

From Pixels to Precision: Techniques for Image Dataset Refinement

Raji.N1, Dr.S. Manohar2, Dr.V. Rajasekar3

¹Research Scholar, Department of Computer Science and Engineering, SRM Institute of Science and Technology, Vadapalani Campus
Chennai, India. E-mail: rn5525@srmist.edu.in

²Assistant Professor, Department of Computer Science and Engineering, SRM Institute of Science and Technology, Vadapalani Campus
Chennai, India. manohars@srmist.edu.in

³Associate Professor, Department of Computer Science and Engineering, SRM Institute of Science and Technology, Vadapalani Campus
Chennai, India. rajasekv2@srmist.edu.in

ARTICLE INFO

ABSTRACT

Received: 11 Oct 2024

Revised: 13 Dec 2024

Accepted: 23 Dec 2024

Preparing and processing image datasets is a crucial step in deep learning and computer vision applications. However, many challenges arise in creating high-quality image datasets that can make deep learning models more accurate and generic. We provide a thorough review of best practices in this document for preparing and processing image datasets, from raw data collection to final dataset creation. We discuss various techniques and tools that can be used to enhance image quality, reduce noise, and address issues such as class imbalance and bias. Additionally, we provide practical tips for data cleaning, normalization, and augmentation, which are essential for improving the accuracy and robustness of machine learning procedures. Consequently, we demonstrate the effectiveness of our approach by applying it to real-time images and converting them to datasets suitable for several computer vision applications, such as object detection and classification.

Keywords: Image, Preprocessing, Deep Learning, Dataset.

INTRODUCTION

The quality of an image dataset is crucial for the performance of deep learning models that rely on it. Image datasets are used to train and test deep learning models that can recognize and classify objects, detect anomalies, and perform a wide range of other tasks. There are many applications that rely on image datasets, identifying objects within an image or video stream, and recognizing their type or category, identifying and recognizing people in photographs or videos based on their facial features, X-rays, Computed tomography, and Magnetic resonance imaging, among others, to identify medical disorders, analyzing images from cameras mounted on self-driving cars to help them navigate and avoid obstacles, monitoring security cameras to detect and identify potential threats or suspicious activity, finding similar images based on a given image or searching for images based on specific features or attributes, overlaying virtual images onto real-world objects or environments, creating immersive environments by rendering 3D images in real-time, recommending clothing items based on user preferences and visual features and generating or manipulating images using AI algorithms to create new artistic expressions.

However, creating high-quality image datasets can be challenging due to various issues such as image noise, class imbalance, and bias. In this paper, we present best practices for preparing and processing [13-17] raw data to create a high-quality image dataset. We begin with data collection techniques, including image selection and data augmentation, to ensure that the dataset is diverse and representative of the target population.

We begin by discussing the raw data collection process, which is an essential step in creating high-quality image datasets. We provide guidance on how to collect representative and diverse data that can capture all possible variations in the target domain. Next, we discuss various techniques and tools that can be used to enhance image quality, reduce noise, and address issues such as class imbalance and bias. We provide examples of how to apply these techniques to different types of image datasets.

We then delve into practical tips for data cleaning, normalization, and augmentation. Data cleaning involves removing outliers and incorrect data points, while normalization involves standardizing image features [23] such as brightness and contrast. Augmentation [19-22,24] involves creating new images from existing ones by applying transformations such as rotation, scaling, and flipping. We provide examples of how these techniques can be used to enhance image datasets for specific computer vision tasks.

METHODOLOGY

While existing image datasets have been instrumental in advancing computer vision research, there are several important reasons why new image datasets[26-28] are still necessary and valuable .Firstly, new image datasets can address gaps in existing datasets. For example, existing datasets may be biased towards certain demographic groups, geographical regions, or types of images. New datasets can provide a more diverse and representative sample, which can help to improve the generalizability of computer vision models.

Secondly, new datasets can reflect emerging trends or new applications in computer vision. For example, as the use of autonomous vehicles becomes more widespread, new datasets that focus on identifying road hazards or pedestrians in different lighting conditions may be needed. Thirdly, new datasets can incorporate new types of information or features. For example, recent datasets have included additional information such as depth maps, 3D models, or audio, which can help to improve the accuracy and robustness of computer vision models.

