

Combating Driver Drowsiness Using Machine Learning

Toby Lowe

N1105086

Supervised by: Dr Archontis Giannakidis

*A dissertation submitted in partial fulfillment of the requirements for
BSc (Hons) Data Science at Nottingham Trent University.*

April 2025

Abstract

Driver drowsiness is a key cause of fatal road accidents, and to combat it effectively, we need robust early detection methods. In this project, we place our focus on preventing these collisions and saving lives using both traditional and deep machine learning approaches.

We employ a traditional machine learning pipeline with custom engineered geometric features — Eye Aspect Ratio (EAR), Mouth Aspect Ratio (MAR), head yaw and tilt — calculated using facial landmarks; achieving $\sim 77\%$ test accuracy. Our deep learning methods consisted of a custom CNN that can detect eye state (open vs closed using a full face) at $\sim 98\%$ test accuracy and a fine-tuned ResNet-18 model that classifies cropped eye states (open vs closed) at $\sim 90\%$ test accuracy.

A functional Python dashboard, tested on the road, integrates an external iPhone camera and alerts when it detects micro-sleep in real-time. Other functionalities include live blink-rate metrics and a lighter-weight pre-drive drowsiness assessment tool, powered by a traditional machine learning stacked ensemble.

By blending traditional machine learning interpret-ability and adaptability with deep learning's superior performance, we achieve balanced and robust performance that can be deployed in our real-time dashboard. This brings to life genuine opportunities for improved driver and road safety.

Contents

1	Introduction	5
2	Literature Review	9
2.1	Drowsiness Detection Techniques Overview	9
2.2	Traditional Machine Learning Approaches	12
2.3	Deep Learning Approaches for Drowsiness Detection	17
3	Mathematical Formulations for More Robust Drowsiness	
	Classification	23
3.1	Data Preparation & Statistical Foundations	24
3.1.1	Person-Based Data Splitting Mathematics	24
3.1.2	Feature Normalisation with StandardScaler	29
3.1.3	Decision Boundary Optimisation Through F1-Score Maximization	33
3.2	Physiological Metrics and Geometrical Calculations	37
3.2.1	Eye Aspect Ratio (EAR)	37
3.2.2	Mouth Aspect Ratio (MAR) and Combined Features	38
3.2.3	Scale Invariance	39
3.3	Neural Network Architecture and Implementation	41
3.3.1	Overview of the Implemented CNN Architecture	41

3.3.2	Mathematical Foundations of Convolutional Layers . . .	44
3.3.3	Mathematical Foundations of Pooling Layers	46
3.3.4	Hierarchical Feature Extraction and Theoretical Frame- work	50
3.4	Understanding the Mathematical Foundations of Transfer Learn- ing when Classifying Eye States.	56
3.4.1	Formalising the Domains and Tasks in Transfer Learning	56
3.4.2	Unifying the Definition of Transfer Learning	59
3.4.3	Mechanisms of Knowledge Transfer in our Eye Classifi- cation Model	60
4	Methodology	63
4.1	Data Collection and Preprocessing	63
4.2	Feature Engineering and Model Architectures	64
4.2.1	XGBoost Pre-Drive Check	64
4.2.2	Full-Face CNN Eye-State Classifier	65
4.2.3	ResNet-18 Cropped-Eye Classifier	66
4.3	Live Drowsiness Detection System	66
4.3.1	System Architecture	66
4.3.2	Real-time Monitoring Features	67
4.3.3	Fallback Mechanisms	68
5	Results	69
5.1	XGBoost Model	69
5.1.1	Without Stacking	69
5.1.2	With Stacking	71

CONTENTS

5.1.3	Feature Importance	72
5.1.4	Notable Observations	74
5.2	Custom CNN for Drowsiness Detection: Results	74
5.2.1	Dataset and Preprocessing	74
5.2.2	Model Architecture	75
5.2.3	Training Details	76
5.3	ResNet Model (Eyes Open vs. Closed)	79
5.3.1	Architecture details (ResNet variant, input size):	79
5.3.2	Preprocessing (normalisation, augmentation):	80
5.3.3	Dataset Sizes	80
5.3.4	Hyperparameters:	81
5.3.5	Performance metrics:	81
5.3.6	Confusion matrix:	82
5.3.7	Notable observations:	83
5.4	Real-World Deployment Results	84
5.4.1	On-Road Test Setup	84
5.4.2	Functionality Performance	85
5.4.3	Qualitative Feedback & Observations	87
5.4.4	Prototype Demonstration	87
5.4.5	Summary of Real-World Deployment	88
6	Discussion	90
7	Conclusion	93
	References	95
.1	Code Listings	100

Acknowledgements

I'd like to express my gratitude to everyone who has supported me through this journey.

With special thanks to my brother, Jay, for road-testing the model whenever I asked. And to my girlfriend, Lucy, for her endless patience, encouragement, and invaluable feedback during my late-night coding sessions.

I'm also grateful to Dr. Giannakidis for his expert supervision and access to on-campus computers, and to my friends and colleagues at E.ON Next for the discussions and moral support.

Finally, thank you to my parents for believing in me—this dissertation would not have been possible without you.

Chapter 1

Introduction

Code/Datasets available at https://github.com/tobyjl/final_drowsiness_detection

<https://drive.google.com/drive/folders/1eVIF1hSIw5T4wncNKKkz5GDfFI74auRv?usp=sharing>

Before delving into the technical details, consider this warning:

“Drink a coffee. Have a rest. And survive.”

– *Road Safety Scotland (Road Safety Scotland, n.d.)*

Driver drowsiness poses one of the greatest yet most overlooked threats on our roads. While hazards such as speeding and drink-driving receive significant public attention, fatigue often escapes notice: unlike alcohol or drugs, there is no roadside test for tiredness, and its symptoms can remain undetected until it is too late. Despite these challenges, drowsiness continues to be a major

contributor to road accidents worldwide. This dissertation addresses this silent danger by developing a system that progresses from classical machine learning to deep learning, culminating in a real-time, deployable solution.

Table 1.1 summarises reported UK road casualties in 2022, with estimated contributions from driver drowsiness. Figure 1.1 visualises the gap between officially recorded cases (4%) and the estimated true scale (20%). Together, these data underscore both the scale of the problem and the mismatch between reported and actual fatigue-related crashes.

This mismatch is even more alarming when contrasted with what is legally permitted. In the UK, drivers may legally operate a vehicle for up to 4.5 hours before taking a 45-minute break, and then continue for another 4.5 hours (RAC, n.d). However, a particularly notable study by Ting et al. (Ting, Hwang, Doong, & Jeng, 2008) used simulated driving experiments to quantify the deterioration of alertness over time. Their methodology included measuring sleepiness ratings, reaction times, and performance scores throughout a 90-minute highway simulation. The findings, using principal component analysis (PCA), revealed a clear decline in driving stability over time. Performance scores at the eighth ($p = 0.009$) and ninth periods ($p = 0.001$) were more than double those at the beginning of the session. The authors concluded that "this experimental finding is strong evidence that optimum duration of safe highway driving is approximately 80 min"—far below the current legal threshold.

Severity	2022 Volume	Driver Drowsiness as a Contributing Factor (est*)†
Killed	1,711	342
Seriously injured	28,031	5,606
KSI	29,742	5,948
Slightly injured	105,738	21,147
All casualties	135,480	27,096

Table 1.1: Estimated casualties involving driver drowsiness in Great Britain (2022). Data source: Department for Transport 2022 ([Department for Transport, 2023](#)); Fatigue estimate: Department for Transport 2011 ([Department for Transport, 2011](#)).

* 20% estimate of crashes involving driver fatigue (RSWP21 ([Department for Transport, 2011](#))).

† Adjusted values for seriously injured, KSI, and slightly injured based on DfT casualty reporting methodology.

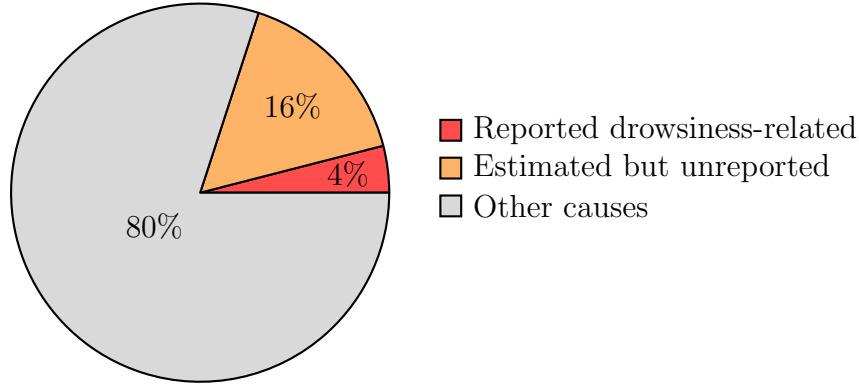


Figure 1.1: Estimated and reported contribution of drowsiness to road crashes in the UK. Officially recorded cases (4%) vs estimated true scale (20%).

Objectives and Contributions

- Develop a geometric feature-based XGBoost model achieving 77% accuracy in drowsiness detection.
- Design a custom CNN achieving $\geq 92\%$ accuracy for eye-state classification.
- Fine-tune ResNet-18 to reach $\geq 90.54\%$ accuracy in eye-state classification.
- Integrate models into a real-time Python-based dashboard with blink-rate metrics, drowsiness scoring, and audio alerts.
- Deploy and evaluate the system using a live iPhone camera, measuring inference latency and detection reliability.

Chapter 2

Literature Review

2.1 Drowsiness Detection Techniques Overview

Drowsiness detection techniques are commonly categorised into three primary classes: physiological, behavioural, and vehicular methods ([Ramzan et al., 2019](#)). Each class offers unique strengths and challenges in terms of accuracy, hardware complexity, and real-world applicability.

