

- [27] “Farm to Fork Strategy - Food Safety - European Commission.” Accessed: Oct. 15, 2025. [Online]. Available: https://food.ec.europa.eu/horizontal-topics/farm-fork-strategy_en

APPENDIX A: EVALUATION METRICS

To ensure consistency across the experimental evaluations presented in this thesis, the performance metrics used in the following technical chapters are defined in this section.

A.1 Classification metrics

For the classification task, four different metrics, accuracy, precision (P), recall (R), and F1-score, were used. The accuracy represents the correct predictions over the total number of predictions. Precision indicates the proportion of correct positive predictions and the total positive predictions, while recall indicates the proportion of correct positive predictions and actual positives. The F1-score shows the balance between precision and recall metrics for a model. These metrics were chosen as they are commonly used in binary classification models and are calculated as shown in Eqs. (A.1) - (A.4):

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (A.1)$$

$$P = \frac{TP}{TP + FP} \quad (A.2)$$

$$R = \frac{TP}{TP + FN} \quad (A.3)$$

$$F1 = \frac{2 \times P \times R}{P + R} = \frac{2 \times TP}{2 \times TP + FP + FN} \quad (A.4)$$

where TP is the number of true positives, TN is the number of true negatives, FP is the number of false positives, and FN is the number of false negatives.

A.2 Segmentation metrics

The segmentation output, which creates a black-and-white image highlighting the target insect on the image, was assessed using four different metrics: Intersection over Union (IoU), Dice Similarity Coefficient (DSC), precision, and recall. IoU and DSC values measure the ratio of the overlap and similarity of the predicted mask and the actual mask; IoU penalizes partial overlaps more than DSC, meaning that it is more sensitive to small mismatches.

Precision (P) and Recall (R) measure the accuracy of the model in classifying the pixels as black or white, considering their respective classes. Precision focuses on the ability of a model to minimize false positives by measuring the ratio of correct predicted positives (white pixels) to all predicted positives. Recall focuses on the ability of the model to minimize false negatives by calculating the ratio of correct positives to total actual positives. These metrics are calculated as shown in Eqs. (A.5) - (A.8):

$$P = \frac{TP}{TP + FP} \quad (A.5)$$

$$R = \frac{TP}{TP + FN} \quad (A.6)$$

$$IoU = \frac{P \times R}{P + R - P \times R} = \frac{TP}{TP + FP + FN} \quad (A.7)$$

$$DSC = \frac{2 \times P \times R}{P + R} = \frac{2 \times TP}{2 \times TP + FP + FN} \quad (A.8)$$

where TP is the number of true positives, FP is the number of false positives, and FN is the number of false negatives.

A.3 Counting metrics

The counting task was evaluated by Mean Squared Error (MSE) and Mean Absolute Error (MAE). MSE measures the mean squared difference between the number of predicted insects in the input image and the actual one. MAE evaluates the counting performance by computing the mean absolute difference between the predicted and actual number of insects. These errors are calculated as shown in Eqs. (A.9) - (A.10):

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (A.9)$$

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (A.10)$$

where n is the number of input images, y_i is the correct value, $\hat{y}_i$ is the predicted value.

In addition, R^2 is also considered for counting evaluation. R^2 measures the degree of how well the predicted insect number matches the actual insect number. This metric is calculated as shown in Eq. (A.11):

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (\text{A.11})$$

where n is the number of input images, y_i is the correct value, $\hat{y}_i$ is the predicted value and $\bar{y}$ is the mean of the correct value.

A.4 Incremental learning (IL) metrics

Incremental learning performance was evaluated using standard metrics, including Average Incremental Accuracy (AIA), Backward Transfer (BWT), and Forward Transfer (FWT). AIA shows the overall performance trend by averaging accuracy after each subset update, giving an indication of how well the model maintains performance as new data is added. BWT measures stability, showing whether training on later subsets helps or harms earlier ones, negative values indicate catastrophic forgetting. FWT reflects plasticity, the ability to generalize knowledge from earlier subsets to unseen ones. Positive values indicate helpful forward transfer, and negative values show interference. These metrics are calculated as shown in Eqs. (A.12) - (A.15):

$$AA_k = \frac{1}{k} \sum_{j=1}^k a_{k,j} \quad (\text{A.12})$$

$$AIA_k = \frac{1}{k} \sum_{i=1}^k AA_i \quad (\text{A.13})$$

$$BWT_k = \frac{1}{k-1} \sum_{j=1}^{k-1} (a_{k,j} - a_{j,j}) \quad (\text{A.14})$$

$$FWT_k = \frac{1}{k-1} \sum_{j=2}^k (a_{j,j} - \tilde{a}_j) \quad (\text{A.15})$$

Where $a_{k,j}$ is the accuracy evaluated on the test set of the j th after incremental learning of the k th ($j \leq k$) and $\tilde{a}_j$ is the classification accuracy of a randomly initialized reference model trained on task j .

A.5 Computational and resource efficiency metrics

In addition to the above-mentioned accuracy-based metrics, the performance of the models was also evaluated in terms of computational and resource efficiency to assess their feasibility for deployment on resource-constrained embedded systems.

The performance of the proposed DL models in terms of computation metrics, including model complexity, size, peak memory usage, and the number of parameters used in the model, was also assessed. These metrics are particularly important when deploying DL models on resource-constrained MCUs.

In this context, model size refers to the flash storage required to store the model on the MCU. RAM usage represents the runtime memory required during model execution. Therefore, peak memory usage, which indicates the maximum memory required by the model during inference, was used to assess the model in terms of its memory requirements on the MCU. The target deployment platform in this study was MCU-based devices with only a few megabytes of flash and RAM, which impose strict constraints on model size and peak memory usage.

In addition, model complexity was evaluated by the number of MACs (Multiply-Accumulate). MACs represent the number of multiplications and additions performed during one inference, which provides an estimate of the computational cost required for model inference.

The performance of the DL models was also assessed in terms of inference time, which represents the time required to process an input sample and produce a prediction on hardware. Moreover, the run time of one operational cycle was measured and evaluated to provide a comprehensive assessment of the whole system.

Finally, considering that the model should run on a battery-powered device, power consumption, as a critical factor affecting device lifetime, was also evaluated. In this thesis, three related metrics were considered when analyzing the energy behavior of the system: current consumption, power consumption, and energy consumption. Current consumption refers to the electrical current drawn by the system during operation and is typically measured in milliamperes (mA). Power consumption refers to the rate at which energy is consumed by the system (i.e., the amount of energy consumed per unit time). It depends on both the supply voltage and the current drawn

by the device and is typically measured in watts (W). Energy consumption refers to the total amount of energy used over a period of time or during a specific task, such as one complete operational cycle of the system, typically measured in joules (J) or watt-hours (Wh). Therefore, current consumption was primarily used to describe the electrical behavior of the hardware components, while energy consumption was used to evaluate the total energy required for specific operations or processing cycles of the proposed system.