

Title	Deep learning pipeline for blood cell segmentation, classification and counting
Authors	Rao, Abhijeet;Kho, Kiang Wei;Mykytiv, Vitaliy;Visentin, Andrea
Publication date	2024
Original Citation	Rao, A., Kho, K. W., Mykytiv, V., and Visentin, A.(2024) 'Deep learning pipeline for blood cell segmentation, classification and counting', 32nd Irish Conference on Artificial Intelligence and Cognitive Science, Dublin, Ireland, December 9-10, 2024.
Type of publication	Conference item
Link to publisher's version	https://ceur-ws.org/Vol-3910/
Rights	© 2022, the Authors. Use permitted under Creative Commons License Attribution 4.0 International [CC BY 4.0] - https://creativecommons.org/licenses/by/4.0/
Download date	2026-08-31 00:33:27
Item downloaded from	https://hdl.handle.net/10468/16874



UCC

University College Cork, Ireland
Coláiste na hOllscoile Corcaigh

Deep Learning Pipeline for Blood Cell Segmentation, Classification and Counting

Abhijeet Rao¹, Kiang Wei Kho², Vitaliy Mykytiv³ and Andrea Visentin^{1,4,*}

¹*School of Computer Science & IT, University College Cork, Ireland*

²*Biophotonics@Tyndall, IPIC, Tyndall National Institute, Ireland*

³*Hematology Unit, Cork University Hospital, Ireland*

⁴*SFI Insight Centre for Data Analytics, University College Cork, Ireland*

Abstract

Blood cell analysis, an important component of modern medicine, provides critical insights into a range of health conditions, right from minor infections to blood disorders such as leukemia. Traditional methods such as flow cytometry have been very effective in blood cell counting but are often reliant on expensive instruments, specialised reagents, and expert operators, thereby limiting their accessibility. To combat these challenges, this research explores the use of advanced deep learning methods to develop a robust and scalable pipeline for blood cell segmentation, classification, and counting. The proposed pipeline automates the entire process of generating high-precision segmentation masks and classifying blood cells into red blood cells or white blood cells by leveraging Meta AI's Segment Anything Model and a Convolutional Neural Network. Apart from employing a host of image processing operations, the study also addresses challenges such as low image quality and data scarcity, which leads to class imbalance, by employing techniques such as image contrast enhancement and data augmentation. The pipeline can provide accurate cell counting from simple blood smears. The pipeline was then adapted to blood smear images collected by our research team through different instrumentation using transfer learning applied to the deep learning classifier. With over 95% accuracy in blood cell counts across different datasets, the pipeline exhibits a high degree of adaptability, which ultimately validates its potential for broader clinical applications.

Keywords

Deep Learning, Segment Anything Model, Blood Smears, Convolutional Neural Network, Transfer learning

1. Introduction

The rise of new and re-emerging diseases has placed immense pressure on modern medicine to develop faster and more precise diagnostics, piling on top of the already existing issue of global health challenges such as accessibility of cheap, quick, and quality healthcare. In this context, blood analysis is an essential diagnostic and monitoring tool for a broad spectrum of illnesses. One of the most common tests is the complete blood test (CBC), which gives vital details on the types and quantities of blood cells, namely red blood cells (RBC), white blood cells (WBC), and platelets. Accuracy in blood cell counting is vital to detect diseases like leukemia, anaemia, and infections, where any deviations from the expected cell counts and morphologies could indicate underlying health issues.

Significant progress has been made in the traditional biotechnological methods of blood analysis, especially in the area of cell cytometry. Counting and analyzing cells was initially a manual process using instruments like hemocytometers, which was labour-intensive and prone to inaccuracy. These techniques have developed into automated systems that leverage optical analysis, electrical impedance, and image cytometry (Vembadi et al.[1]). By enhancing accuracy and throughput, these technologies have radically transformed the cell counting process, thereby playing a crucial role in both clinical and research settings. One popular method which permits the detailed multi-parametric analysis of thousands of cells per second is flow cytometry. It enables accurate cell identification and quantification by using lasers to excite fluorochromes attached to cells and detectors to measure their fluorescence and

32nd Irish Conference on Artificial Intelligence and Cognitive Science, December 09–10, 2024, Dublin, Ireland

✉ 123112200@umail.ucc.ie (A. Rao); kiangwei.kho@tyndall.ie (K. W. Kho); Vitaliy.Mykytiv@hse.ie (V. Mykytiv); andrea.visentin@ucc.ie (A. Visentin)

ORCID 0009-0000-7545-833X (A. Rao); 0000-0003-3702-4826 (A. Visentin)



© 2022 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

scatter properties. Despite its high accuracy and throughput, flow cytometry is often constrained by its high cost, the requirement for specialised reagents, and the need for skilled operators to operate the equipment and interpret the results. This restricts its scalability and accessibility, especially in settings with limited resources or smaller laboratories with limited funding. This situation emphasises the need for more accessible, user-friendly, and scalable blood analysis solutions.

In the meantime, the field of artificial intelligence (AI) grew by leaps and bounds, with remarkable progress in the subfields of machine learning and deep learning. This naturally prompted the adoption of AI in healthcare, more particularly haematological analysis. Early attempts at digital automation employed classical image processing techniques such as thresholding, edge detection, and morphological operations, which, although useful, often struggled with variations in cell shapes, sizes, and staining. Then, machine learning (ML) provided a more sophisticated means of analyzing blood cells by learning from labelled datasets [2, 3].