Lastly, new datasets can also provide an opportunity for researchers to benchmark in addition to performance comparison of diverse computer vision models. As new datasets are introduced, researchers can use it to assess the results of existing methodologies and develop new models that perform better on the new data. Overall, new image datasets can help to address existing limitations, reflect new applications, incorporate new features [11,12], and provide opportunities for benchmarking and comparison, all of which can help to advance computer vision research and applications.

A. STEPS IN PREPARING A DATASET

Preparing an image dataset involves several steps, which may vary depending on the specific use case and the type of images being used. Here are some general steps you can follow:

- 1.Determine the objective of the dataset: Decide on the goal of the dataset, such as classification, detection, or segmentation. Identify the types of images needed for the task and the number of images required.
- 2.Collect images: Collect images that are relevant to the objective of the dataset. Images can be collected through various sources such as public image repositories, search engines, social media platforms, or by capturing new images.
- 3.Clean the images: Clean the images by removing duplicates, blurry images, and low-quality images. The images should be standardized in terms of size, format, and resolution.
- 4.Label the images: Assign labels to each image based on the objective of the dataset. This is typically done manually by humans, but can also be done using automated labelling techniques.
- 5.Split the dataset: Create training, validation, and testing sets from the dataset. The validation set is utilized to smooth the model once it has been trained, while the testing set serves as a way to assess the model's effectiveness.
- 6.Augment the images: Augment the images by using various techniques to generate additional images. This can include techniques such as rotation, flipping, cropping, and changing the brightness and contrast.
- 7.Pre-process the images: Pre-process the images by applying transformations such as resizing, normalization, and cropping. This step ensures that the images are in a format that can be used by the model.
- 8.Save the dataset: Save the dataset in a format that can be easily loaded by the model. Common formats include JPEG, PNG, or TensorFlow's native format.
- 9.Document the dataset: Document the dataset by providing information about the source of the images, the labelling process, and any other relevant information. This information will be useful for other researchers who may want to use the dataset.

B. COLLECTING AND CLEANING THE IMAGES

Images are collected from a rice field located in Alagoor, a village near Padapai, Kancheepuram district and also from Theerthanagiri, a village from Cuddalore district. There are several methods to clean images in a dataset, depending on the specific needs of the dataset and the desired outcome. Here are some common methods for cleaning images:

1. Manual inspection and removal: This method involves visually inspecting each image and manually removing any images that do not meet the necessary criteria. This can be time-consuming but is often the most accurate method.
2. Automatic removal using image processing techniques: This method involves using image processing techniques such as edge detection, color thresholding, and image segmentation to automatically identify and remove images that do not meet the necessary criteria. This method can be faster than manual inspection, but may also remove some images that should be kept.
3. Machine learning-based classification: This method involves using machine learning algorithms to classify images as either relevant or irrelevant to the task at hand. Images that are classified as irrelevant can then be removed from the dataset. This method can be very effective but requires a large, accurately labelled dataset for training the machine learning algorithm.
4. Combination of manual and automatic methods: This method involves using a combination of manual inspection and automatic image processing techniques to clean the dataset. This can help to speed up the cleaning process while still maintaining a high level of accuracy.

C. LABELLING THE IMAGES

Labelling the images is an important step in preparing a dataset for training a machine learning model. Labelling involves assigning a specific class or category to each image. Images can be labelled in a few different ways. A human annotator manually applies a label to each image as part of the manual labelling process. This approach takes a lot of effort and can make mistakes, yet it can be useful for ensuring high-quality labels and for tasks where there is no existing dataset available.

Semi-automatic labelling involves using a combination of human annotators and automated tools to label images. For example, an annotator might manually label a subset of the images, and then use a machine learning model to automatically label the rest of the images based on the patterns learned from the manually labeled subset. Automatic labelling involves using automated tools to label images based on pre-defined rules or patterns. For example, a model might be trained to recognize specific objects in images and automatically assign labels based on the presence of those objects.