Physiological-Based Methods

EEG-Based Detection

Physiological methods involve tracking internal biological signals that reflect alertness levels. One of the most studied approaches uses electroencephalog-

raphy (EEG) to monitor brain activity. Arif et al. (Arif, Munawar, & Ali, 2023) proposed a passive brain-computer interface (pBCI) using EEG signals collected from six channels in the prefrontal and occipital cortices during simulated driving tasks. Their system achieved 85.6% accuracy, 89.7% recall, and an AUC of 91%, highlighting strong spatial localisation of drowsiness, though with the downside of high hardware complexity.

Heart Rate Variability Analysis

An alternative physiological approach utilises electrocardiogram (ECG) signals to monitor heart rate variability (HRV), particularly changes in R-R intervals. Fujiwara et al. (Fujiwara, Iwamoto, Hori, & Kano, 2024) implemented a self-attention autoencoder (SA-AE) model that identified fatigue-linked HRV abnormalities. Tested on 20 participants in a simulated driving setup, the method reached 88% recall with a low false positive rate of 0.6/hour, showing promise for wearable, non-invasive real-world applications.

Vehicular-Based Methods

Steering Pattern Analysis

Vehicular methods derive fatigue indicators from driving behaviour. Chai et al. (Chai, Li, Sun, Guo, & Huang, 2019) analysed eleven steering wheel-related parameters, selecting four key features through variance analysis. Their multilevel ordered logit (MOL) model achieved 72.92% accuracy, outperforming

support vector machines (SVM) and backpropagation (BP) neural networks at 63.86% and 62.10% respectively. The advantage stemmed from accounting for individual differences in steering patterns.

Behavioural-Based Methods

Eye Closure Monitoring

Behavioural methods infer drowsiness from physical cues such as eye closure, yawning, or facial movement. Junaedi et al. (Junaedi & Akbar, 2018) measured PERCLOS (Percentage of Eye Closure) using iris detection via circular Hough transforms and morphological operations on real-world driving footage. They tested thresholds P60, P70, and P80, finding values such as 83.66% accuracy for paired iris detection during talking and 74.41% during yawning, though performance dropped when the driver's face was out of frame.

Yawn Detection Systems

In another behavioural approach, Kielty et al. (Kielty, Dilmaghani, Ryan, Lemley, & Corcoran, 2023) proposed using neuromorphic vision systems—based on event camera technology—to detect yawning as an indicator of drowsiness. A custom CNN with self-attention and a recurrent head achieved 95.9% precision and 94.7% recall on a real yawning dataset, validating its effectiveness as a real-time input to driver monitoring systems. These three categories

form the foundational landscape of drowsiness detection systems. In the next section, we will examine how traditional machine learning techniques have historically tackled these challenges, with a particular focus on handcrafted feature pipelines.

2.2 Traditional Machine Learning Approaches

Geometric Facial Features

In traditional driver drowsiness detection, there is often a reliance on specific geometric facial features, such as the Eye Aspect Ratio (EAR), Mouth Aspect Ratio (MAR), and head pose angles. Eye Aspect Ratio (EAR) measures how open an eye is by making a comparison between vertical and horizontal distances between eye landmarks (Dewi et al., 2022).

$$\text{EAR} = \frac{\|p_2 - p_6\| + \|p_3 - p_5\|}{2, \|p_1 - p_4\|} \quad (2.1)$$

EAR scales upwards when eyes are open and approaches zero the more the eyes close (Dewi et al., 2022). By applying a threshold to this EAR, we are able to detect blinks. Furthermore, in a temporal method, if EAR stays below a threshold for multiple frames of a video, this would be a definitive indicator of microsleep (Dewi et al., 2022).

Mouth and Head Pose Analysis

Similarly, Mouth Aspect Ratio (MAR) quantifies mouth openness by taking distances between lip landmarks and comparing them: this gives us the ability to identify a yawn. MAR increases as you open your mouth to yawn (Xu, Huang, Liu, & Li, 2025), so by tracking MAR over time, it allows us to differentiate between normal speech (short mouth openings) and yawning (sustained high MAR), ideally using an adaptive threshold (Xu et al., 2025). Another key cue is head pose; a drowsy driver's head will likely tilt or nod. Head pose can be represented by Euler angles (e.g., pitch, yaw) computed via a Perspective-n-Point (PnP) algorithm such as POSIT, using facial landmarks (Xu et al., 2025). These features are extracted from video frames in real time; if done correctly, the computational cost remains low.

Classification Algorithms

These geometric features are brought to life in the drowsiness-detection ecosystem by utilising traditional machine learning classifiers to decide between drowsy and non-drowsy states. The most common algorithms include Support Vector Machines (SVM), K-Nearest Neighbors (KNN), Decision Trees, and simple Neural Networks. Researchers compute the features above (and others) over a time window and feed them into the chosen classifier (Xu et al., 2025). The review by Albadawi et al. (2022) provides a strong starting point for discovering recent developments in traditional machine learning approaches to drowsiness detection, helping identify emerging methodologies in the field

(Albadawi, Tahruri, & Awad, 2022).

Key Studies in Traditional ML Approaches

Real-Time EAR-Based Detection

In a notable study, Maier et al. (Souto Maier, Moura, Santana, & Lins, 2020) developed a real-time drowsiness detection system using EAR sequences from a webcam feed. They trained three classifiers—Multi-Layer Perceptron, SVM, and Random Forest—to distinguish prolonged blinks (indicative of drowsiness) from standard ones. The SVM achieved the best performance (94.9% test accuracy), outperforming both the neural net and decision forest. A key finding was that the SVM balanced accuracy with efficiency, making it suitable for practical, real-time applications, whereas the MLP showed overfitting and high computational cost. They also introduced a "personal feedback" mechanism: during use, the system gathers EAR data from the current user and retrains a new SVM, addressing person-based blink variability and significantly improving prediction accuracy in ambiguous eye states.

PERCLOS and Geometric Feature Systems

Bamidele et al. (Bamidele et al., 2019) developed a non-intrusive drowsiness detection system using geometric eye features extracted from the NTHU-DDD video dataset. After detecting face and eye regions, they computed PERCLOS, maximum eye-closure duration, and blink frequency over sliding windows,

then trained four classifiers (KNN, SVM, Logistic Regression, ANN). Their KNN (72.25%) and ANN (71.61%) models nearly matched the performance of a more complex fusion architecture (IAA: 73.06%) and outperformed several deep baselines (AlexNet, VGG-FaceNet, LRCN). This underscores that simple, well-chosen geometric descriptors plus lightweight classifiers can outperform heavier deep models—provided video quality is high—though the approach remains sensitive to lighting, head pose, and inter-individual blink variability.

Environment-Specific Detection

In a recent tunnel-driving study, Yuan et al. (Yuan et al., 2025) first used paired t-tests and Spearman correlation to identify three core fatigue indexes—blink frequency (BF), blink duration (TBD), mean blink-closure duration (MVBD)—along with two environment-specific metrics (cumulative and continuous tunnel proportions). They then fed these five interpretable features into an XGBoost classifier, training and evaluating it with 5 fold cross-validation. The model achieved 98% accuracy on unseen long-tunnel driving data (no drop across folds), demonstrating high predictive power and robustness. Feature-importance analysis confirmed that MVBD and continuous tunnel proportion were the strongest predictors, highlighting XGBoost’s ability to capture nonlinear interactions among heterogeneous inputs. This shows how rigorous statistical feature selection plus XGBoost’s efficient gradient-boosting can yield a lightweight yet highly generalizable fatigue detector—precisely the qualities we seek when extending drowsiness systems beyond simple eye-closure cues.

Summary and Limitations

These traditional approaches demonstrate respectable accuracy (70–95%) using relatively simple features and models, with very low computational requirements. This makes them ideal for embedded or in-car systems. Their main pitfall is handling complex real-world scenarios (e.g., varied lighting or facial occlusions), motivating researchers to explore deep learning methods for higher accuracy and generalisation.

Study	Extracted Features	Classifiers Evaluated	Best Model & Accuracy
Maior et al.	EAR sequences (sliding windows of EAR)	SVM, Random Forest, MLP	SVM: 94.9%
Bamidele et al.	PERCLOS, max eye-closure duration, blink frequency	KNN, SVM, Logistic Regression, ANN	KNN: 72.25%; ANN: 71.61%
Yuan et al.	BF, TBD, MVBD, continuous & cumulative tunnel proportions	XGBoost	XGBoost: 98%

Table 2.1: Summary of some traditional machine learning approaches for drowsiness detection and the top accuracies.

2.3 Deep Learning Approaches for Drowsiness Detection

With the developments made in the deep learning landscape, particularly in the past decade, there have been significant advantages brought to the table that allow greater capability than that of traditional machine learning for drowsiness detection. In our previous discussion, researchers had to manually define features such as eye aspect ratio and head pose - highly explainable in terms of results but a failure point is the inability to capture the complexity of driver behaviours; both drowsy and non-drowsy.

In contrast, deep learning methods automatically learn hierarchical feature representations from raw data. This means the subtle and complex patterns that make all the difference in accuracy and precision - the very patterns that traditional machine learning would typically miss, are now able to be captured (Knaak et al., 2021). A great example of this is that deep neural networks have the capability to learn very nuanced facial cues of fatigue. This can range from drooping eyelids, to micro-expressions, to even potentially things that the human eye cannot perceive, and it can do all of this without it being explicitly programmed. Indeed, in some cases, we do not have the available research, knowledge, or skill to explicitly program it if that is what is required. Knaak et al. notes that “deep learning models are capable of extracting more refined and complex image characteristics and are therefore expected to provide higher classification accuracies than conventional approaches based on feature engineering and traditional classifiers”.