In recent years, deep learning (DL) has revolutionized the field of medical image analysis, particularly through the use of Convolutional Neural Networks (CNNs). These models have the ability to automatically learn and extract hierarchical features from images, leading to highly accurate classification and segmentation of blood cells [4]. Building on these advancements, this study explores the integration of the Segment Anything Model (SAM)[5], the latest open-source state-of-the-art deep learning based segmentation tool created by Meta AI, and the previously proven CNN classifier to develop a comprehensive pipeline for blood cell analysis. This approach not only automates the process of blood cell segmentation and classification but also addresses challenges such as class imbalance and the need for contrast enhancement in low-quality images. By leveraging transfer learning and data augmentation, the proposed method aims to provide a highly scalable and accurate module for blood cell analysis.

This study was motivated by the persistent challenges in medical diagnostics, particularly in the field of haematology. Accurate measurement and classification of blood cells are necessary for the diagnosis of many diseases, from minor infections to serious conditions like leukemia. Despite their effectiveness, traditional methods have certain disadvantages. They are expensive, necessitate specific equipment, and depend on skilled laboratory technicians. These challenges may result in delayed diagnosis and suboptimal patient outcomes, especially in areas with limited resources. Artificial intelligence and deep learning have the potential to promote access to high-quality blood analysis through the development of more scalable and economical systems. Therefore, the goal is to use these techniques to alleviate the challenges encountered in the healthcare industry.

2. Literature Review

Significant advancements in blood cell image analysis have been made possible by the incorporation of cutting-edge machine learning algorithms, which have fuelled the creation of exact and effective diagnostic instruments. Tai et al.[2] pioneered a hierarchical classification technique for blood cell identification and categorisation that makes use of multi-class SVM. Their method extracts geometric elements from digital images to successfully segment blood cells and aid in distinguishing between distinct cell types such as erythrocytes, leukocytes, and thrombocytes. This approach showed the potential to be used in clinical diagnostics because it not only automated the labour-intensive and historically labour-intensive blood cell counting process, but it also showed excellent accuracy when compared to expert evaluations. Huang et al.[3] concentrated mainly on the segmentation and identification of leukocyte nuclei within blood-smear pictures, building on the groundwork established by these methods. Their technique uses multilevel Otsu thresholding to segment the image after enhancing the leukocyte nucleus with a specialised enhancer. For classification, a genetic algorithm-based k-means clustering method is employed, with a dimensionality reduction step using Principal Component Analysis (PCA). This method achieves a high recognition rate and shows promise in automating the leukocyte differential counting procedure, which is critical for identifying many diseases. It also proves to be particularly resistant to fluctuations in staining conditions and image quality. Ruberto et al.[6] overcame the difficulties caused by different lighting and colour circumstances by creating an automated system

specifically designed to segment and count WBCs in blood smear images. The method divides the blood cells into groups according to pixel-wise characteristics in a non-linear feature space by using a SVM. After that, the WBCs are precisely counted using the Circular Hough Transform even when cells are clumped together. This system was tested on several public datasets and demonstrated its robustness and effectiveness in various imaging situations with a remarkable accuracy of 99.2%. Finally, Tavakoli et al.[7] made further advancements in this area by creating a unique strategy combining feature extraction and segmentation exclusively for WBCs in peripheral smear pictures. Their technique presents a novel algorithm for precisely identifying the cytoplasm and segmenting the nucleus, after which form and colour information are extracted. This method, which uses a SVM model for classification, performs better than conventional CNN-based techniques, exhibiting high accuracy and good generalisation over a variety of datasets. This development emphasises how useful it could be as a clinical diagnostic tool.

The merging of deep learning and sophisticated machine learning algorithms has led to notable progress in blood cell image analysis in recent years. Mohamed et al.[8] employed a two-stage hybrid approach that blends conventional machine learning classifiers with pre-trained deep learning models to improve WBC categorisation. With a classification accuracy of 97.03%, this method showed exceptional accuracy, especially when combining MobileNet-224 with logistic regression. The study demonstrates how well deep learning characteristics may be combined with traditional machine learning models to improve WBC classification reliability. Kaza et al.[9] proposed the use of conditional Generative Adversarial Networks (cGAN) for virtual staining of single channel deep-ultraviolet (UV) images to mimic Giemsa staining. For the segmentation of WBCs, two independent CNNs are used to segment cellular and nuclear regions, and for classification, a pre-trained ResNet-18 model is used to classify WBCs into five subtypes. Using the recently created RV-PBS dataset, the research by Pal et al.[10] presents a novel method for automating the segmentation and categorisation of WBCs. The research tackles the issues of class imbalance and small sample numbers by utilising Mask R-CNN, for instance, segmentation and domain adaption strategies. Significant gains in precision, recall, and F1-score were obtained when the RV-PBS and PBC datasets were combined, highlighting the method's potential to increase diagnosis accuracy in clinical settings.