Automatic labelling of images is the process of using machine learning algorithms to automatically assign labels to images based on pre-defined rules or patterns. This can be useful for tasks where there are a large number of images to be labelled, or for tasks where manual labelling is not feasible. There are different approaches to automatic image labelling, depending on the specific problem and the available data.

Object detection involves training a model to detect specific objects in images and assigning labels based on the presence or absence of those objects. For example, a model might be trained to detect cars in images and label images as "car" or "not car" based on whether a car is detected.

An picture is broken down into sections or areas by image segmentation and it also involves assigning labels to each segment based on the content of that region. For example, a model might be trained to segment images of clothing and assign labels based on the type of clothing in each segment (e.g., shirt, pants, shoes).

Using trained models that have been trained on large datasets to automatically label new images is part of transfer learning. For example, a pre-trained model that has been trained to recognize faces can be used to automatically label images of faces in a new dataset.

D. SPLITTING THE DATASET

A crucial step in getting a dataset ready for machine learning is dividing it into training, validation, and testing sets. To be sure that the model can generalise well to fresh, untested data, the data is separated rather than just memorizing

the training data. The typical approach to splitting the data is to use a proportion of 60:20:20 is maintained training, testing, and validation, respectively. This means that 60% of the data is used for training the model, 20% is used for validating the model during training, and 20% is used for testing the final model after training.

The validation set is used to assess the model's performance while training and modify the model's input variables to enhance performance. The training set serves as the basis for training the model on the data. The testing set is used to assess the model's ultimate performance on fresh, untested data. Making sure the data are accurately representing the population the model will be applied to is crucial when splitting the data. This means that the data should be randomly sampled and should include a diverse range of examples from each class or category in the dataset.

Additionally, it's important to ensure that the data is split in a way that avoids any data leakage between the training, validation, and testing sets. Data leakage can occur when information from the testing or validation sets is inadvertently used to train the model, which can lead to overfitting and poor generalization performance. To avoid data leakage, the data should be split before any preprocessing or feature engineering steps are taken, and the same preprocessing steps should be applied consistently across all sets.

E. IMAGE AUGMENTATION

By producing additional images using the original photographs, image augmentation artificially increases the size of a dataset. The purpose of image augmentation is to upsurge the range of the dataset and improve the performance of the model by preventing overfitting. An picture can be rotated either horizontally or vertically to produce a new image with a similar appearance to the original but a different orientation. This can help to increase the diversity of the dataset.

An image can be rotated by a tiny amount to produce a new representation that is identical to the original but distinct orientation. This can help to increase the diversity of the dataset. Cropping an image to focus on a specific region of interest can create a new image that is identical to the original but distinct composition. This can help to increase the diversity of the dataset.

Scaling an image by a small amount can create a new image that is identical to the original but distinct size. This can help to increase the diversity of the dataset. A new image that is identical to the original but with a little different texture can be produced by adding random noise to an existing image. This can help to increase the diversity of the dataset.

A slight change in an image's colour palette might result in a new image that resembles the original while using a different colour scheme. This may contribute to the dataset's becoming more diverse. A new image that is similar to the original but has minor distortions that imitate the organic changes in real-world photographs can be produced by applying elastic transformations to an image.

By using these techniques for image augmentation, it is possible to create a larger, more diverse dataset that can expand the working nature of the model. However, it is significant to avoid over-augmenting the images, as this can lead to a decrease in performance.

F. IMAGE PREPROCESSING

Image preprocessing refers to the set of techniques used to prepare an image for further analysis or processing. Image preprocessing techniques are applied to improve the quality of the image, reduce noise, and enhance the features of interest in the image. By averaging the pixel values in a small area, picture smoothing is a technique for reducing noise in images. This can be achieved using filters such as the mean filter, median filter, or Gaussian filter. Image smoothing is often used as a first step in image preprocessing to improve the signal-to-noise ratio.

