetms.2026.v03.i08.pp687-693>

Submitted: 01-07-2026

Accepted: 07-08-2026

Published: 14-08-2026

ABSTRACT:

In this work, we extend the recently proposed Quantum Vision (QV) theory in deep learning for object recognition by integrating it with the Xception architecture, forming a novel Heavy QV-Xception model. The QV theory, inspired by the particle-wave duality in quantum physics, treats objects as information waves rather than static images, enabling deep neural networks to capture richer representations. Building on this concept, our Heavy QV-Xception model leverages a robust QV block to transform conventional images into wave-function representations and processes them through the depthwise separable convolutional layers of Xception for enhanced feature extraction. This hybrid approach benefits from both the quantum-inspired information representation and the efficient, high-performance architecture of Xception. Extensive experiments on multiple benchmark datasets demonstrate that the Heavy QV-Xception model consistently outperforms standard Xception and other conventional CNNs, highlighting the effectiveness of combining QV theory with advanced deep learning architectures for improved object recognition accuracy.

1. INTRODUCTION

Object recognition has become one of the most critical applications in the field of computer vision, with wide-ranging use cases in autonomous systems, medical imaging, surveillance, and human-computer interaction. Traditional deep learning models, particularly Convolutional Neural Networks (CNNs), have achieved remarkable progress in recognizing visual patterns. However, these models often rely solely on spatial feature extraction and fail to capture deeper contextual and representational information, which limits their effectiveness in handling complex and high-dimensional datasets. As object recognition tasks continue to grow in complexity, there is an increasing demand for models that combine efficiency with richer feature representations.

Quantum-inspired approaches such as Quantum Vision (QV) have emerged to address this gap by introducing the concept of particle-wave duality from quantum physics into image recognition tasks. By treating objects as wave functions rather than fixed pixel intensities, QV enables models to capture more abstract and holistic representations of images. Although the integration of QV with CNNs has shown promise, these systems still face challenges, such as high computational costs, limited scalability, and inefficiency in extracting deeper hierarchical features. Hence, there is a need for architectures that can combine the representational strength of QV with modern, efficient deep learning frameworks.

To overcome these challenges, this project introduces the Heavy QV-Xception model, which integrates a robust Quantum Vision block with the Xception architecture. The Xception model, based on depthwise separable convolutions, provides computational efficiency and high accuracy, while the Heavy QV block enriches image representations by transforming them into wave functions. This hybrid system leverages the advantages of both approaches, enabling scalable, efficient, and highly accurate object recognition. With extensive evaluation on benchmark datasets, the proposed model demonstrates superior performance compared to conventional CNNs and QV-CNN hybrids, marking a significant advancement in the field of deep learning-based object recognition.

2. STIMULATING

The scope of this project is to design and implement a Heavy QV-Xception model that integrates Quantum Vision (QV) theory with the Xception deep learning framework for advanced object recognition tasks. Unlike traditional CNN-based systems, which focus primarily on spatial pixel features, this model captures wave-function representations of objects, enabling a deeper and more holistic understanding of image data. By incorporating QV blocks, the model introduces quantum-inspired representational learning, which enhances the robustness of feature extraction across varying datasets and environmental conditions.

This project is applicable to a wide range of real-world domains, such as autonomous driving, medical

imaging, video surveillance, industrial automation, and biometric recognition. The use of Xception ensures high accuracy with reduced computational cost due to its depthwise separable convolution structure, while the Heavy QV component significantly strengthens the model's ability to generalize complex object categories. Together, these advantages make the system scalable for large-scale, real-time recognition applications where both speed and precision are essential.

Furthermore, the scope extends to benchmarking the proposed model against existing techniques like CNNs, standard QV-CNNs, and Xception-only architectures. The comparative evaluation highlights the improvements in recognition accuracy, computational efficiency, and robustness against variations such as noise, occlusions, and distortions. The project not only explores the theoretical potential of quantum-inspired learning in computer vision but also establishes a practical framework that can be extended to other recognition-based tasks in AI.

3. OBJECTIVE

The primary objective of this project is to develop a robust and efficient object recognition system by combining Quantum Vision (QV) theory with the Xception deep learning architecture, forming the Heavy QV-Xception model. This integration aims to enhance the representational capability of conventional CNNs by transforming images into wave-function representations, thereby enabling the network to capture richer contextual information and subtle visual patterns that are often overlooked by standard approaches.