Translating this into our problem, we find that deep learning models are delivering more impressive results than classical machine learning approaches, in the dual aspects of accuracy and generalisation. Importantly, in many cases and especially in the EEG signal space, deep learning has even discovered biological features we weren't previously aware of - Jian Cui et al. discovered that Alpha spindles and Theta burst in EEG were actually evidence of the drowsy state (Cui et al., 2022), a novel discovery we can attribute namely to the combined power of human engineering and deep learning potential.

Re-orienting to images and their classification, studies are finding that deep learning models are increasing accuracy through their ability to extract these more complex and effectively more informative image characteristics.

There have been multiple deep learning techniques relevant to the driver drowsiness detection solution research. One key approach (taken also in our work) is Convolutional Neural Networks (CNNs).

Convolutional Neural Networks

CNNs are by far the most common deep learning models for drowsiness detection via Image. In driver monitoring, CNNs automatically detect patterns such as eye state, yawning and other fatigue induced facial expressions (Albadawi et al., 2022). Pre-trained CNN architectures such as ResNet (an approach we employ) are being fine-tuned on driver datasets to improve their accuracy (Albadawi et al., 2022).

Research is showing that CNNs are achieving strong performance, the Al-

badawi et al. review finds CNN powered image systems are reporting accuracies spanning from 72% to 99% and above (Albadawi et al., 2022), hence they are forming the backbone of a bulk of the computer vision based drowsiness detectors due to their feature learning effectiveness in the face and eye region images.

Zhao et al. proposed in 2021 a fully automated fatigue detection method using a customised CNN architecture (Zhao et al., 2020). Their approach is compelling, and benchmarks their model against the state-of-the-art models out there such as AlexNet, VGG-16, GoogLeNet and ResNet-50. Their eye-mouth-CNN achieves the highest overall performance with an accuracy of 93.623%, sensitivity of 93.643% and specificity of 60.882%. This actually outperforms other networks across all metrics. The sensitivity score is more than important in the safety centered space that is driver drowsiness detection, where correctly identifying the drowsy instances is crucial and potentially life-saving.

3D Convolutional Neural Networks

Going a step further, 3D CNNs can be introduced to the equation. A 3D CNN extends upon the already established improvements of a CNN's exploration of spatial dimensions, but now factors in time. Existing 3D CNNs like I3D and one completed by Wijnands et al. (Wijnands, Thompson, Nice, et al., 2020) are able to pick up on motion based cues. Physical dynamics like the ones observed in blinking, yawning or head nodding, are ones that a 2D network simply cannot capture. A significant discovery is that when the eyes are

partially obscured, for example by sunglasses, temporal features still remain visible in a video stream so the 3D model maintains performance at $\sim 75.4\%$ accuracy.

Thanks to the architectural improvements, for example more efficient 3D kernels or depth-wise separable convolutions, and the ever-advancing power of mobile GPUs (and their ability to support video inference), 3D CNNs have their foot in the door in the future of Drowsiness Detection via Deep Learning.

Challenges in Deep Learning for Drowsiness Detection

Of course, deep learning based Drowsiness Detection is not without its challenges. A key advantage of deep learning—the scope and scalability to utilise big data—finds its own bottleneck; data availability and more so data quality. We are seeing very impressive results recorded, but Albadawi et al. notes that “those results do not necessarily represent real-life driving situations” as the results are being recorded in virtual and simulated environments ([Albadawi et al., 2022](#)), or at best carefully curated and optimised-for-success recording scenarios.

Additionally, deep learning is unable to escape the problems complex machine-learning faces in general; overfitting and generalisation. With dozens of layers, and parameters in the thousands to millions, by definition deep machine learning is complex and is susceptible to memorising training data as opposed to learning generalisable patterns. This issue was one that my study was

heavily geared to combat, but reviewing the notebooks working with the Driver Drowsiness Dataset available on Kaggle ([Nasri, 2022](#)), and we see implementations plagued by overfitting.

Sayed holds the most highly rated notebook associated with the dataset, but reports a validation accuracy of 99.9% ([Sayed, 2024](#)). This level of accuracy would be groundbreaking, but clearly shows signs of overfitting and data leakage. The proposed implementation fails to combat data-leakage and we see the case-in-point—the highly complex CNN is recognising the characteristics of each video used for the dataset and classifying using the wrong features. My implementation successfully approaches this issue.

Furthermore, there are the real-time performance and hardware constraints, where deploying deep learning models in an actual vehicle becomes problematic. High performing models in research are of course computationally heavy, for example a large CNN or CNN-LTSM ensemble, and are evaluated on desktop GPUs which are often not available on a typical driver’s home computer, nevermind in their car. Without acceleration, a deep model could introduce delays that actually make the warning too late to be useful ([Albadawi et al., 2022](#)).

Wijnands et al. notes in their future directions that they recommend the integration of AI chips within smartphones to allow machine learning functionality ([Wijnands et al., 2020](#)). This is significant and relevant to our implementation, where we use an iPhone camera connected to a Macbook. Further deployment of AI chips in iPhones creates future scope for my Deep Learning Model and even more advanced ones to be integrated successfully into the iPhone itself,

meaning the incredibly common setup of a phone attached to a windshield is suddenly a playground for driver-drowsiness detection. Even Apple’s budget option (the iPhone 16e), is “incorporating the latest A18 chipset” (Chiew & Hong, 2025).

In the meantime, deep learning models struggle to fit into deployable real-time driver-drowsiness detection systems and dashboards due to this constraint. We explore what they now look like in practice in the next section.

Chapter 3

Mathematical Formulations for More Robust Drowsiness Classification

Introduction to Mathematical Foundations

This chapter aims to establish the mathematical foundations that back the developed driver drowsiness system. Transforming raw facial data into binary classifications of drowsy and non-drowsy is supported by mathematics at essentially every step, and this project comprehensively inducts statistical, geometric, and algorithmic formulations to enable us to robustly detect across a diverse array of individuals.

The chapter begins by looking at statistical foundations for data prepara-

tion - this includes person-based validation methodologies as well as feature standardisation techniques.

Then, we will explore the geometric relationships between facial landmarks which quantify physiological indicators of drowsiness.

To conclude, we break down and formalise our neural network architectures, and explore the mathematics behind Transfer Learning.

Through this chapter, we will illustrate how this theoretical rigor entails higher performance in drowsiness detection, and what this does to ensure a robust and generalised approach in a safety-critical scenario.

3.1 Data Preparation & Statistical Foundations

3.1.1 Person-Based Data Splitting Mathematics

Person-Based Data Splitting: Mathematical Formulation

Theoretical Foundation

We can build on Stone's (1974) (Stone, 1974) foundational work on cross-validation in order to formalise our approach involving person-based data

splitting. Stone defined cross-validators assessment as:

$$C^n(\alpha) = \frac{1}{n} \sum_{i=1}^n L[y_i, f(x_i; \alpha, S_{(i)})] \quad (3.1)$$

Where:

- $S_{(i)}$ denotes the sample with the i -th item omitted
- $f(x_i; \alpha, S_{(i)})$ is the prediction for the i -th item using a model parameterized by α and trained on $S_{(i)}$
- L is a loss function that measures prediction accuracy

Adaptation for Person-Based Data Partitioning

In contrast to a traditional leave-one-out cross-validation, we partitioned our subjects in a fixed manner which addresses the identity-based data leakage in our drowsiness detection system. Let us define:

- $P = \{p^1, p^2, \dots, p^M\}$ be the set of M distinct persons
- $D_j = \{(x_i, y_i) | \text{person}(i) = p^j\}$ be all of the samples belonging to person j

We partition P into three disjoint subsets:

$$P = P_{train} \cup P_{val} \cup P_{test} \quad (3.2)$$

$$P_{train} \cap P_{val} = P_{train} \cap P_{test} = P_{val} \cap P_{test} = \emptyset \quad (3.3)$$

Where we define:

- P_{train} contains approximately $r_{train} \cdot M$ persons (in our implementation, $r_{train} = 0.7$)
- P_{val} contains approximately $r_{val} \cdot M$ persons ($r_{val} = 0.15$)
- P_{test} contains approximately $r_{test} \cdot M$ persons ($r_{test} = 0.15$)

This creates the following data partitions:

$$D_{train} = \bigcup_{j \in P_{train}} D_j \quad (3.4)$$

$$D_{val} = \bigcup_{j \in P_{val}} D_j \quad (3.5)$$

$$D_{test} = \bigcup_{j \in P_{test}} D_j \quad (3.6)$$

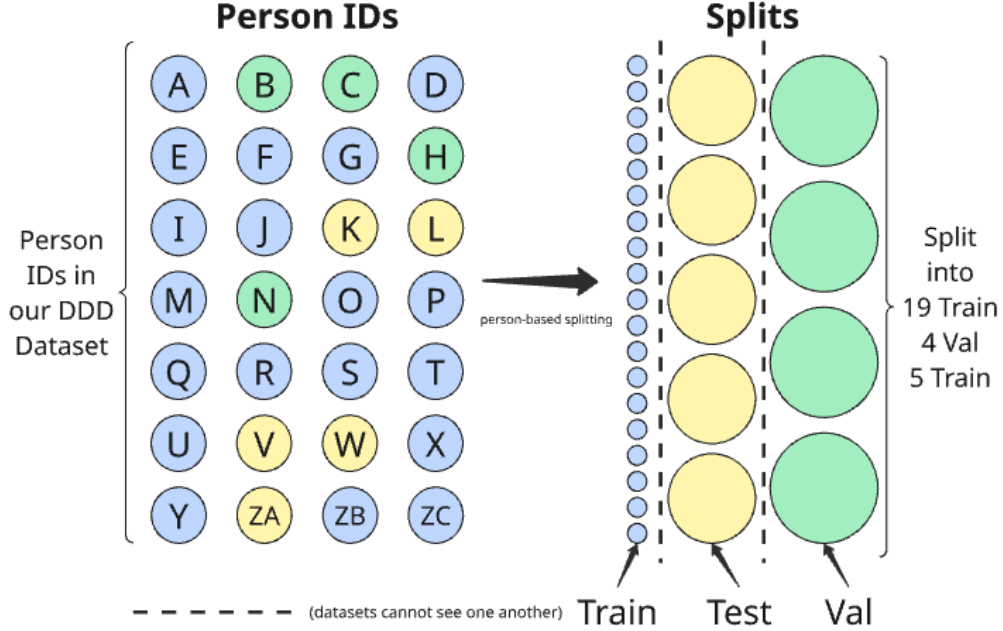


Figure 3.1: Illustration of our Person-Based Splitting Strategy. Unique Person IDs in our drowsiness dataset (DDD Dataset) are split in a way that no individual appears in more than one subset, preventing identity-based data leakage to allow our the model to generalise to unseen individuals. Split is randomised per run.