The Gaussian filter is a linear filter that convolves the input image with a Gaussian kernel. The Gaussian kernel is a matrix that contains weights that correspond to the values of the Gaussian function for each pixel in the filter. Dimensions of the kernel and Gaussian function's standard deviation determine the amount of smoothing that is applied to the image. The Gaussian filter works by convolving the kernel with the image. Each pixel in the output image is calculated by taking the weighted sum of the surrounding pixels in the input image, with the weights determined by the values in the Gaussian kernel. This process effectively blurs the image and reduces the high-frequency noise.

Gaussian filtering is a widely used technique in image processing and computer vision, with applications ranging from edge detection and feature extraction to image enhancement and noise reduction. The effectiveness of the technique depends on dimensions of the kernel and Gaussian function's standard deviation, which must be carefully chosen to balance the amount of smoothing and the preservation of important image features.

By enhancing the image's contrast, brightness, or sharpness, image enhancement [1,11] methods are utilised to enhance the visual quality of the image.. These methods can be applied globally or locally to specific regions of interest in the image. Commonly used techniques include histogram equalization, contrast stretching, and unsharp masking. A technique used in the processing of digital images to improve contrast in an image is histogram equalisation.

Histogram equalisation aims to redistribute an image's intensity values in such a way that they are spread more evenly throughout the whole intensity value range. An image's histogram is a graph that displays the frequency at which each intensity value appears in the image.. In an ideal image, the histogram should be evenly distributed across the entire range of intensity values. However, in many real-world images, the histogram is often skewed towards the low or high end of the intensity range, resulting in poor contrast and visibility of details.

Using a transformation function, histogram equalisation converts an image's intensity values to a new set of values. The transformation function is made to distribute the image's intensity values across the entire range of values, producing an output image with a more equal distribution of intensity values.

Unsharp masking is a technique used in image processing to sharpen digital images by enhancing edges and increasing contrast. The process involves subtracting a blurred version of the image from the original image, resulting in an image with enhanced edges and increased contrast. Instead of amplifying the high-frequency components of the image, the process of unsharp masking reduces the low-frequency components, thereby enhancing the edges and making the image appear sharper.

Image normalization techniques are used to normalize the pixel values in an image to a common scale, typically to improve the robustness of subsequent processing steps. This can be achieved by subtracting the mean value of the image or dividing by the standard deviation. Image resizing techniques are used to resize an image to a desired resolution or scale. This can be achieved using interpolation techniques such as nearest-neighbor, bilinear, or bicubic interpolation.

The process of segmenting an image into various areas or segments, each of which refers to a different unit or portion of the image, is known as image segmentation [5-10]. Image segmentation aims to reduce complexity and/or transform an image's representation into something more relevant and understandable. There are various techniques and algorithms for image segmentation. Using the thresholding technique, each pixel in the image is classified either as foreground or background depending upon if the intensity value is higher or lower than the threshold. This technique is simple and effective for images with high contrast, but may not work well for images with complex or uneven lighting.

Region-based segmentation technique involves grouping pixels into regions based on their similarity in color, texture, or other features. This technique can be more robust than thresholding for images with complex lighting, but may be computationally intensive. Edge-based segmentation technique involves detecting edges in the image and using them to delineate object boundaries. This technique can be effective for images with well-defined edges, but may not work well for images with low contrast or noise.

Watershed segmentation [2,3] technique involves treating the image as a topographic map and flooding it from multiple sources to create separate basins, each of which resembles to a different object or region. This technique can be effective for images with irregular shapes and multiple objects, but may be sensitive to noise and other artifacts.

Image cropping techniques are used to remove unwanted portions of an image or focus on specific regions of interest. This can be achieved by selecting a rectangular region of interest or using more advanced techniques such as object detection or segmentation to identify the region of interest. Image registration techniques are used to align multiple images of the same scene or object taken at different times or from different viewpoints. This can be achieved by finding corresponding features in the images and transforming them to a common coordinate system.

RESULTS

Images captured using various smart devices were subjected to cleaning, labelling, splitting, augmentation and preprocessing and were tested with various algorithms to measure the quality of the image dataset. To compare the performance of the classification algorithms, it's common to use valuation metrics such as accuracy, precision, recall, F1 score, and AUC.