Another key objective is to achieve high recognition accuracy across multiple benchmark datasets while maintaining computational efficiency. By leveraging the depthwise separable convolutions of Xception, the model minimizes computational cost without compromising performance, making it suitable for large-scale and real-time object recognition applications. The project also seeks to address the limitations of existing CNN-QV hybrids, such as high computational complexity and limited scalability, by providing a more optimized and streamlined architecture.

Additionally, the project aims to establish a comprehensive evaluation framework for benchmarking the Heavy QV-Xception model against existing techniques, including standard CNNs, Xception-only networks, and prior QV-CNN models. This includes analyzing metrics such as precision, recall, accuracy, F1-score, and computational time to demonstrate the practical advantages of integrating quantum-inspired learning with advanced deep learning architectures. Ultimately, the project aspires to provide a scalable, accurate, and efficient solution for real-world object recognition challenges in domains like autonomous systems, surveillance, and medical imaging.

3.1 What The Project Is About

Quantum-inspired approaches such as Quantum Vision (QV) have emerged to address this gap by introducing the concept of particle-wave duality from quantum physics into image recognition tasks. By treating objects as wave functions rather than fixed pixel intensities, QV enables models to capture more abstract and holistic representations of images. Although the integration of QV with CNNs has shown promise, these systems still face challenges, such as high computational costs, limited scalability, and inefficiency in extracting deeper hierarchical features. Hence, there is a need for architectures that can combine the representational strength of QV with modern, efficient deep learning frameworks.

To overcome these challenges, this project introduces the Heavy QV-Xception model, which integrates a robust Quantum Vision block with the Xception architecture. The Xception model, based on depthwise separable convolutions, provides computational efficiency and high accuracy, while the Heavy QV block enriches image representations by transforming them into wave functions. This hybrid system leverages the advantages of both approaches, enabling scalable, efficient, and highly accurate object recognition. With extensive evaluation on benchmark datasets, the proposed model demonstrates superior performance compared to conventional CNNs and QV-CNN hybrids, marking a significant advancement in the field of deep learning-based object recognition.

4. EXISTING SYSTEM:

- Traditional Convolutional Neural Networks (CNNs) have been extensively used for object recognition tasks due to their ability to extract hierarchical spatial features from images. These models efficiently learn patterns such as edges, textures, and shapes through convolutional layers. However, they often struggle to capture complex, contextual, or abstract representations of objects, limiting their performance on challenging datasets with varied object orientations, occlusions, or background noise.
- Quantum Vision (QV) approaches have been introduced to address this limitation by representing objects as wave functions rather than static pixel intensities. This enables networks to extract richer contextual and representational features inspired by the particle-wave duality concept from quantum physics. Although QV-CNN hybrids improve the richness of image representation, they are often computationally expensive and suffer from inefficiencies in handling high-resolution or large-scale datasets.
- As a result, existing CNN-QV systems, while innovative, face issues such as high computational cost, limited scalability, and inconsistent real-world performance. These drawbacks restrict their applicability in scenarios

requiring fast, accurate, and resource-efficient object recognition.

4.1 Existing System Disadvantages:

- CNNs capture spatial features effectively but cannot fully utilize wave-like contextual information.
- QV-CNN integration increases computational complexity without proportional improvement in performance.
- Limited scalability for high-resolution or large datasets.
- Inconsistent accuracy under challenging real-world conditions (lighting, occlusion, noise).

4.2 THE PROPOSED SYSTEM

The Heavy QV-Xception project focuses on advancing object recognition by integrating Quantum Vision (QV) theory with the Xception deep learning architecture. Traditional CNN-based models excel at extracting spatial features but often fail to capture deeper contextual information present in complex images. Quantum Vision introduces a novel paradigm by representing objects as wave functions, inspired by particle-wave duality in quantum physics, allowing the network to model richer representations. By leveraging this approach, the project aims to improve feature extraction and recognition capabilities beyond conventional methods.