The difficulty of domain variability in blood cell categorisation was further addressed by Li et al.[11] by introducing a domain-invariant representation learning (DoRL) framework. Their DoRL approach combines a cross-domain autoencoder (CAE) for domain-invariant feature extraction with a LoRA-based Segment Anything Model (LoRA-SAM) for segmentation. This approach greatly outperformed previous models and offered a reliable means of classifying blood cells in a variety of clinical settings, where variations in imaging circumstances and laboratory protocols have historically weakened model performance. Abdulkadir et al.[12] presented a hybrid approach for the classification of WBCs in blood smear images that further combines machine learning and image processing. Their method uses a combination of deep learning models and conventional image processing approaches, making use of both shape-based data and deep features taken from a CNN model that has already been trained. After that, a Long Short-Term Memory (LSTM) network is used for classification. With an accuracy of 85.7%, this hybrid approach proved how well-suited it is to combine conventional approaches with state-of-the-art deep learning techniques for WBC categorisation. When taken as a whole, these developments mark a substantial improvement in blood cell analysis automation and accuracy, opening the door to more dependable and effective clinical diagnostics. Lastly, Hemalatha et al.[13] introduced an Enhanced Convolutional Neural Network (ECNN) that targets both WBCs and RBCs in an effort to increase the accuracy of blood cell classification. Their ECNN model, which uses sophisticated feature extraction and pre-processing methods, including K-means clustering for segmentation, produced a classification accuracy of 95%. This method shows how deep learning may greatly improve blood cell analysis's precision and effectiveness, making it a useful tool for medical applications requiring early detection.

3. Methodology

This section initially provides a brief explanation of the proposed pipeline along with an overview of the datasets used. Then, each stage of the pipeline is broken down into separate sections, starting with preprocessing, proceeding with segmentation and classification, and terminating with counting and metrics used to evaluate the pipeline's efficacy.

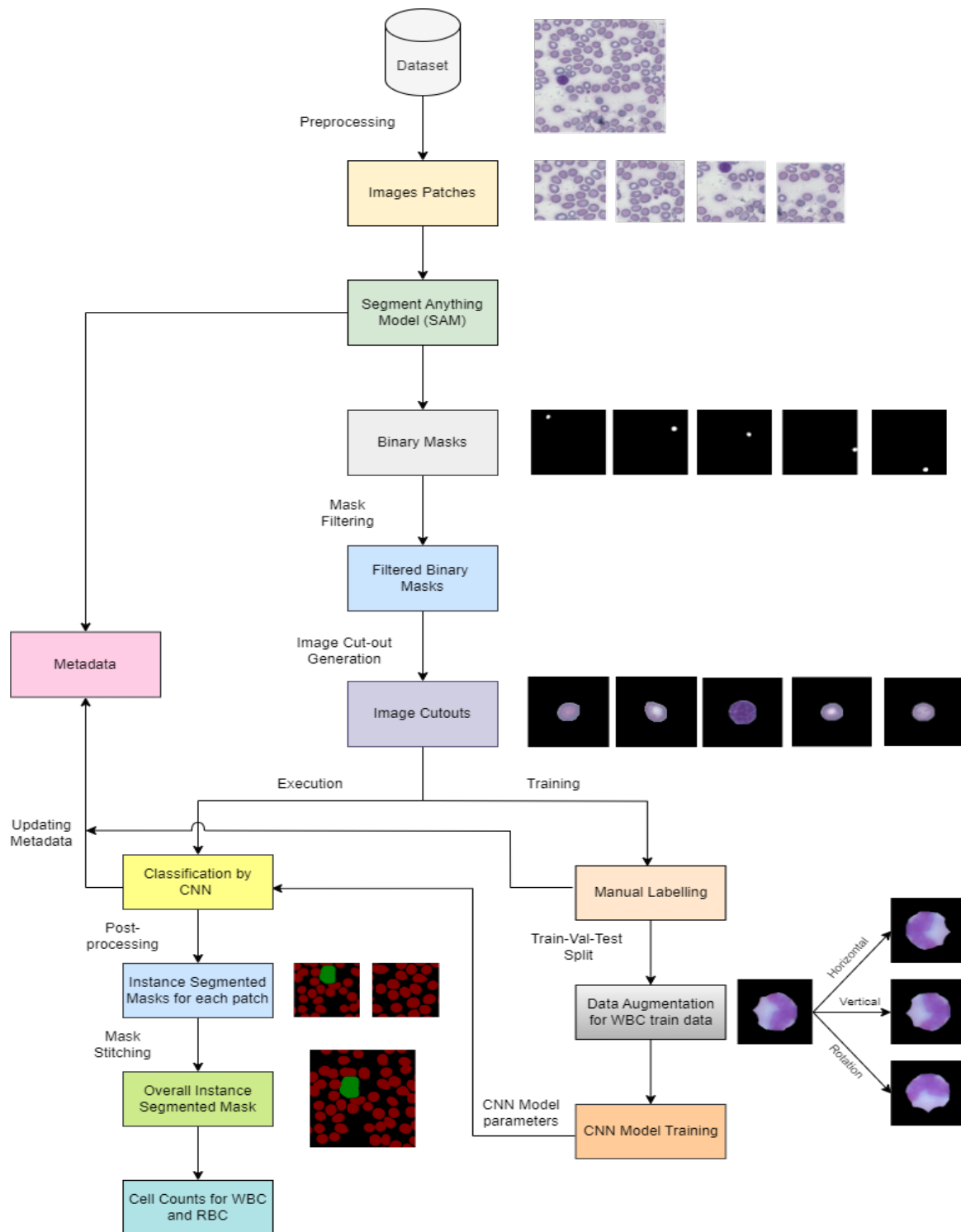


Figure 1: Flowchart of the pipeline process

The pipeline proposed for automated analysis of blood smear analysis integrates both a training and execution branch as indicated in Figure 1. This pipeline uses a common pre-processing and segmentation stage to ensure consistency across the entire workflow. The goal of this pipeline is to generate an instance segmentation mask for the blood smears and accurately count RBCs and WBCs. It begins with the preprocessing step of dividing the input image into smaller, non-overlapping patches using a grid-based approach. This step is crucial as it aids in detailed analysis and segmentation of individual patches rather than the entire image at once. SAM operates on these individual patches and generates binary masks to identify cellular regions. Now, Mask filtering is performed by removing repetitive and overlapping masks. These refined masks serve as a foundation for training the classification model and executing the instance segmentation and cell counting process.