Then, the model f is trained using parameters α on D_{train} and evaluated on D_{val} to yield a performance metric:

$$\text{Performance}(\alpha) = \frac{1}{|D_{val}|} \sum_{(x_i, y_i) \in D_{val}} L[y_i, f(x_i; \alpha, D_{train})] \quad (3.7)$$

Using the above formulation, we are mathematically capturing how our approach successfully prevents person-based data leakage. It does this by

completely ensuring separation of the many different individuals between training, validation, and testing sets.

What impact does this have on the model’s evaluation?

By partitioning using a person-based approach, we are gaining a far more realistic estimate of the model’s performance in terms of its ability to generalise. We are explicitly accounting for variation between the different individuals as opposed to just between samples. In drowsiness detection, this is especially crucial, as individual differences in facial expressions or just physical characteristics can indeed significantly affect the feature distribution. This is computationally far more efficient than a full leave-one-person-out cross-validation approach, but maintains the principle of directly testing the model’s ability to generalise to previously unseen data and therefore individuals. In practice, this translated to its impressive ability to perform accurately in real-life road testing.

3.1.2 Feature Normalisation with StandardScaler

Mathematical Formulation of Standardisation

In the context of machine learning applications, feature standardisation is the process of transforming variables to a common scale with both zero mean and unit variance. This follows the standard normal distribution $\mathcal{N}(0, 1)$. This process (which is also known as z-score normalisation) can be mathematically defined like this ([scikit-learn](#), n.d.):

$$z = \frac{x - \mu}{\sigma} \quad (3.8)$$

Where we define such that:

- z represents our standardised feature value
- x represents the original feature value
- μ represents the mean of the training samples
- σ represents the standard deviation of our training samples

When we apply this to a dataset with multiple features like our xg-boost model, standardisation will operate independently on each of the feature columns. So, for a feature matrix $X \in \mathbb{R}^{n \times d}$ with n samples and d features we can have the standardisation for the j -th feature defined as follows:

$$X_{\text{standardised}}^{(j)} = \frac{X^{(j)} - \mu^{(j)}}{\sigma^{(j)}} \quad (3.9)$$

The parameters $\mu^{(j)}$ and $\sigma^{(j)}$ are then estimated from the training data:

$$\mu^{(j)} = \frac{1}{n_{train}} \sum_{i=1}^{n_{train}} X_{train,i}^{(j)} \quad (3.10)$$

$$\sigma^{(j)} = \sqrt{\frac{1}{n_{train}} \sum_{i=1}^{n_{train}} (X_{train,i}^{(j)} - \mu^{(j)})^2} \quad (3.11)$$

We compute these statistics during the fitting phase and store them for subsequent transformation of any data. This does include both validation and test sets ([scikit-learn](#), n.d.).

Implementation in Drowsiness Detection Framework

Pivoting to our drowsiness detection system and feature standardisation is applied to the extracted facial landmark features before training the model:

```
scaler = StandardScaler().fit(X_train)
X_train_s = scaler.transform(X_train)
X_val_s = scaler.transform(X_val)
X_test_s = scaler.transform(X_test)
```

You can split the standardisation into 2 phases:

- **Fitting Phase:** The `fit()` method computes the mean $\mu^{(j)}$ and standard deviation $\sigma^{(j)}$ for every feature in our training set.

- **Transformation Phase:** The `transform()` method then applies the standardisation formula using the previously mentioned stored statistics in order for us to transform the training, validation, and test data.

It is important to mention that parameters estimated from our training data are applied consistently to all of the datasets, as this preserves the relative relationships.

What are the theoretical implications for model performance?

Standardisation provides us with a plethora of mathematical and computational advantages in our context:

- **Scale-Invariant Feature Importance:** It follows common sense that the primary features in drowsiness detection (i.e. eye aspect ratio, mouth aspect ratio and more) will naturally operate on different scales. By using standardisation, we get to lock in that every feature will contribute proportionally to the distance calculations. This means that the features with larger magnitudes will not be able to dominate the model ([scikit-learn](#), n.d.).
- **Stability in Optimisation:** In a gradient-based implementation like our XGBoost model, it assumes that the features are centred around zero with similar variances. The error surface will become more spherical when this is met which allows for more efficient optimisation and this means we have a consistent learning rate across all dimensions.

- **Statistical Distribution Alignment:** Many machine learning algorithms will actually perform better when the input features follow a more Gaussian distribution. So in the context of our very diverse facial features, we're transforming them to more closely resemble this distribution. In their raw form, they will naturally take on different scales and distributions, for example eye aspect ratio will typically range from 0.15 to 0.35, whereas head position angles can range from negatives to positives. By performing standardisation, we bring all of the different measurements into a common, statistical framework, allowing us to put them in more effective direct comparison.

In simpler terms, standardisation gives the model a common reference frame whereby a unit change in any feature will represent a movement of one standard deviation within that feature's distribution.

3.1.3 Decision Boundary Optimisation Through F1-Score Maximization

Mathematical Formulation of F1-Score

We use the F1-score as a harmonic mean of precision and recall and it is most commonly used for evaluating binary classifiers such as ours. It is especially geared for dealing with imbalanced datasets. It can be mathematically expressed as the ratio of true positives, to the sum of true positives plus half the sum of false positives and false negatives, or notated as:

$$F1 = \frac{2tp}{2tp + fp + fn} \quad (3.12)$$

where we can say tp , fp , and fn represent true positives, false positives, and false negatives respectively.

Optimal Threshold for F1 Maximisation

There is a theoretical relationship (for classifiers that produce probabilistic outputs) between the maximum achievable F1 value, and the optimal decision threshold. As [Lipton, Elkan, & Naryanaswamy, 2014](#) demonstrate in their paper “Optimal Thresholding of Classifiers to Maximize F1 Measure”, when a classifier outputs calibrated probabilities $s = p(t = 1|score)$, we assign an instance to the positive class iff

$$s \geq \frac{F^*}{2} \quad (3.13)$$

where F^* is the maximum F1 measure achieved by this optimal decision rule. From this we imply that the optimal threshold $\tau^* = F^*/2$, is therefore half

the maximum achievable F1 score.

In a more general case, for any classifier that outputs real-valued scores, [Lipton et al., 2014](#) establish in Theorem 1 that an example with score s should be assigned to the positive class iff:

$$\frac{b \cdot p(s|t=1)}{(1-b) \cdot p(s|t=0)} \geq J \quad (3.14)$$

where $J = \frac{tp}{fn+tp+fp}$ is the Jaccard index of the optimal classifier. This is a value which is directly related to the F1 score. b is the base rate $p(t=1)$, and $p(s|t=1)$ and $p(s|t=0)$ are the likelihoods of score s given the true class.

Implementation in Drowsiness Detection Framework

We implemented an empirical search in our drowsiness detection system in order to find optimal threshold using the validation data:

```
# Threshold optimisation
best_f1, best_thresh = 0, 0.5
for t in np.arange(0.5, 0.9, 0.05):
    f1 = f1_score(y_val, (probs>=t).astype(int))
    if f1 > best_f1:
        best_f1, best_thresh = f1, t

# Final prediction with the optimised threshold
probs_test = model.predict_proba(X_test_s)[: ,1]
pred_test = (probs_test>=best_thresh).astype(int)
```

Listing 3.1: Threshold Optimisation and Final Prediction

This implementation aligns with the theoretical findings of Lipton et al. in several ways...

Firstly, we iterated through a range of candidate thresholds (0.5, 0.55, ..., 0.85) and evaluated the F1 score achieved by each threshold on the validation set.

The threshold that gave us the maximum F1 score was selected as the optimal threshold $\tau^* = \text{best_thresh}$.

By doing this approach, we address what Lipton et al. refer to as the “batch observation” (Lipton et al., 2014). This is the insight where we can say optimal thresholding for F1 depends not only on a specific instance’s probability, but also on the distribution of probabilities across all instances in the batch.

What are the Theoretical Implications for Model Performance?

There are many different and crucial theoretical implications for our drowsiness detection model using the threshold optimisation process:

Calibration Assessment: If we find that our XGBoost model produces well calibrated probabilities, then we can estimate the theoretical optimal threshold should be approximately $\tau^* \approx F^*/2$. By comparing our empirically found `best_thresh` with `best_f1/2`, we get insights into our model’s calibration quality.

Handling Class Imbalance: Lipton et al., 2014 illustrates that the F1 score is especially useful for imbalanced datasets. We adapt to our actual class

distribution using our empirical threshold search in the drowsiness detection task rather than using a suboptimal value of 0.5 we have assigned essentially blindly.