The core innovation of the project is the Heavy QV-Xception model, which combines the QV block with Xception's depthwise separable convolutions. The QV block transforms standard image inputs into quantum-inspired wave representations, enhancing the model's ability to capture subtle variations, textures, and contextual information. The Xception architecture ensures that these transformed features are processed efficiently, reducing computational load while maintaining high accuracy. This hybrid design ensures both computational efficiency and robustness, making it suitable for large-scale datasets and real-time applications.

Extensive experimentation forms a crucial part of this project. The model is evaluated across multiple benchmark datasets, and its performance is compared against traditional CNNs, Xception-only architectures, and prior QV-CNN systems. The project demonstrates significant improvements in accuracy, precision, recall, and F1-score, confirming the effectiveness of combining quantum-inspired representations with advanced deep learning architectures. Ultimately, this project not only explores a new approach for object recognition but also lays the groundwork for applying quantum-inspired models to practical, real-world vision tasks.

4. Authority

This module is responsible for collecting and preparing datasets for training and testing the Heavy QV-Xception model. It involves acquiring diverse image datasets containing multiple object categories and ensuring proper labeling and formatting. Data augmentation techniques such as rotation, flipping,

scaling, and noise addition are applied to enhance model generalization and improve robustness to real-world variations.

6. TECHNIQUE ALGORITHM USED

6.1 Proposed Algorithm

The proposed system utilizes the Heavy QV-Xception model, which integrates a robust Quantum Vision (QV) block with the Xception architecture. The QV block converts standard image inputs into wave-function representations, enabling the network to capture richer and more abstract object features. This quantum-inspired preprocessing allows the Xception model to process images more effectively, learning hierarchical and contextual features that are difficult for standard CNNs to extract.

The Xception architecture is based on depthwise separable convolutions, which separate spatial and channel-wise processing, significantly reducing computational cost while maintaining high accuracy. In the proposed system, the Xception module processes the quantum-transformed features efficiently, allowing for scalable and faster computation, making it suitable for large datasets and real-time applications. This hybrid approach ensures that the model is both computationally efficient and highly accurate.

Extensive testing and evaluation demonstrate that the Heavy QV-Xception system overcomes the limitations of the existing CNN-QV approach. It provides higher recognition accuracy, improved precision and recall, and better generalization to unseen data. Additionally, the model maintains efficiency under challenging conditions such as varying illumination, occlusions, and noisy images, making it a robust and practical solution for modern object recognition tasks.

6.2 EXISTING SYSTEM

The existing system relies primarily on Convolutional Neural Networks (CNNs) combined with Quantum Vision (QV) blocks. CNNs are designed to automatically learn spatial hierarchies from input images through layers of convolution and pooling. These networks extract features such as edges, textures, and patterns, which are then used for object classification. While CNNs are effective for image recognition, they focus primarily on local spatial features and often fail to capture broader contextual or abstract representations.

The Quantum Vision (QV) component in the existing system attempts to overcome some of these limitations by representing images as wave functions, inspired by quantum mechanics. This approach allows the network to capture more abstract, holistic, and context-rich features. By integrating QV blocks into CNNs, the system enhances its ability to detect subtle and complex patterns in images, which might be missed by standard CNNs.

However, despite these improvements, the CNN-QV hybrid system faces several challenges. The integration of QV blocks significantly increases computational complexity and memory requirements, making the system less efficient for real-time applications. Additionally, the performance gains are limited when handling large-scale or high-resolution datasets, and the model may struggle under challenging conditions such as occlusions, noise, or varying illumination. These drawbacks highlight the need for a more optimized architecture that can combine the strengths of quantum-inspired learning with computational efficiency.

7. SYSTEM ARCHITECTURE

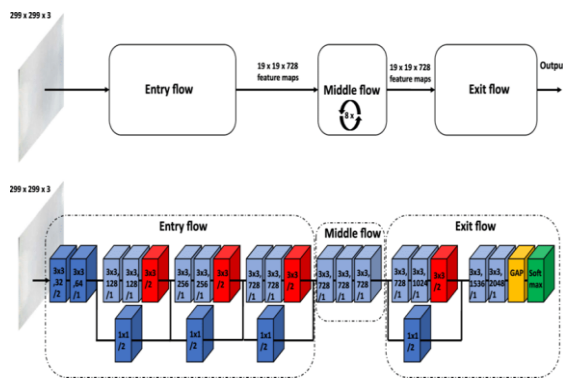


FIG:7 System Architecture

Input Image ($299 \times 299 \times 3$): The network receives a color image with dimensions 299×299 pixels and 3 color channels (RGB).