Training the classification model is necessary prior to executing the overall pipeline. So, in the training branch, the binary masks generated by SAM are used to create image cutouts, which are essentially smaller images containing isolated cellular regions. To build the training dataset, the cutouts are labelled as either RBCs or WBCs. Data augmentation is applied specifically to the WBC cutouts to enhance the diversity of the training set. The augmented and original data are then used to train a CNN classifier to distinguish between RBCs and WBCs.

In the execution branch, image cutouts are generated from these refined masks following the same pre-processing step, mask generation by SAM, and mask filtering. The cutouts are then passed through the trained CNN classifier, which labels each cutout as either RBC or WBC. Finally, in the post-processing stage, these labelled cutouts are compiled to produce a final segmentation mask, and the RBC count is adjusted by analyzing the distribution of cell areas. The final output, including the segmented image, RBC and WBC counts, are made available to the user. This chapter deals with each stage of the pipeline individually. For the sake of reproducibility, the code implementing the full pipeline has been made available ¹.

3.1. Datasets

Two datasets were utilized in the overall implementation of this project. The first dataset is sourced from the study performed by Ojaghi et al.[14], where the researchers focus on developing a method for analysing blood cells without the need for staining or labelling, which is common in traditional haematology techniques. This dataset is publicly available². This method employs ultraviolet (UV) microscopy to perform haematological analysis, offering a label-free approach. However, to ensure the validity and comparability of their results, the researchers also included brightfield images in their dataset. These traditional microscopy-created images provide a benchmark for comparison.

The blood smears used to create the brightfield images had their cells and constituent parts dyed with Giemsa dye, which allowed for their visualisation under a brightfield microscope. As is common in conventional blood smear examinations, the samples were fixed with methanol before staining. These rich, colourized images are crucial for verifying the unique UV microscopy approach since they make cellular features easily visible. An example of a brightfield blood smear from this dataset is shown in Figure 2a.

The second dataset is used for testing and adapting the pipeline, and includes 118 images, each of dimensions 3264x1840 was collected internally. Anticoagulant added Whole Blood (WB) sample was collected from a patient. Blood smears were prepared by dropping approximately 10-20ul of the WB on an uncoated quartz slide, then evenly spread out with a second spreading slide to form a monolayer of non-overlapping blood cells. The sample was then stained with polychrome stains (Wright-Giemsa) using Sysmex's automated slide preparation unit SP-50 so as to accentuate the nuclei structure for cell differentiation. Images were then acquired with a MICON BM2 microscope using a x20 objective lens under white-light illumination. Images from this dataset will be referred as the new dataset of blood smears for the rest of this study. The pipeline classifier is adapted to this new dataset which allows the evaluation of the generalisability and effectiveness of the pipeline to this data. A typical blood smear

¹<https://github.com/Abhi21298/Thesis-Blood-Smears>

²<https://osf.io/ayw4j/>

image from this new dataset is shown in figure 2b. This study has been reviewed and approved by CREC (Clinical Research Ethics Committee) under protocol number ECM 4 (m) 10/09/2024.

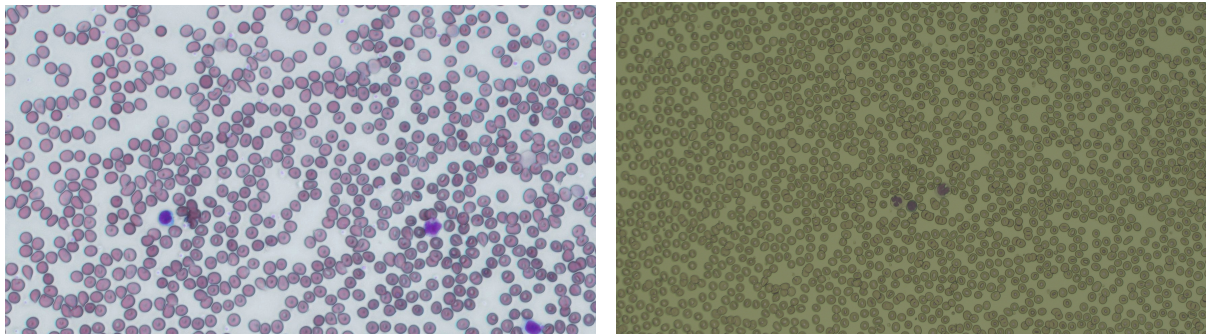


Figure 2: a) Giemsa stained blood smear from Brightfield microscope [14] b) New dataset blood smear image

3.2. Preprocessing

This is the first stage of the pipeline wherein an input image of $W \times H$ dimensions is divided into smaller patches. Here, W and H denote the width and height of the image, respectively. This step aims to handle large images effectively and ensure that each patch is small enough to be handled efficiently. The input image is split into non-overlapping patches to avoid processing repetitive information and allow focused analysis. Each patch is typically of the size 512×512 pixels. However, for images whose width and height are not perfectly divisible, the remaining portions of the patch are extracted to facilitate further analysis. Figure 1 provides an example of an image and its corresponding patches.