Mitigating the Data Shift Effects: A common challenge in machine learning is when the statistical properties of the data change on deployment in comparison to training. In our case, drowsiness patterns may change through different times of day, or lighting, camera angles, and other factors might differ too. By using a separate validation set to optimise the threshold, we better represent the real-world deployment data shift that we expect. This means we’re optimising our data closer to the real-world conditions we are deploying in.

Comparing with Alternative Approaches

There are alternative approaches we could have taken, such as fixing the threshold at 0.5 (which would incorrectly assume we have equal misclassification costs and balanced class distributions), or a precision-recall curve analysis, which would allow us to balance the precision recall according to a more application specific requirement. However, by using the F1 maximisation approach, we are justified by Lipton et al.’s theoretical foundation and also the very nature of drowsiness detection where false positives and false negatives have significant and non-symmetrical costs.

3.2 Physiological Metrics and Geometrical Calculations

3.2.1 Eye Aspect Ratio (EAR)

Arguably the most critical metric for drowsiness detection, we have EAR (Eye Aspect Ratio). We calculate this using specific facial landmarks extracted using dlib python library. The purpose of the eye aspect ratio is to capture the relationship between the eye's height and the eye's width ([Varun Shiva Krishna Rupani & Tejith, 2024](#)).

$$EAR = \frac{||p_2 - p_6|| + ||p_3 - p_5||}{2 \cdot ||p_1 - p_4||} \quad (3.15)$$

Our implementation processes this calculation through the `calculate_eye_aspect_ratio()` function:

```
def calculate_eye_aspect_ratio(eye_points):  
    a = distance.euclidean(eye_points[1], eye_points[5])  
    b = distance.euclidean(eye_points[2], eye_points[4])  
    c = distance.euclidean(eye_points[0], eye_points[3])  
    return 0 if c == 0 else (a + b) / (2.0 * c)
```

Listing 3.2: Eye Aspect Ratio Calculation

Using this, we can compute the ratio between how vertically open the eye is, and how wide the eye is horizontally. We measure the vertical openness by calculating the euclidean distances between the upper and lower eyelids,

and for the eye width it is the euclidean distance between the horizontal endpoints. According to [Varun Shiva Krishna Rupani & Tejith, 2024](#), this ratio significantly decreases when eyes close during drowsiness episodes, providing a reliable indicator for detection.

3.2.2 Mouth Aspect Ratio (MAR) and Combined Features

To capture yawning behaviour which [Varun Shiva Krishna Rupani & Tejith, 2024](#) cites as another significant indicator of drowsiness, we extend upon the EAR by bringing in MAR (Mouth Aspect Ratio). We capture this in our implementation through this function:

```
def calculate_mouth_aspect_ratio(mouth_points):  
    a = distance.euclidean(mouth_points[2], mouth_points  
        [10])  
    b = distance.euclidean(mouth_points[3], mouth_points  
        [9])  
    c = distance.euclidean(mouth_points[4], mouth_points  
        [8])  
    d = distance.euclidean(mouth_points[0], mouth_points  
        [6])  
    return 0 if d == 0 else (a + b + c) / (3.0 * d)
```

Listing 3.3: Mouth Aspect Ratio Calculation

This measures 3 vertical distances between lower lips and upper lips, averages

them, and then divides by the horizontal mouth width. This creates a scale-invariant ratio. As you can therefore imagine, when you yawn, the MAR value will increase significantly and correspondingly (Varun Shiva Krishna Rupani & Tejith, 2024).

We’ve further enhanced detection accuracy by introducing a relatively novel combined metric—the Eye-Mouth Ratio (EMR):

```
features['emr'] = avg_ear / (mar + 0.01)
# 0.01 addition avoids division by zero
```

Listing 3.4: Eye-Mouth Ratio (EMR) Feature Calculation

Using this ratio, we’re leveraging the inverse relationship between eye closure and mouth openness that become especially apparent during drowsiness episodes. This metric is more sensitive than the individual use of either of the others.

3.2.3 Scale Invariance

We take a multi-faceted approach to make sure we handle scale invariance further:

We normalise the head pose features against face width, giving us consistency across the variety of both different face sizes and also distances between metrics:

```
face_width = distance.euclidean(points[0], points[16])

features['head_pitch'] = \
    distance.euclidean(nose_tip, chin) / face_width

features['head_yaw'] = abs(left_dist - right_dist) /
    face_width
```

Listing 3.5: Head Pose Feature Extraction

For us to mitigate noise and individual eye variations - we computed the EAR for *both* eyes and take the average value.

```
left_ear = calculate_eye_aspect_ratio(points[36:42])
right_ear = calculate_eye_aspect_ratio(points[42:48])
avg_ear = (left_ear + right_ear) / 2.0
```

Standardisation Pipeline: As previously discussed, our preprocessing pipeline applies StandardScaler to normalise all features before model training.

Varun Shiva Krishna Rupani & Tejith, 2024 identify a major challenge in drowsiness detection systems: how can we maintain consistently high performance across different subjects where lighting conditions, physiological variation, and camera positions? By bringing in these dynamic and flexible geometrical calculations, and further addressing scale invariance, we help bridge this issue.

3.3 Neural Network Architecture and Implementation

In processing image data, CNNs have proven their success as a technique in comparison to regular neural networks. This is likely due to their ability to automatically learn hierarchical representations of features (Goodfellow, Bengio, & Courville, 2016). In this section, we detail my specific customised CNN and provide ourselves with the mathematical foundations of its feature extraction process. To explain the increase in the channels and the spatial dimension reductions, we aim to trace the evolution of feature maps through the different network layers using established mathematical principles. Throughout this, I draw and build on the foundational works of Goodfellow et al. (Goodfellow et al., 2016), and then break into more advanced theoretical insights from Wiatowski & Bölcskei, 2018.

3.3.1 Overview of the Implemented CNN Architecture

My system's core is the custom CNN, and I implement it using PyTorch. The architecture, which I define in the `DrowsinessDetectionCNN` class, has been tailored to process grayscale images of full faces. The target is to classify from a full face picture whether their eyes are open or their eyes are closed. Any input image is resized to a standard dimension of 224x224 (height x width), with a single channel to use grayscale.

CHAPTER 3: MATHEMATICAL FORMULATIONS FOR MORE ROBUST DROWSINESS CLASSIFICATION

Network Structure For reference, my network has a structure made up of 3 main convolutional blocks, followed by fully connected layers to perform classification. In more detail we have:

Block	Layer	Output / Params	Notes
Block 1	Conv2d:	(8, 224, 224)	—
	$1 \rightarrow 8, 3 \times 3, \text{pad} = 1$		
	BatchNorm2d: 8	(8, 224, 224)	—
	ReLU	(8, 224, 224)	—
	MaxPool2d: 2×2 , stride 2	(8, 112, 112)	—
Block 2	Conv2d:	(16, 112, 112)	—
	$8 \rightarrow 16, 3 \times 3, \text{pad} = 1$		
	BatchNorm2d: 16	(16, 112, 112)	—
	ReLU	(16, 112, 112)	—
	MaxPool2d: 2×2 , stride 2	(16, 56, 56)	—
Block 3	Conv2d:	(32, 56, 56)	—
	$16 \rightarrow 32, 3 \times 3, \text{pad} = 1$		
	BatchNorm2d: 32	(32, 56, 56)	—
	ReLU	(32, 56, 56)	—
	MaxPool2d: 2×2 , stride 2	(32, 28, 28)	—

Continued on next page...

Block	Layer	Output / Params	Notes
Classification Head	Flatten:	25088	—
	$32 \times 28 \times 28 \rightarrow 25088$		
	Dropout (p)	25088	—
	FC1: $25088 \rightarrow 64$	64	BN, ReLU, Dropout
	FC2: $64 \rightarrow 32$	32	BN, ReLU
	FC3: $32 \rightarrow 2$	2	(Binary out)

Table 3.1: CNN Architecture Overview

Implementation Details My architecture contains standard components like Batch Normalisation to stabilise its learning and uses ReLU activations to introduce non-linearity. The complete implementation of the DrowsinessDetectionCNN class is provided in Listing 1 in Appendix A. The implementation follows the architecture described in Table 3.1, with three convolutional blocks (each containing Conv2d, BatchNorm2d, ReLU, and MaxPool2d), followed by a classification head with three fully connected layers for the eye state classification task.

3.3.2 Mathematical Foundations of Convolutional Layers

Cross-Correlation Operation The core of the vast majority of CNNs is the convolution operation. This is the important component that allows our network to learn the different spatial hierarchies of features. When we put this in the context of machine learning, the operation that is implemented and referred to as 'convolution' is technically cross-correlation ([Goodfellow et al., 2016](#)). Let us take an input image I and a 2D kernel K , the discrete cross-correlation output S at location (i, j) can be defined as follows:

$$S(i, j) = (K \star I)(i, j) = \sum_m \sum_n I(i + m, j + n)K(m, n) \quad (3.16)$$

$$\text{(adapted from ([Goodfellow et al., 2016](#)) Equation 9.6)} \quad (3.17)$$

In the above, the kernel K slides across the input I , and output $S(i, j)$ (a point in the output feature map) is then computed as the sum of element-wise products between the kernel and the overlapping input region. While this is mathematically distinct from true convolution, where we would involve flipping the kernel, cross correlation is standard in deep learning libraries. As we know the kernel K contains learnable parameters, the learning algorithm will appropriately adapt the weights in order to make the functional difference often negligible in practice as per [Goodfellow et al., 2016](#), p. 329.