These metrics can deliver insight into the trade-offs between different algorithms in terms of their ability to correctly classify positive and negative examples, handle imbalanced data, and avoid overfitting. It's also important to consider factors such as computational complexity, interpretability, and scalability when selecting an algorithm for a particular task.

Table1. Accuracy and Precision comparison

Algorithm	Accuracy	Precision
Logistic Regression	0.85	0.86
Decision Trees	0.78	0.79
Random Forests	0.88	0.88
SVM	0.87	0.87
KNN	0.80	0.81
Neural Networks	0.90	0.90

Table2. Recall, F1 Score and Area Under Curve (AUC) Comparison

Algorithm	Recall	F1 Score	AUC-ROC
Logistic Regression	0.84	0.85	0.91
Decision Trees	0.77	0.78	0.83
Random Forests	0.88	0.88	0.94
SVM	0.86	0.87	0.92
KNN	0.79	0.80	0.87
Neural Networks	0.90	0.90	0.95

DISCUSSION

Preparing an image dataset for use in computer vision tasks is a critical and challenging process. To ensure high-quality results, it is important to follow best practices for image collection, annotation, cleaning, and preprocessing. This includes paying attention to factors such as image resolution, lighting, color balance, and data augmentation techniques. Additionally, selecting an appropriate dataset size and composition, and ensuring that the dataset is representative of the target population, are essential considerations for achieving accurate and reliable results.

Careful attention to these factors can help to ensure that the resulting image dataset is well-suited to the intended application and can help to maximize the performance of computer vision models. Our approach provides best practices for addressing common issues in image dataset preparation, including noise, class imbalance, and bias. We hope that our guidance will help researchers and practitioners in creating high-quality image datasets that can be used to train and test deep learning models in a wide range of applications.