3.3. Segmentation

SAM [5] is utilized to generate semantic segmentation binary masks for each blood smear image patch acquired previously. It analyzes each image patch to identify and segment individual cells, particularly RBCs and WBCs and creates distinct binary masks automatically for each detected cell, allowing for the separation of individual instances, even when cells with distinct boundaries are clustered together. For each patch, SAM produces several binary masks that outline the boundaries of each detected cell. Example binary masks are shown in Figure 1. Alongside these masks, a metadata file is generated, containing critical information such as its ID, bounding box coordinates, cell area, and Intersection over Union (IoU) values for each mask. This metadata is essential for further filtering and analysis. Once the segmentation is complete, redundant and overlapping masks are removed from the binary masks created for the image patches using filtering. This ensures that each segmented cell is represented only once, improving the accuracy of further analysis.

3.4. Classification

For the classification task, a CNN model is used and manual labelling is required to assign each image cutout as either RBC or WBC. Since cutouts are directly generated from corresponding image patches, their labels are stored in the appropriate metadata. Once labelled, the dataset is split into training, validation, and test sets in a 75:15:10 ratio, and it is applied separately to both RBC and WBC image cutouts to ensure the CNN can generalise well. Owing to SAM's zero-shot nature, segmentation does not require data splitting and training, unlike CNN, which trains on labelled data. The architecture consists of 3 convolutional layers, along with a ReLU activation function for each convolution layer. The convolution layer extracts hierarchical features by producing feature maps, and ReLU adds non-linearity for learning complex patterns. Each convolution layer is followed by a Max pooling layer with a window size of 2×2 . Max pooling reduces the spatial dimensions of the feature maps while retaining important features [4]. Following the convolutional and pooling layers, the feature maps are flattened and passed

through a fully connected layer. This dense layer has a ReLU activation function as well. This layer helps combine all the extracted features and aids in the final decision-making process. A dropout layer is applied to prevent overfitting, randomly setting a fraction of the neurons in the dense layer to zero during training. This encourages the network to not rely heavily on specific features. The final layer consists of a single output neuron without an activation function, allowing a linear output before applying the sigmoid activation function for binary classification [4]. The CNN architecture is displayed in Figure 3.

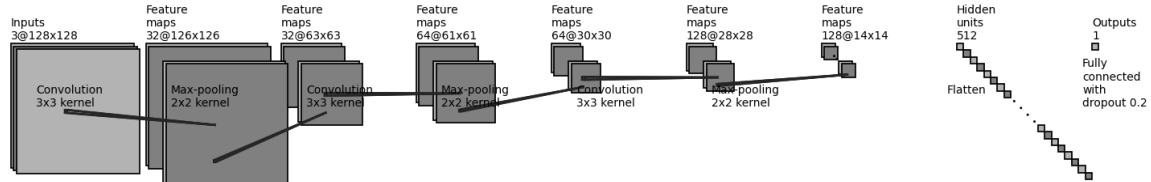


Figure 3: CNN classifier architecture

Hyperparameters which are external mode configurations need to be tuned properly before training the model. They are fixed and govern the behaviour of the learning process. Important hyperparameters here include the number of epochs, learning rate, batch size, and dropout rate. These parameters were tuned to optimal values to extract best classification results.

Data Augmentation

Data augmentation is a machine learning approach to perform several modifications on existing data and artificially increase the size and variety of data ([15]). This is applied exclusively to WBC cutouts due to the significant imbalance between RBCs and WBCs in the dataset, as RBCs outnumbered WBCs by approximately 600:1. Techniques such as horizontal flips, vertical flips and rotations (up to 90 degrees) were applied to inflate the WBC training data by tenfold. Then, some RBC cutouts were randomly excluded from the training data to maintain an RBC to WBC ratio of 2:1 in the training dataset. These steps were necessary to eliminate the bias towards predicting the dominant class by the CNN model. Data augmentation is performed only on the training set to prevent data leakage, leaving the validation and testing dataset untouched. This ensures proper evaluation of the model to unseen data. An example of data augmentation is illustrated in Figure 1.

3.5. Post-processing

This process is part of the pipeline execution and is essential for generating instance segmentation masks that accurately represent the cells identified in the image. After classification, the results are updated in the metadata by associating each cutout's classification label with its corresponding binary mask and patch from which it was derived. After the binary masks for each image patch have been labelled as either RBC or WBC, the next step involves creating a visual representation that combines these masks into an interpretable format. This is done by creating instance segmentation masks for each patch first and then "stitching" these masks together to form an overall instance segmentation mask for the input image. During this process, each mask is assigned a colour based on its label. RBCs are represented in red, while WBCs are represented in green. The resultant mask is a colour-coded instance mask for each patch, which distinguishes between different types of cells, making it easier to interpret the segmented results. The individual patch masks generated previously are stitched together to create a full instance segmentation mask for the entire input image. This reverses the preprocessing step of dividing the image into patches by combining the patches to recreate the original image with all segmented cells represented in appropriate colours.