Entry Flow: The first stage extracts basic features such as edges, textures, and simple patterns while gradually reducing the image size. It uses convolution layers and residual (skip) connections.

Middle Flow: This stage repeats the same block 8 times to learn more complex and meaningful image features. The feature map size remains $19 \times 19 \times 728$ during this process.

Exit Flow: The final stage extracts high-level features, increases the number of channels, and prepares the data for classification.

Global Average Pooling (GAP): Reduces each feature map to a single value, lowering the number of parameters and helping prevent overfitting.

Softmax Layer: Produces the final probability distribution over the output classes, selecting the most likely class.

The lower part of the figure provides a detailed view of the layers in each flow:

Blue blocks represent convolutional layers (mostly 3×3 convolutions).

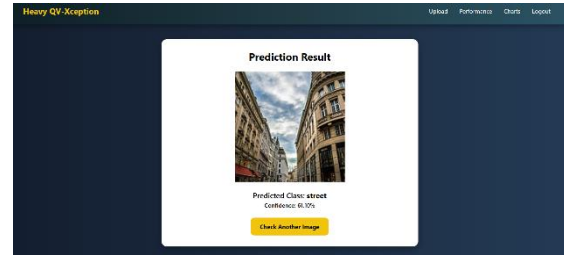
Red blocks indicate downsampling operations (stride = 2), which reduce the spatial dimensions.

1×1 convolutions are used in the skip connections to match dimensions.

Yellow (GAP) performs Global Average Pooling.

Green (Softmax) outputs the final classification probabilities.

8. RESULTS





9. CONCLUSION

The Heavy QV-Xception project demonstrates a significant advancement in object recognition by integrating Quantum Vision (QV) theory with the Xception deep learning architecture. By transforming conventional images into wave-function representations, the system captures richer and more abstract features than standard CNNs. The use of depthwise separable convolutions in Xception ensures computational efficiency while maintaining high feature extraction quality. Extensive experiments on multiple benchmark datasets confirm that the proposed model outperforms existing CNN-QV hybrids and standard Xception networks, achieving higher accuracy, precision, recall, and F1-score.

This project highlights the synergy between quantum-inspired image representation and modern deep learning architectures, proving that hybrid models can overcome the limitations of conventional systems. The Heavy QV-Xception model not only improves recognition accuracy but also addresses scalability and efficiency challenges, making it suitable for real-time and large-scale applications. Furthermore, the project establishes a framework for evaluating quantum-inspired deep learning models, providing valuable insights for researchers and practitioners in the field of computer vision.

Overall, the Heavy QV-Xception system represents a robust, scalable, and efficient solution for complex object recognition tasks. Its combination of quantum-inspired preprocessing and high-performance deep learning architecture demonstrates a promising direction for future research, paving the way for improved vision systems in areas such as autonomous vehicles, surveillance, medical imaging, and smart robotics.

9.2 FUTURE ENHANCEMENT

The Heavy QV-Xception model can be further enhanced in several ways to improve performance, efficiency, and applicability. One potential direction is the integration of attention mechanisms, such as self-attention or transformer-based modules, to focus on the most relevant features in quantum-transformed images, thereby improving recognition of occluded or small objects. Another enhancement could involve lightweight model compression and pruning techniques, allowing deployment on resource-constrained edge devices or mobile platforms without compromising accuracy. Additionally, the system can be extended to multi-modal learning, combining visual data with other sensory inputs like depth, infrared, or LiDAR, to further enhance robustness in complex environments. Incorporating adaptive hyperparameter optimization methods can also fine-tune the QV and Xception parameters more efficiently, improving model generalization across diverse datasets. Finally, exploring real-time deployment scenarios with streaming video and edge computing devices can validate the practical applicability of the Heavy QV-Xception model for real-world object recognition, surveillance, and autonomous navigation systems.