REFERENCES

- [1] Mira Hayati, Kahlil Muchtar, Roslidar, Novi Maulina, Irfan Syamsuddin, Gregorius Natanael Elwirehardja, Bens Pardamean, Impact of CLAHE-based image enhancement for diabetic retinopathy classification through deep learning, *Procedia Computer Science*, Volume 216, 2023, Pages 57-66, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2022.12.111>.
- [2] Yuchao Lyu, Yinghao Xu, Xi Jiang, Jianing Liu, Xiaoyan Zhao, Xijun Zhu, AMS-PAN: Breast ultrasound image segmentation model combining attention mechanism and multi-scale features, *Biomedical Signal Processing and Control*, Volume 81, 2023, 104425, ISSN 1746-8094, <https://doi.org/10.1016/j.bspc.2022.104425>.
- [3] Albert Clèrigues, Sergi Valverde, Joaquim Salvi, Arnau Oliver, Xavier Lladó, Minimizing the effect of white matter lesions on deep learning based tissue segmentation for brain volumetry, *Computerized Medical Imaging and Graphics*, Volume 103, 2023, 102157, ISSN 0895-6111, <https://doi.org/10.1016/j.compmedimag.2022.102157>.
- [4] Selen Pehlivan, Jorma Laaksonen, Improved action proposals using fine-grained proposal features with recurrent attention models, *Journal of Visual Communication and Image Representation*, Volume 90, 2023, 103709, ISSN 1047-3203, <https://doi.org/10.1016/j.jvcir.2022.103709>.
- [5] Shenglong Chen, Yoshiki Ogawa, Chenbo Zhao, Yoshihide Sekimoto, Large-scale individual building extraction from open-source satellite imagery via super-resolution-based instance segmentation approach, *ISPRS Journal of Photogrammetry and Remote Sensing*, Volume 195, 2023, Pages 129-152, ISSN 0924-2716, <https://doi.org/10.1016/j.isprsjprs.2022.11.006>.
- [6] Sandra Jardim, João António, Carlos Mora, Image thresholding approaches for medical image segmentation - short literature review, *Procedia Computer Science*, Volume 219, 2023, Pages 1485-1492, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2023.01.439>.
- [7] Krishna Chaitanya, Ertunc Erdil, Neerav Karani, Ender Konukoglu, Local contrastive loss with pseudo-label based self-training for semi-supervised medical image segmentation, *Medical Image Analysis*, 2023, 102792, ISSN 1361-8415, <https://doi.org/10.1016/j.media.2023.102792>.
- [8] Robert Mendel, David Rauber, Luis A. de Souza, João P. Papa, Christoph Palm, Error-Correcting Mean-Teacher: Corrections instead of consistency-targets applied to semi-supervised medical image segmentation, *Computers in Biology and Medicine*, Volume 154, 2023, 106585, ISSN 0010-4825, <https://doi.org/10.1016/j.compbimed.2023.106585>.
- [9] Guangju Li, Dehu Jin, Qi Yu, Meng Qi, IB-TransUNet: Combining Information Bottleneck and Transformer for Medical Image Segmentation, *Journal of King Saud University - Computer and Information Sciences*, Volume 35, Issue 3, 2023, Pages 249-258, ISSN 1319-1578, <https://doi.org/10.1016/j.jksuci.2023.02.012>.
- [10] Marco Trombini, David Solarna, Gabriele Moser, Silvana Dellepiane, A goal-driven unsupervised image segmentation method combining graph-based processing and Markov random fields, *Pattern Recognition*, Volume 134, 2023, 109082, ISSN 0031-3203, <https://doi.org/10.1016/j.patcog.2022.109082>.
- [11] Pedro Bocca, Adrian Orellana, Carlos Soria, Ricardo Carelli, On field disease detection in olive tree with vision systems, *Array*, Volume 18, 2023, 100286, ISSN 2590-0056, <https://doi.org/10.1016/j.array.2023.100286>.
- [12] Ansam A. Abdulhussien, Mohammad F. Nasrudin, Saad M. Darwish, Zaid Abdi Alkareem Alyasseri, Feature selection method based on quantum inspired genetic algorithm for Arabic signature verification, *Journal of King Saud University - Computer and Information Sciences*, Volume 35, Issue 3, 2023, Pages 141-156, ISSN 1319-1578, <https://doi.org/10.1016/j.jksuci.2023.02.005>.
- [13] Khabir Uddin Ahamed, Manowarul Islam, Ashraf Uddin, Arnisha Akhter, Bikash Kumar Paul, Mohammad Abu Yousuf, Shahadat Uddin, Julian M.W. Quinn, Mohammad Ali Moni, A deep learning approach using effective preprocessing techniques to detect COVID-19 from chest CT-scan and X-ray images, *Computers in Biology and Medicine*, Volume 139, 2021, 105014, ISSN 0010-4825, <https://doi.org/10.1016/j.compbimed.2021.105014>.
- [14] Sushreeta Tripathy, Tripti Swarnkar, Unified Preprocessing and Enhancement Technique for Mammogram Images, *Procedia Computer Science*, Volume 167, 2020, Pages 285-292, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2020.03.223>.