3.6. Counting

When an image is divided into patches, it is highly likely that some cells are split across multiple patches and have separate binary masks respectively. So, basing the count on image cutouts alone would lead to inflated cell counts, as some cells would be counted multiple times. Thus, cell counting is done after image reconstruction to avoid such errors as the final instance mask resolves this issue by merging split cells back into a single entity for accurate counting. For WBCs, the counting process is straightforward as they are distinct and well-separated in the mask. Each green-coloured region is counted as one WBC. For RBCs however, are high in number in comparison to WBC and have a high tendency of overlapping. This leads to larger regions being counted as a single entity because overlapped cells are captured as a single mask. To correct this, RBC areas are analysed. Larger areas likely represent multiple cells overlapping and are therefore counts are inflated appropriately providing a more accurate RBC count.

3.7. Metrics

The metrics calculated for evaluating the CNN classifier's performance include Binary Cross Entropy (BCE) and Classification Accuracy. The BCE function is used as a loss function for binary classification problems where the target outcome is a probability between 0 and 1. The model aims to reduce this value during training to improve prediction accuracy. Classification accuracy measures the proportion of correct predictions made by the model out of the total number of predictions. This accuracy is calculated for training, validation and test datasets. The count accuracy of each blood cell type is attributed to overall pipeline performance. It is computed as the ratio between the predicted number of blood cells and the actual count.

4. Experimental Results

4.1. Classification Accuracy

The process described from Section 3.2 to Section 3.4 is followed to have the image cutout data ready for training the CNN model. Before training, pixel values for the image cutouts belonging to the train dataset are scaled to values between 0 and 1. Hyperparameter tuning is performed to optimise the model's performance. To find the configuration that yields the best results, various values for the batch size, learning rate, the number of epochs and dropout rate are tested using the grid search technique. The search space for each hyperparameter for implementing the grid search technique and the optimal hyperparameters determined are tabulated in Table 1 and Table 2 respectively.

Hyperparameter	Values
Batch size	[2, 4, 8]
Epochs	[10, 20, 30]
Dropout Rate	[0.1, 0.2, 0.5]
Learning Rate	[$1e^{-5}$, $1e^{-4}$, $1e^{-3}$, 0.01]

Table 1
Hyperparameter search space

Hyperparameter	Assigned Value
Batch Size	4
Epochs	20
Dropout Rate	0.2
Learning Rate	$1e^{-4}$

Table 2
Optimal Hyperparameters

Utilising the optimal hyperparameters, the CNN model yielded a validation accuracy of 98.5% and a training accuracy of 100%. To ensure the validity of the model's training, the trained model was tested on the test image cutout dataset. Here, the model achieved an accuracy of 97.14%. This validates the fact that the model learned the required features from the images to distinguish RBC and WBC effectively. Table 3 summarizes the number of image cutouts present in training, validation and test datasets for each type of blood cell, along with the corresponding predictions by the trained model. The table shows that only one WBC image cutout was misclassified as RBC in the unseen or the test dataset.

Dataset Type	Cutout Data		Model Predictions	
	RBC	WBC	RBC	WBC
Train (75%)	500	225 (After data augmentation)	500	225
Validation (15%)	101	10	102	9
Test (10%)	66	4	67	3

Table 3

Data split and model performance for CNN

4.2. Pipeline Execution

The pipeline is executed on the blood smear image shown in Figure 2a. Preprocessing is performed to generate patches and each patch is fed into SAM to generate binary masks along with its metadata. Mask filtering is performed to remove repeated and overlapping masks. Then, image cutouts are generated by using a bitwise operation between the binary mask and the corresponding image patch it is associated with. The mask has varying dimensions dictated by the bounding box coordinates presented in the metadata before being zero-padded to achieve a uniform size of 128x128 pixels. This process is demonstrated visually in Figure 4.

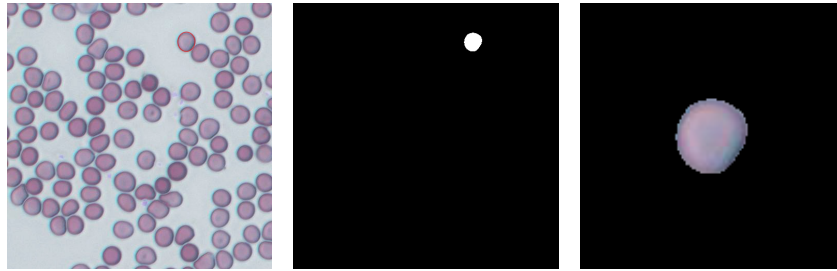


Figure 4: a) patch image b) cell binary mask c) final cutout

Now, the trained CNN model classifies all the image cutouts. Since these cutouts are generated from specific image patches, the classification results are mapped to their corresponding metadata file generated previously by SAM. Later, post-processing is performed to create instance segmentation masks. This is generated by assigning every binary mask a colour based on the label assigned by the CNN (Red - RBC, Green - WBC) and combining all masks associated with a single patch. This process is repeated across all patches. Ultimately, these instance segmentation masks are stitched together to create an overall segmentation mask shown in figure 5.

Finally, cell counting is performed on the basis of this overall instance segmentation mask. While WBC counting is straightforward, with each green-coloured region representing a single WBC, for RBCs, there is a distinct possibility of multiple cells overlapping and posing as a single entity. To address this, a histogram (Figure 6) of the segmented RBC segmented regions is generated with a gaussian kernel density estimation curve (KDE) to smoothen the distribution. The KDE curve helps highlight the key peaks that represent typical cell sizes.