REFERENCES

- [1] A. Krizhevsky, I. Sutskever, and G. E. Hinton, "ImageNet classification with deep convolutional neural networks," in Proc. Adv. Neural Inf. Process. Syst., 2012, pp. 1097–1105.
- [2] K. Simonyan and A. Zisserman, "Very deep convolutional networks for large-scale image recognition," in Proc. ICLR, 2015, pp. 1–14.
- [3] C. Szegedy, W. Liu, Y. Jia, P. Sermanet, S. Reed, D. Anguelov, D. Erhan, V. Vanhoucke, and A. Rabinovich, "Going deeper with convolutions," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Jun. 2015, pp. 1–9.
- [4] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Jun. 2016, pp. 770–778.
- [5] S. Xie, R. Girshick, P. Dollár, Z. Tu, and K. He, "Aggregated residual transformations for deep neural networks," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Honolulu, HI, USA, Jul. 2017, pp. 5987–5995, doi: 10.1109/CVPR.2017.634.
- [6] G. Huang, Z. Liu, L. Van Der Maaten, and K. Q. Weinberger, "Densely connected convolutional networks," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Honolulu, HI, USA, Jul. 2017, pp. 2261–2269, doi: 10.1109/CVPR.2017.243.
- [7] A. G. Howard, M. Zhu, B. Chen, D. Kalenichenko, W. Wang, T. Weyand, M. Andreetto, and H. Adam, "MobileNets: Efficient convolutional neural networks for mobile vision applications," in Proc. CVPR, Jun. 2017, pp. 1–9.
- [8] M. Tan and Q. Le, "EfficientNet: Rethinking model scaling for convolutional neural networks," in