- [15] Larissa Ferreira Rodrigues, Murilo Coelho Naldi, João Fernando Mari, Comparing convolutional neural networks and preprocessing techniques for HEp-2 cell classification in immunofluorescence images, *Computers in Biology and Medicine*, Volume 116, 2020, 103542, ISSN 0010-4825, <https://doi.org/10.1016/j.compbimed.2019.103542>.
- [16] Yiwei Fan, Chunyu Guo, Yang Han, Weizheng Qiao, Peng Xu, Yunfei Kuai, Deep-learning-based image preprocessing for particle image velocimetry, *Applied Ocean Research*, Volume 130, 2023, 103406, ISSN 0141-1187, <https://doi.org/10.1016/j.apor.2022.103406>.
- [17] Kouta Hirotaki, Shunsuke Moriya, Tsunemichi Akita, Kazutoshi Yokoyama, Takeji Sakae, Image preprocessing to improve the accuracy and robustness of mutual-information-based automatic image registration in proton therapy, *Physica Medica*, Volume 101, 2022, Pages 95-103, ISSN 1120-1797, <https://doi.org/10.1016/j.ejmp.2022.08.005>.
- [18] Burak Tasci, Irem Tasci, Deep feature extraction based brain image classification model using preprocessed images: PDRNet, *Biomedical Signal Processing and Control*, Volume 78, 2022, 103948, ISSN 1746-8094, <https://doi.org/10.1016/j.bspc.2022.103948>.
- [19] Reewos Talla-Chumpitaz, Manuel Castillo-Cara, Luis Orozco-Barbosa, Raúl García-Castro, A novel deep learning approach using blurring image techniques for Bluetooth-based indoor localisation, *Information Fusion*, Volume 91, 2023, Pages 173-186, ISSN 1566-2535, <https://doi.org/10.1016/j.inffus.2022.10.011>.
- [20] Aashu Kumar, Peeta Basa Pati, Offline HWR Accuracy Enhancement with Image Enhancement and Deep Learning Techniques, *Procedia Computer Science*, Volume 218, 2023, Pages 35-44, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2022.12.399>.
- [21] Samuel Fernández-Mendiña, Fernando Pérez-González, Miguel Masciopinto, Source camera attribution via PRNU emphasis: Towards a generalized multiplicative model, *Signal Processing: Image Communication*, Volume 114, 2023, 116944, ISSN 0923-5965, <https://doi.org/10.1016/j.image.2023.116944>.
- [22] Mahesh S Patil, Satyadhyam Chickerur, C Abhimalya, Anishma Naik, Nidhi Kumari, Shashank Maurya, Effective Deep Learning Data Augmentation Techniques for Diabetic Retinopathy Classification, *Procedia Computer Science*, Volume 218, 2023, Pages 1156-1165, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2023.01.094>.
- [23] Ilknur Tuncer, Prabal Datta Barua, Sengul Dogan, Mehmet Baygin, Turker Tuncer, Ru-San Tan, Chai Hong Yeong, U. Rajendra Acharya, Swin-textural: A novel textural features-based image classification model for COVID-19 detection on chest computed tomography, *Informatics in Medicine Unlocked*, Volume 36, 2023, 101158, ISSN 2352-9148, <https://doi.org/10.1016/j.imu.2022.101158>.
- [24] Tomé Albuquerque, Luís Rosado, Ricardo Cruz, Maria João M. Vasconcelos, Tiago Oliveira, Jaime S. Cardoso, Rethinking low-cost microscopy workflow: Image enhancement using deep based Extended Depth of Field methods, *Intelligent Systems with Applications*, Volume 17, 2023, 200170, ISSN 2667-3053, <https://doi.org/10.1016/j.iswa.2022.200170>.
- [25] Tiwalade Modupe Usman, Yakub Kayode Saheed, Djitog Ignace, Augustine Nsang, Diabetic retinopathy detection using principal component analysis multi-label feature extraction and classification, *International Journal of Cognitive Computing in Engineering*, Volume 4, 2023, Pages 78-88, ISSN 2666-3074, <https://doi.org/10.1016/j.ijcce.2023.02.002>.
- [26] Laurens Versluis, Mehmet Cetin, Caspar Greeven, Kristian Laursen, Damian Podareanu, Valeriu Codreanu, Alexandru Uta, Alexandru Iosup, Less is not more: We need rich datasets to explore, *Future Generation Computer Systems*, Volume 142, 2023, Pages 117-130, ISSN 0167-739X, <https://doi.org/10.1016/j.future.2022.12.022>.
- [27] Shahid Ahmad Wani, S.M.K. Quadri, Evaluation of Computational Methods for Single Cell Multi-Omics Integration, *Procedia Computer Science*, Volume 218, 2023, Pages 2744-2754, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2023.01.246>.
- [28] Jiaxin Yu, Florian Wellmann, Simon Virgo, Marven von Domarus, Mingze Jiang, Joyce Schmatz, Bastian Leibe, Superpixel segmentations for thin sections: Evaluation of methods to enable the generation of machine learning training data sets, *Computers & Geosciences*, Volume 170, 2023, 105232, ISSN 0098-3004, <https://doi.org/10.1016/j.cageo.2022.105232>.