The first peak in the KDE curve corresponds to the average area of a single RBC. Applying thresholds in between the KDE curve's successive peaks allows for the handling of larger regions in the segmentation masks, which represent overlapping cells. These thresholds take into consideration the potential for overlaps to cause many cells to appear as one. Since it is uncommon for more than four RBCs to overlap, regions with areas over the first threshold are counted as two RBCs, those exceeding the second threshold as three RBCs, and so on, up to four cells. This approach ensures the accuracy of both RBC and WBC counts, providing reliable results even in cases of overlapping cells. In the test set, the pipeline correctly identified 100% of the WBC and counted 2384 of the 2441 RBC (97.66% accuracy).

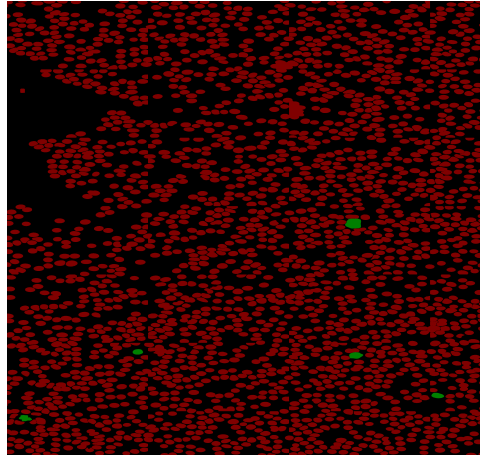


Figure 5: Overall Instance segmentation mask for Figure 2a

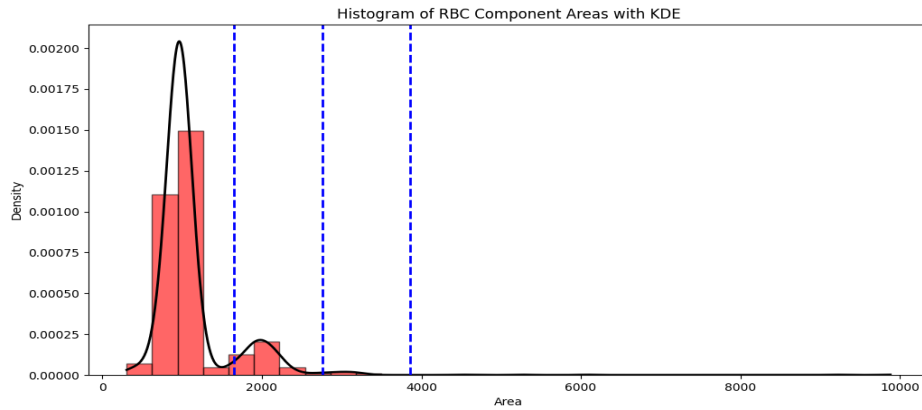


Figure 6: Area histogram of RBC cells in Figure 2a

4.3. Deployment to a new dataset

Adapting the pipeline to a new dataset involves two key modifications: contrast enhancement and transfer learning.

Contrast Enhancement - SAM's ability to effectively discern cell borders is restricted by the new dataset's lower brightness images. This is addressed by applying contrast enhancement, which brightens the photos below a preset threshold of 140. After experimenting with several values, it was determined that 140 would provide the best brightness for identifying the cells. Hence, this threshold was manually chosen. If an image's brightness is below this threshold, its pixel intensity is adjusted to meet the threshold. This ensures that SAM generates accurate segmentation masks for RBCs and WBCs, improving segmentation quality even in low-light conditions. Figure 7 shows the contrast-enhanced version.

Transfer Learning - Once image cutouts are available from images belonging to the new dataset, they are used to retrain the CNN classifier using transfer learning. Here, the CNN trained on the original dataset is fine-tuned using the new dataset image cutouts. The model retains the pre-trained weights and uses the same optimal hyperparameter values to train the model instead of training from scratch. This procedure yielded a BCE loss of 0 for both training and validation data and correspondingly yielded a 100% classification accuracy in train and validation datasets. The performance for training, validation and test datasets is tabulated in Table 4. There was only one misclassification where one WBC was identified as RBC in the test dataset. So, the test accuracy here is 96.5%.

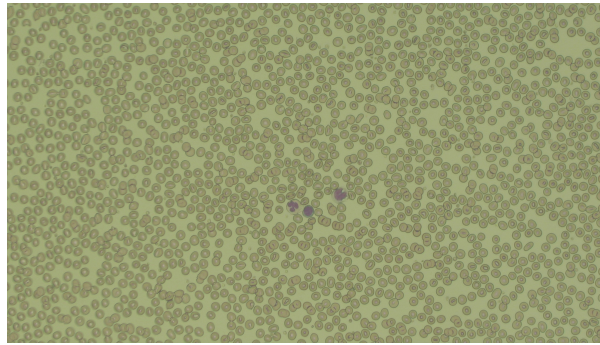


Figure 7: Contrast enhanced version of Figure 2b

Dataset type	Cutout Data		Model Predictions	
	RBC	WBC	RBC	WBC
Train (75%)	400	206 (After data augmentation)	400	206
Validation (15%)	80	7	80	7
Test (10%)	53	5	54	4

Table 4

Data split and Model performance for CNN after Transfer Learning

In the updated pipeline, this CNN model replaces the original CNN classifier to implement the pipeline on blood smear images belonging to the new dataset. The final instance segmentation mask is generated when the image is fed into the pipeline, and the cell count accuracy is determined. The pipeline correctly identified all WBC and counted 1818 red blood cells, while a manual counting returned a 1864 value, an accuracy of 97.54%.