Multi-Channel Convolution Modern CNNs will generally operate with multiple-channel inputs and multi-channel outputs. In our implementation, the input to the second convolutional layer has 8 channels and produces an

output with 16 channels, involving a 4D kernel tensor. Let V be the input tensor with the following dimensions; (C_{in}, H_{in}, W_{in}) (Input Channels, Height, Width). Let K be the 4D kernel tensor with dimensions $(C_{out}, C_{in}, F_h, F_w)$ (Output Channels, Input Channels, Kernel Height, Kernel Width). The output tensor Z with dimensions $(C_{out}, H_{out}, W_{out})$ is computed. For clarity in the formula, we assume cross correlation and 1-based indexing (as in Goodfellow et al.) where the element $Z_{i,j,k}$ (output channel is i , row is j , and column is k) can be given as:

$$Z_{i,j,k} = \sum_l \sum_{m,n} V_{l,j+m-1,k+n-1} K_{i,l,m,n} \quad (3.18)$$

$$(\text{adapted from (Goodfellow et al., 2016) Equation 9.7}) \quad (3.19)$$

Using this formula, we are highlighting that each output channel i is derived from a sum over all of the input channels l and weighted by the corresponding kernel slices $K_{i,l,:}$. By performing this summation across the input channels, we are essentially allowing the layer to connect the information from the different feature maps which are generated by the previous layer.

Spatial Dimensions Calculation The output feature map's spatial dimensions (H_{out}, W_{out}) directly depend on the dimensions of the input (H_{in}, W_{in}) , the size of our kernel (F_h, F_w) , the padding (P_h, P_w) , and the stride (S_h, S_w) applied during the convolution as per Goodfellow et al., 2016, Ch. 9.5. In general, the formulae for calculating our output spatial dimensions are:

$$H_{out} = \text{floor} \left(\frac{H_{in} - F_h + 2P_h}{S_h} \right) + 1 \quad (3.20)$$

$$W_{out} = \text{floor} \left(\frac{W_{in} - F_w + 2P_w}{S_w} \right) + 1 \quad (3.21)$$

All convolutional layers in our architecture (conv1-3) make use of a kernel sized $F_h = F_w = 3$, padding size $P_h = P_w = 1$, and a stride size of $S_h = S_w = 1$. So if we apply these parameters to Equation 3.20 and Equation 3.21 to give us:

$$H_{out} = \text{floor} \left(\frac{H_{in} - 3 + 2 \times 1}{1} \right) + 1 = H_{in}$$

$$W_{out} = \text{floor} \left(\frac{W_{in} - 3 + 2 \times 1}{1} \right) + 1 = W_{in}$$

Therefore prior to the pooling, the convolutional operations themselves will preserve the spatial dimensions of their input feature maps. Primarily, the change in representation within these layers originates from the transformation across the channels, as defined by Equation 3.18. An example of this is how conv1 transforms the $(1, 224, 224)$ input into an $(8, 224, 224)$ feature map.

3.3.3 Mathematical Foundations of Pooling Layers

Pooling Operation Although not nearly the core component the convolutional layer is, pooling layers remain a common component of a CNN. We employ them to reduce the spatial dimensions of the feature map, with benefits such as decreasing the computational complexity, and introducing some degree of local translation invariance as per Goodfellow et al., 2016, Ch. 9.3 and (Wiatowski & Bölskei, 2018). If we have a neighbourhood of features in the input map, a pooling function will summarise them into a singular value in the output map. In my architecture, I utilise Max Pooling (MaxPool2d). This operation when applied with a kernel/window of size (F_{ph}, F_{pw}) and stride (S_{ph}, S_{pw}) will output the maximum value within each window as it slides across the input feature map. Keep in mind that that

pooling will operate independently on each channel, hence it preserves the channel's depth ($C_{out_pool} = C_{in_pool}$).

Pooling Dimensions Calculation The calculation of the output spatial dimensions follows the same principle as convolution (Equation 3.20 and Equation 3.21), typically assuming zero padding for the standard pooling operations like our `MaxPool2d`. So if we take pooling layer with kernel size F_p and stride S_p :

$$H_{out_pool} = \text{floor} \left(\frac{H_{in_conv} - F_{ph} + 2P_{ph}}{S_{ph}} \right) + 1 \quad (3.22)$$

$$W_{out_pool} = \text{floor} \left(\frac{W_{in_conv} - F_{pw} + 2P_{pw}}{S_{pw}} \right) + 1 \quad (3.23)$$

In our implementation, the pooling layers use $F_{ph} = F_{pw} = 2$ and $S_{ph} = S_{pw} = 2$ (stride 2 is the default for `MaxPool2d(2)`). Then if we apply these with $P_p = 0$ to Equation 3.22 and Equation 3.23 we get:

$$H_{out_pool} = \text{floor} \left(\frac{H_{in_conv} - 2}{2} \right) + 1$$

$$W_{out_pool} = \text{floor} \left(\frac{W_{in_conv} - 2}{2} \right) + 1$$

Since in each stage in our network the input dimensions H_{in_conv} and W_{in_conv} are even in our network, we can simplify these formulas to become $H_{out_pool} = H_{in_conv}/2$ and $W_{out_pool} = W_{in_conv}/2$. Therefore, each pooling layer in our architecture, in effect, will half the height and the width of the feature map, while still retaining the channel's depth.

Practical Implications for Eye State Classification There are important practical implications for our eye state classification task. By reducing

the spatial dimensions, we get the dual benefit of decreased computational requirements and increase in the effective receptive field of the following layer. This means that as the information flows deeper and deeper into the network, the many neurons can 'see' bigger portions of the original image. This means we can detect relationships between features further apart. For complex patterns like distinguishing between open and closed eyes within the context of a full face, this becomes critical.

Feature Map Evolution in the CNN

By applying the mathematical operations described above, we can trace the evolution of the feature map dimensions as data propagates through the convolutional blocks of our implemented `DrowsinessDetectionCNN`.

Using this step-by-step evolution, we can show how the spatial resolution progressively reduces from 224x224 to 112x112 to 56x56 to 28x28, and how the channel depth grows simultaneously from 1 to 8 to 16 to 32. This architectural pattern is fundamental in the way that CNNs build the hierarchical representations.

Figure 3.2 is a flowchart illustrating the evolution of feature map dimensions (Channels, Height, Width) through the convolutional blocks and flattening stage of my implemented CNN architecture. The process starts with input X_0 and results in the final feature vector X_3 before the classification head.

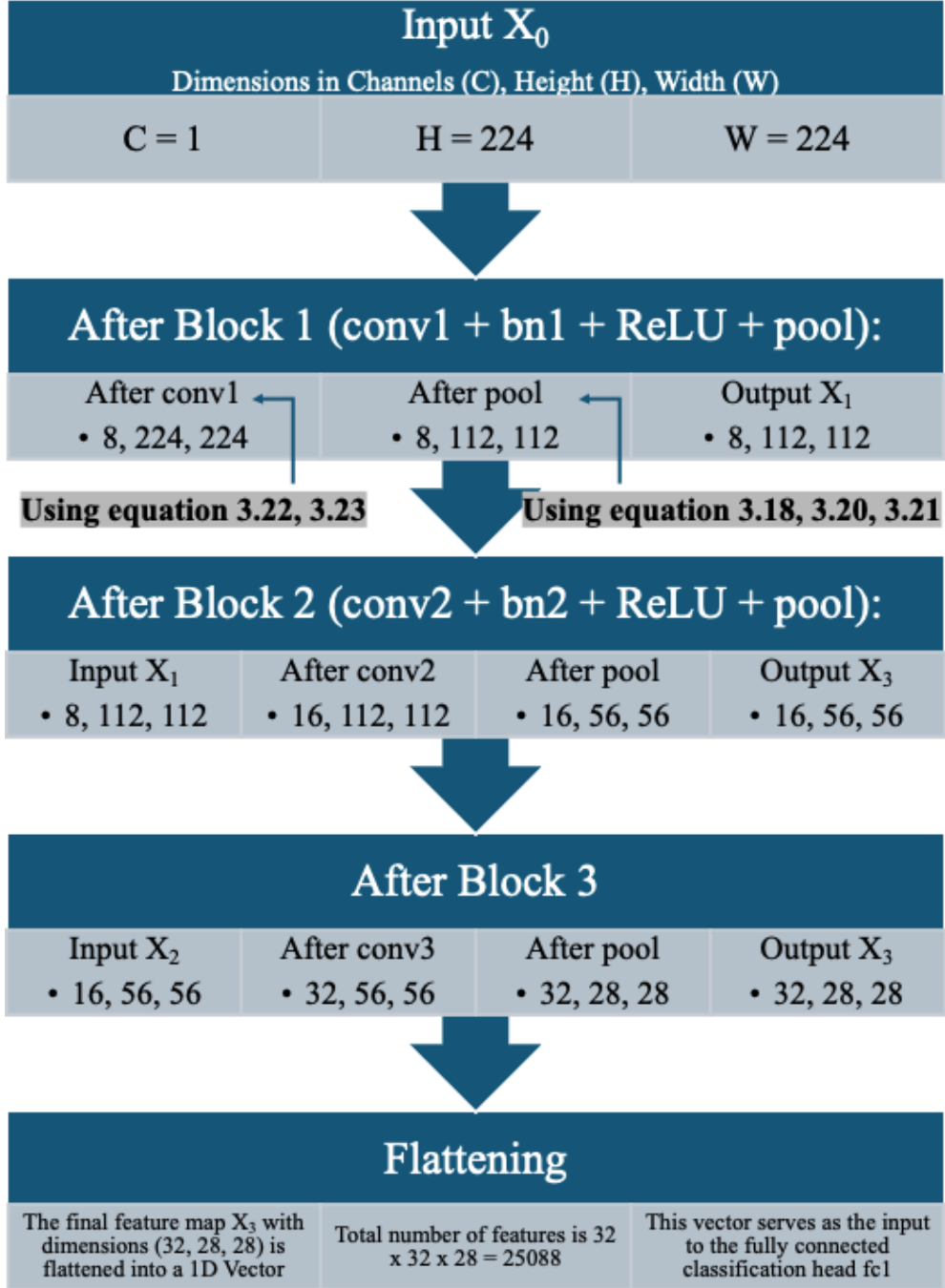


Figure 3.2: Flowchart of CNN implementation (described above)

3.3.4 Hierarchical Feature Extraction and Theoretical Framework

Key Principles of CNNs A CNN is designed structurally to take advantage of key principles to be effective in feature learning from grid-like data, and our eye state classification model is a great example of this. These principles include sparse interactions, parameter sharing, and equivariant representations as per [Goodfellow et al., 2016](#), Ch. 9.2. By using these principles, the CNN gains efficiency both statistically and computationally - allowing it to learn hierarchical features.