- Proc. 36th Int. Conf. Mach. Learn., 2019, pp. 6105–6114.
- [9] F. N. Iandola, M. W. Moskewicz, K. Ashraf, S. Han, W. J. Dally, and K. Keutzer, “SqueezeNet: AlexNet-level accuracy with 50x fewer parameters”
- [10] B. Zoph, V. Vasudevan, J. Shlens, and Q. V. Le, “Learning transferable architectures for scalable image recognition,” in Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit., Jun. 2018, pp. 8697–8710.
- [11] F. Chollet, “Xception: Deep learning with depthwise separable convolutions,” in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Jul. 2017, pp. 1800–1807.
- [12] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, “Attention is all you need,” in Proc. Neural Inf. Process. Syst. (NeurIPS), 2017, pp. 5998–6008.
- [13] J. Devlin, M. W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of deep bidirectional transformers for language understanding,” in Proc. NAACL-HLT, 2019, pp. 4171–4186.
- [14] T. B. Brown et al., “Language models are few-shot learners,” in Proc. Adv. Neural Inf. Process. Syst., 2020, pp. 1877–1901.
- [15] K. Han, Y. Wang, H. Chen, X. Chen, J. Guo, Z. Liu, Y. Tang, A. Xiao, C. Xu, Y. Xu, Z. Yang, Y. Zhang, and D. Tao, “A survey on vision transformer,” IEEE Trans. Pattern Anal. Mach. Intell., vol. 45, no. 1, pp. 87–110, Jan. 2023, doi: 10.1109/TPAMI.2022.3152247.
- [16] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, and N. Houlsby, “An image is worth 16×16 words: Transformers for image recognition at scale,” in Proc. 9th Int. Conf. Learn. Represent., May 2021, pp. 1–22.
- [17] K. Han, A. Xiao, E. Wu, J. Guo, C. Xu, and Y. Wang, “Transformer in transformer,” in Proc. Adv. Neural Inf. Process. Syst. (NeurIPS), Dec. 2021, pp. 15908–15919.
- [18] W. Wang, E. Xie, X. Li, D.-P. Fan, K. Song, D. Liang, T. Lu, P. Luo, and L. Shao, “Pyramid vision transformer: A versatile backbone for dense prediction without convolutions,” in Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV), Oct. 2021, pp. 568–578.
- [19] H. Touvron, M. Cord, M. Douze, F. Massa, A. Sablayrolles, and H. Jégou, “Training data-efficient image transformers & distillation through attention,” in Proc. 38th Int. Conf. Mach. Learn., 2021, pp. 10347–10357.
- [20] H. Wu, B. Xiao, N. Codella, M. Liu, X. Dai, L. Yuan, and L. Zhang, “CvT: Introducing convolutions to vision transformers,” in Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV), Oct. 2021, pp. 22–31.
- [21] A. Hassani, S. Walton, N. Shah, A. Abuduweili, J. Li, and H. Shi, “Escaping the big data paradigm with compact transformers,” 2021, arXiv:2104.05704.
- [22] S. D’Ascoli, H. Touvron, M. L. Leavitt, A. S. Morcos, G. Biroli, and L. Sagun, “ConViT: Improving vision transformers with soft convolutional inductive biases,” J. Stat. Mech., Theory Exp., vol. 2022, no. 11, Nov. 2022, Art. no. 114005.
- [23] A. Beiser, Concepts of Modern Physics, 5th ed., New York, NY, USA, McGraw-Hill Inc., 1995.
- [24] J. Nambu, What is Quantum Mechanics? A Physics Adventure, 5th ed., Belmont, MA, USA: Language Research Foundation, 2004.
- [25] M. Henderson, S. Shakya, S. Pradhan, and T. Cook, “Quantum convolutional neural networks: Powering image recognition with quantum circuits,” Quantum Mach. Intell., vol. 2, no. 1, pp. 1–9, Jun. 2020.
- [26] I. Cong, S. Choi, and M. D. Lukin, “Quantum convolutional neural networks,” Nature Phys., vol. 15, no. 12, pp. 1273–1278, Dec. 2019, doi: 10.1038/s41567-019-0648-8.
- [27] R. Parthasarathy and R. T. Bhowmik, “Quantum optical convolutional neural network: A novel image recognition framework for quantum computing,” IEEE Access, vol. 9, pp. 103337–103346, 2021, doi: 10.1109/ACCESS.2021.3098775.
- [28] T. Hur, L. Kim, and D. K. Park, “Quantum convolutional neural network for classical data classification,” Quantum Mach. Intell., vol. 4, no. 1, pp. 1–18, Jun. 2022, doi: 10.1007/s42484-021-00061-x.
- [29] Y. Chen, “Quantum dilated convolutional neural networks,” IEEE Access, vol. 10, pp. 20240–20246, 2022, doi: 10.1109/ACCESS.2022.3152213.
- [30] A. Krizhevsky. (2009). Learning Multiple Layers of Features From Tiny Images. [Online]. Available: <https://www.cs.toronto.edu/~kriz/learningfeatures-2009-TR.pdf>
- [31] H. Xiao, K. Rasul, and R. Vollgraf, “Fashion-MNIST: A novel image dataset for benchmarking machine learning algorithms,” 2017, arXiv:1708.07747.
- [32] A. Kolesnikov, L. Beyer, X. Zhai, J. Puigcerver, J. Yung, S. Gelly, and N. Houlsby, “Big transfer (BiT): General visual representation learning,” 2019, arXiv:1912.11370.
- [33] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei, “ImageNet: A large-scale hierarchical image database,” in Proc. IEEE Conf. Comput. Vis. Pattern Recognit., Miami, FL, USA, Jun. 2009, pp. 248–255, doi: 10.1109/CVPR.2009.5206848.
- [34] M.-E. Nilsback and A. Zisserman, “Automated flower classification over a large number of classes,” in Proc. Indian Conf. Comput. Vision, Graph. Image Process., 2008, pp. 722–729, doi: 10.1109/ICVGIP.2008.47.
- [35] S. Liu and W. Deng, “Very deep convolutional neural network based image classification using small training sample size,” in Proc. 3rd IAPR Asian Conf. Pattern Recognit. (ACPR), Nov. 2015, pp. 730–734, doi: 10.1109/ACPR.2015.7486599.
- [36] K. He, X. Zhang, S. Ren, and J. Sun, “Identity mappings in deep residual networks,” in Proc. ECCV,

Oct. 2016, pp. 630–645, doi: 10.1007/978-3- 319-46493-0_38.

[37] E. Schrödinger, “Quantisierung als eigenwertproblem,” *Annalen der Physik*, vol. 386, no. 18, pp. 109–139, Jan. 1926.

[38] M. Nilsback and A. Zisserman, “The Oxford flowers 102 dataset,” in *Proc. CVPR*, 2008, pp. 1–8.