5. Conclusions

In this research, we set out to address the challenges posed by traditional blood cell analysis methods by developing a comprehensive and automated deep-learning pipeline for blood cell segmentation, classification, and counting. This work adds to the evolution of medical diagnostics by addressing the shortcomings of conventional approaches and expanding on information obtained from earlier studies. It lays the groundwork for future developments by offering a viable, accessible, and reliable alternative to existing techniques. A thorough review of previous research in this discipline revealed the advantages and disadvantages of a number of techniques. Building on this foundation, we incorporated state-of-the-art methodologies to overcome these limitations. SAM completely reshaped the segmentation process by eliminating the necessity for manual mask creation, which is a common bottleneck in conventional techniques. The pipeline implemented was both scalable and flexible because of SAM's zero-shot generalisation capability, which generated extremely accurate segmentation masks without the need for additional training. The CNN model provides an accurate classification of RBCs and WBCs. Addressing challenges such as class imbalance through data augmentation and improving image quality via contrast enhancement were critical steps in ensuring the CNN model could perform reliably across varied datasets. Additionally, the study demonstrated the pipeline's adaptability through transfer learning, which enabled the model to maintain high accuracy when applied to new blood smear datasets with extremely limited amount of images.

This study's outcomes of cell count accuracy reaching over 95% across different image datasets. Given the variety of datasets comprising images made available from different preparation techniques, this success is worthy of note. The pipeline's effective implementation showcases its adaptability and scalability, making it a promising tool for improving diagnostic accuracy and efficiency in haematology.

Acknowledgments

This work was conducted with the financial support of Science Foundation Ireland under Grant Nos. 12/RC/2289-P2 and 16/RC/3918 which are co-funded under the European Regional Development Fund. This research was partially supported by the EU's Horizon Digital, Industry, and Space program under grant agreement ID 101092989-DATAMITE. For the purpose of Open Access, the author has applied a CC BY public copyright license to any Author Accepted Manuscript version arising from this submission.

References

- [1] A. Vembadi, A. Menachery, M. A. Qasaimeh, Cell cytometry: Review and perspective on biotechnological advances, *Frontiers in Bioengineering and Biotechnology* 7 (2019) 147.
- [2] W.-L. Tai, R.-M. Hu, H. C. Hsiao, R.-M. Chen, J. J. Tsai, Blood cell image classification based on hierarchical svm, in: 2011 IEEE International Symposium on Multimedia, 2011, pp. 129–136.
- [3] D.-C. Huang, K.-D. Hung, Leukocyte nucleus segmentation and recognition in color blood-smear images, in: 2012 IEEE International Instrumentation and Measurement Technology Conference Proceedings, 2012, pp. 171–176.
- [4] R. Yamashita, M. Nishio, R. K. G. Do, K. Togashi, Convolutional neural networks: an overview and application in radiology, *Insights into Imaging* 9 (2018) 611–629.
- [5] A. Kirillov, E. Mintun, N. Ravi, H. Mao, C. Rolland, L. Gustafson, T. Xiao, S. Whitehead, A. C. Berg, W.-Y. Lo, P. Dollár, R. Girshick, Segment anything, 2023. [arXiv:2304.02643](https://arxiv.org/abs/2304.02643).
- [6] C. D. Ruberto, A. Loddo, L. Putzu, A leucocytes count system from blood smear images, *Machine Vision and Applications* 27 (2016) 1151–1160.
- [7] S. Tavakoli, A. Ghaffari, Z. M. Kouzehkanan, et al., New segmentation and feature extraction algorithm for classification of white blood cells in peripheral smear images, *Scientific Reports* 11 (2021) 19428.
- [8] E. Mohamed, W. El-Behaidy, G. Khoriba, J. Li, Improved white blood cells classification based on pre-trained deep learning models, *Journal of Communications Software and Systems* 16 (2020).
- [9] N. Kaza, A. Ojaghi, F. E. Robles, Virtual staining, segmentation, and classification of blood smears for label-free hematology analysis, *BME Frontiers* (2022) 9853606.
- [10] J. B. Pal, A. Bhattacharyea, D. Banerjee, B. T. Maharaj, Advancing instance segmentation and wbc classification in peripheral blood smear through domain adaptation: A study on pbc and the novel rv-pbs datasets, *Expert Systems with Applications* 249 (2024) 123660.
- [11] Y. Li, L. Cai, Y. Lu, C. Lin, Y. Zhang, J. Jiang, G. Dai, B. Zhang, J. Cao, X. Zhang, X. Fan, Domain-invariant representation learning via segment anything model for blood cell classification, 2024. [arXiv:2408.07467](https://arxiv.org/abs/2408.07467).
- [12] A. Şengür, Y. Akbulut, Ümit Budak, Z. Cömert, White blood cell classification based on shape and deep features, in: 2019 International Artificial Intelligence and Data Processing Symposium (IDAP), 2019, pp. 1–4.
- [13] B. Hemalatha, B. Karthik, C. V. K. Reddy, A. Latha, Deep learning approach for segmentation and classification of blood cells using enhanced cnn, *Measurement: Sensors* 24 (2022) 100582.
- [14] A. Ojaghi, G. Carrazana, C. Caruso, A. Abbas, D. R. Myers, W. A. Lam, F. E. Robles, Label-free hematology analysis using deep-ultraviolet microscopy, *Proceedings of the National Academy of Sciences* 117 (2020) 14779–14789.
- [15] K. Alomar, H. I. Aysel, X. Cai, Data augmentation in classification and segmentation: A survey and new strategies, *Journal of Imaging* 9 (2023) 46.