Sparse Interactions Sparse interactions are enforced by the convolutional layers by using kernels smaller than the input feature map. A unit in the output map is connected only to a local region (known as the receptive field) in the input map. Our 3x3 kernels in practice ensure that every output neuron in convolutional layers 1 to 3 will only directly process a 3x3 patch from each input channel. This is very much in contrast to the fully connected layers, where every output unit interacts with every input unit. This sparsity significantly reduces the number of parameters and computations ([Goodfellow et al., 2016](#), p. 330).

Parameter Sharing Parameter Sharing: In convolution we employ parameter sharing by applying the same kernel equally across every location of the input map instead of learning separate weights for each input to output connection pair - a single kernel per output channel is learned and then re-used.

In my network, the 8 kernels learned in conv1 are slid across the total input image, heavily reducing the volume of learnable parameters in comparison to say, a locally connected layer without weight sharing (Goodfellow et al., 2016, pp. 331–333). By doing this, we inherently assume that any useful features at one spatial location, such as an edge detector, are probably going to be useful at other locations.

Equivariant Representations Equivariant Representations: Because we have employed parameter sharing (as discussed earlier), convolutional layers can exhibit ‘equivariance’ to translation. This means that if the input is shifted/translated, the output feature map will dynamically shift by the same amount with the features detected remaining the same (Goodfellow et al., 2016, p. 334). As an example, let I be the input, and g be the translation function. We can go on to say that $\text{Conv}(g(I)) = g(\text{Conv}(I))$. This is a very valuable property for tasks where the location of a feature might vary, such as our eye state classification task where the eyes might appear in slightly different locations within the full face image.

Pooling and Invariance Pooling and Invariance: While convolution provides us with equivariance, pooling layers will introduce an approximate invariance to minor translations. Max pooling is illustrated well by (Goodfellow et al., 2016, Fig. 9.8), showing us that if a feature shifts slightly within the pooling window, the maximum activation will often remain the same. This means that the pooled output is unchanged. This invariance is highly beneficial when we’re in a scenario whereby the precision of a feature holds

less significance than the actual presence within a region.

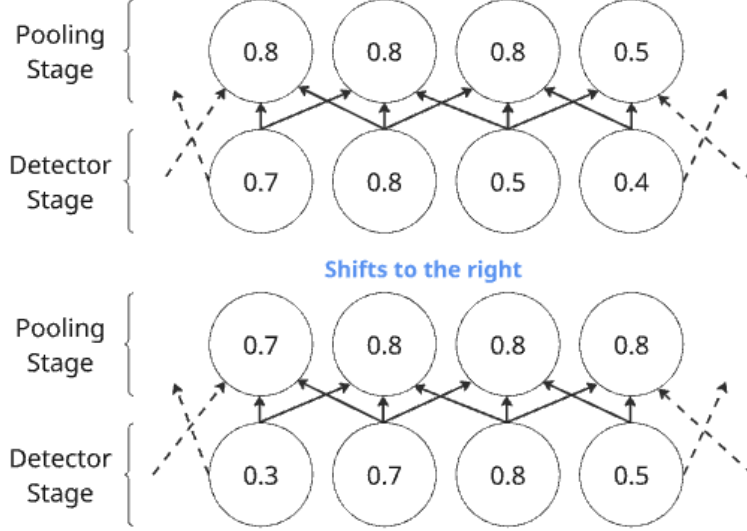


Figure 3.3: Illustration showing max pooling invariance using custom values (adapted from the principle shown in Goodfellow et al., 2016, Fig. 9.8). Top: State before input shift. Bottom: State after input shift to the right. Note how the second and third pooling units keep the value 0.8 despite changes in the underlying detector stage values.

Unified Mathematical Framework Wiatowski & Bölcskei, 2018 captures all of this mathematical theory and encompasses each operation within a unified framework for feature extraction DCNNS. Their work generalises beyond specific filter types such as wavelets, and includes general convolutional transforms (i.e. semi discrete frames) $\Psi_n = \{g_{\lambda_n}\}_{\lambda_n \in A_n}$ and general Lipschitz-continuous non-linearities M_n similar to our ReLU, Lipschitz-continuous pooling operators P_n with pooling factors S_n , which corresponds to our `MaxPool2d` with factor $S_n = 2$. With this, they formulate the n -th network

layer's (U_n) operation that acts on an input f , with a specific filter g_{λ_n} as per [Wiatowski & Bölcskei, 2018](#) Equation 10.

$$U_n[\lambda_n]f := S_n^{d/2} P_n (M_n(f * g_{\lambda_n})) (S_n \cdot) \quad (10)$$

In this $*$ denotes convolution, and the equation captures the sequence: convolution, non-linearity ($M_n = \text{ReLU}$), pooling ($P_n = \text{MaxPool}$), and downsampling ($S_n = 2$). Paths $q = (\lambda_1, \dots, \lambda_n)$ define the propagation through multiple layers, which represents a sequence of filter choices. The overall transformation is $U[q]f = U_n[\lambda_n] \cdots U_1[\lambda_1]f$ ([Wiatowski & Bölcskei, 2018](#), Eq. 13). The complete feature extractor $\Phi_\Omega(f)$ is the union of features extracted at all depths n ([Wiatowski & Bölcskei, 2018](#), Eq. 15).

Translation Invariance Theory Their theory formalises the role of pooling in achieving translation invariance. Theorem 1 (ii) ([Wiatowski & Bölcskei, 2018](#), Eq. 21) provides us with a bound that shows the difference between the features of a translated input $T_t f$ and the original features $\Phi_\Omega^n(f)$ decreases as the product of pooling factors $S_1 \cdots S_n$ increases:

$$||\Phi_\Omega^n(T_t f) - \Phi_\Omega^n(f)|| \leq \frac{2\pi|t|K}{S_1 \cdots S_n} \|f\|_2 \quad (3.24)$$

(adapted from [Wiatowski & Bölcskei, 2018](#), Equation 21)

$$(3.25)$$

This "vertical translation invariance" is a mathematical demonstration of the notion that features will become progressively and increasingly invariant

to translations with increasing network depth, driven directly by the spatial downsampling ($S_n > 1$) which is introduced by pooling. In our network where $S_1 = S_2 = S_3 = 2$, the invariance factor increases by 2 at each block.

Practical Application to Eye State Classification The architecture’s channel expansion (1 to 8, to 16, to 32) aligns perfectly with our goal of hierarchical feature extraction. The early layers such as conv1, with fewer channels, learn generic features such as edges and basic textures. Meanwhile, the deeper layers receive their input from multiple channels from the previous layers. This allows them to combine simpler features to learn more complex and abstract patterns that the eye state classification task requires, such as the specific appearance of open or closed eyes, including eyelids, eyelashes, and surrounding facial features (Goodfellow et al., 2016). The Wiatowski & Bölskei, 2018 framework accounts for this hierarchical structure by allowing different filter sets Ψ_n (potentially of different sizes $|A_n|$).

Real-world Benefits This theoretical foundation manifests itself practically in our scenario and can be clearly observed in how the CNN processes the full face images. The increasing invariance to small translations means the model will still function effectively despite the head moving in a minor way, or the camera position varying as it naturally would do depending on where different people might place the phone camera in the car.

Conclusion In conclusion, our CNN takes advantage of:

- The principles of convolution - for sparse, shared and equivariant feature detection
- Pooling - for dimensionality reduction and invariance

In practice, this design allows our network to transform the raw pixelated input $(1, 224, 224)$, into a compact but high level feature representation $(32, 28, 28)$, appropriate for our classification task. We have demonstrated this is grounded both in established CNN practices ([Goodfellow et al., 2016](#)) and advanced mathematical theory ([Wiatowski & Bölskei, 2018](#)).

3.4 Understanding the Mathematical Foundations of Transfer Learning when Classifying Eye States.

Transfer learning refers to the process of applying insights and patterns learned from one task or dataset to improve performance on a different, yet related, task or dataset (Olivas, Guerrero, Sober, Benedito, & Lopez, 2009) (Murel, n.d.). In our drowsiness detection task, it forms the backbone behind how we classify eyes open and eyes closed state, and gives us the ability to take advantage of proven and pre-existing knowledge to apply to our unique application. In order for us to properly understand how it applies in our project, we reference the formal mathematics established in literature.

3.4.1 Formalising the Domains and Tasks in Transfer Learning

Work completed by Pan & Yang, 2010 gives us a clear structure by defining the core components of transfer learning; domains and tasks.

What is a Domain?

According to (Pan & Yang, 2010), a domain D is conceptualised as consisting of two fundamental components itself. These are a feature space X and a marginal probability distribution $P(X)$.

Starting with a feature space X ; this is what represents any and all possible data instances. If you take for example $\{x_1, \dots, x_n\}$, then each x_i can be seen as a specific data point within that space. For the marginal probability distribution $P(X)$; this is the distribution giving us the likelihood of observing any of those data instances in the feature space.

If we put this in the context of our eye state classification model:

The source domain (D_S) is the ImageNet dataset (this is what the ResNet-18 model was initially trained on) ([MathWorks, n.d.](#)). The feature space (X_S) for ImageNet is made up of over one million general images (from animals and objects to scenes), and is typically represented as multi-dimensional tensors of pixel values after initial processing. The marginal probability distribution reflects the entire distribution of these natural images.

The target distribution is our custom dataset of eye images, managed by the `EyesDataset` class in our code and then processed through the `prepare_data` function.

Hence, the feature space (X_T) is made up of the images of human eyes. Although the images are transformed into tensors similarly to imageNet images, the visual nature of the content (focused on eyes alone) defines this specific feature space.

The marginal probability distribution ($P_T(X)$) is specific to these eye images. Expectedly, this will be different to $P_S(X)$; for example we know that the way colours and edges are shaped in an eye image will be very different to that of the majority of the ImageNet dataset. As ([Pan & Yang, 2010](#)) notes,

”if two domains are different, then they may have different feature spaces or different marginal probability distributions”. In our case, $P(X)$ clearly differs.

Defining a Task

Given a specific domain $D = \{X, P(X)\}$, (Pan & Yang, 2010) defines that a task T is made up of a label space Y and an objective predictive function $f(\cdot)$, denoted as $T = \{Y, f(\cdot)\}$.

The label space is the set of labels for the task, whereas the objective predictive function is not observed directly. This is learned from the training data $\{(x_i, y_i)\}$, where $x_i \in X$ and $y_i \in Y$, has the goal of predicting our label $f(x)$ for a new instance x . If we look at this probabilistically, $f(x)$ can be represented as $P(y|x)$.

Let us bring this back to our project – The source task (T_S) is whatever original task for which the ResNet-18 model was trained and their label space (Y_S) consists of around 1000 distinct object categories (MathWorks, n.d.).

The predictive function ($f_S(\cdot)$) is the original ResNet-18 architecture trained to map an image $x_S \in X_S$ to one of these 1000 labels $y_S \in Y_S$.

The target task (T_T) is therefore our eye state classification, and our label space (Y_T) is the binary pair; ”eyes open” and ”eyes closed”. This is reflected in our `EyesModel` by `num_classes=2` and in our data preparation where the labels are 0 or 1 (e.g., `train_openlbls = [1]`, `train_closedlbls = [0]`).

The predictive function is our `EyesModel` class. The `forward` method takes an input eye image tensor $x_T \in X_T$, and produces outputs that are used to predict the probability $P(y_T|x_T)$ for whether the eyes are open or not.

3.4.2 Unifying the Definition of Transfer Learning

Table 3.2: Transfer Learning Components in Eye State Classification

Component	Source (ImageNet)	Target (Eye Classification)
Domain (D)	$D_S = \text{ImageNet dataset}$	$D_T = \text{Custom eye images dataset}$
Feature Space (X)	$X_S = \text{General images (animals, objects, scenes)}$	$X_T = \text{Human eye images}$
Probability Distribution	$P_S(X) = \text{Broad distribution of natural images}$	$P_T(X) = \text{Narrow distribution specific to eye images}$
Task (T)	$T_S = \text{ImageNet classification}$	$T_T = \text{Eye state classification}$
Label Space (Y)	$Y_S = 1000 \text{ object categories}$	$Y_T = 2 \text{ classes (open, closed)}$
Predictive Function	$f_S(\cdot) = \text{Original ResNet-18}$	$f_T(\cdot) = \text{Modified EyesModel}$

[Pan & Yang, 2010](#) provides a unified definition in "Definition 1 (Transfer Learning)". Given a source domain D_S and learning task T_S , a target domain

D_T and learning task T_T , transfer learning aims to help improve the learning of the target predictive function $f_T(\cdot)$ in D_T using the knowledge in D_S and T_S , where $D_S \neq D_T$, or $T_S \neq T_T$.”

Our project is mathematically dependent on that very definition, because as we have established the marginal probability distributions $P_S(X)$ (general images) and $P_T(X)$ (eye images) are distinct, despite the fact that both domains involve images. Certainly, we know they are not identical.

$T_S \neq T_T$: The label spaces are clearly different (Y_S has 1000 classes, Y_T has 2 classes). This directly implies that $Y_S \neq Y_T$, and hence the conditional probability distributions $P(Y_S|X_S)$ and $P(Y_T|X_T)$ that each predictive function is trying to model are also distinct.

So as we have proven that both conditions ($D_S \neq D_T$ and $T_S \neq T_T$) are met, we can confidently state that the application of a pre-trained ResNet-18 to eye state classification is an explicit instance of transfer learning, as has been formally defined. We are using the ”knowledge” in D_S and T_S (the learned features and weights of the ResNet-18 model from ImageNet) to improve $f_T(\cdot)$ (our model) to classify the different eye states.

3.4.3 Mechanisms of Knowledge Transfer in our Eye Classification Model

The way that the knowledge is transferred from the source to the task can take a variety of forms, and [Sharma, 2024](#) outlines these concepts. The most

relevant to our ResNet-18 approach are Parameter Transfer and Feature Representation.

Parameter Transfer

Parameter Transfer is essentially reusing the model weights. Sharma states that Parameter Transfer "transfers parameters or parts of models (e.g., layers in a neural network) from the source task to the target task" ([Sharma, 2024](#)).

When we use the line `self.backbone = models.resnet18(pretrained=True)`, this is exactly what we actually achieve. In our example, the pre-trained weights are the parameters of the many convolutional layers of ResNet-18, and they are learned on D_S for T_S . These parameters are directly loaded into our model.

In the following step in the implementation, we modify the final classification layer. This is done by:

```
in_features = self.backbone.fc.in_features
self.backbone.fc = nn.Linear(in_features, num_classes)
```

The parameters of the early layers will typically fine tune during training on D_T , which allows them to adjust to the nuances of pictures of an eye while still having the advantage of the general features they learned from ImageNet.

Feature Representation Transfer

To try and simplify feature representation, it is at its core building on the learned hierarchies. Sharma states that this process "focuses on finding a common feature representation for both the source and target domains. The goal is to learn a new feature space where the distributions of the source and target domains are similar" (Sharma, 2024).

My implementation is not explicit in employing MMD to minimise the distribution differences, but by using ResNet-18 in the first place we are inherently taking advantage of this concept. The convolutional layers of ResNet-18 have already learned hierarchies of features from ImageNet with increasing complexity into the deeper layers, where it figures out complex object parts.

If we transfer these layers, we are assuming that this learned feature space $\Phi(X_S)$ is general and useful enough to be used by our target domain D_T . The fine-tuning process then will progressively adapt the feature space to specialise in dividing "eyes open" from "eyes closed". In effect, it is adjusting the feature representation $\Phi(\cdot)$ so that the distributions $P(\Phi(X_S))$ and $P(\Phi(X_T))$ are utilisable for classification by $f_T(\cdot)$.

In conclusion, these mechanisms make up the building blocks that make using a pre-trained model an applicable approach for classifying whether the eyes are open or the eyes are closed.

Chapter 4

Methodology

In this section, the development of the full system is detailed. This includes the datasets, techniques for preprocessing, model architectures and the strategies I undertook for the full implementation.

4.1 Data Collection and Preprocessing

Three different datasets were used for training the models:

- **Driver Drowsiness Dataset (DDD):** 2,000 static images (1,000 per class) derived from the UTA-RLDD dataset, cropped using Viola-Jones algorithm and pre-labeled as drowsy/non-drowsy across 28 subjects.
- **MRL Eye Dataset:** 84,902 infrared eye images captured under various lighting conditions (41,948 closed, 42,954 open) used for the ResNet-18

model.

- **Closed Eyes in the Wild (CEW)**: 27,202 full-face images (14,337 eyes open, 12,865 eyes closed) with each image representing a unique subject.

For facial landmark detection in both the dashboard and the XG-Boost training, I utilised dlib, which was also used in the dashboard for identifying the face. I implemented person based splitting for the XGBoost model in order to make sure the system had the ability to generalise as opposed to learning person-specific characteristics. I used an 80% to 20% for the MRL dataset as this was recommended by the creators. For both the XGBoost (and stacked) as well as the CNN, I used a 70/15/15 split.

4.2 Feature Engineering and Model Architectures

4.2.1 XGBoost Pre-Drive Check

Our XGBoost model (and stacked) extracts eight features using the detected facial landmarks:

- Left and right eye aspect ratios (EAR)
- Average eye aspect ratio

- Difference between left and right eye ratios
- Mouth aspect ratio (MAR)
- Eye-mouth ratio (EMR)
- Normalised head pitch angle
- Normalised head yaw angle

I configured the model with 1,000 estimators with a maximum depth of 5 and a learning rate of 0.1. I implemented a stacked ensemble on top of the original XGBoost Model, with the ensemble being made up of XGBoost, LightGBM, and RandomForestClassifier. In that ensemble, I used a Logistic Regression meta-learner using 5-fold cross-validation in order to prevent data leakage.

4.2.2 Full-Face CNN Eye-State Classifier

The CNN architecture was developed to have three convolutional blocks and three fully connected layers, which accept the input dimensions of (1, 224, 224). I used CrossEntropyLoss during training with a label smoothing of 0.1. Adam optimiser was used with an initial learning rate 0.001, and a ReduceLROnPlateau scheduler. I used batch size 32 and implemented an early stopping mechanism.

I made a late switch to this approach (classifying the eyes open and eyes closed on a full face) after I found an earlier CNN I had been developing throughout this year, thinking that it had been properly implementing person-based

splitting, had a data leakage problem, progressively mixing the person-id up during its person-based splitting phase on every run, giving false improvements on accuracy that I had attributed to hyperparameter tuning. This explains its issues in the implementation on the dashboard.

4.2.3 ResNet-18 Cropped-Eye Classifier

For the cropped-eye-state classification, I used a pretrained ResNet-18 model. I replaced the final fully connected layer with a new linear layer (512 input features, 2 output classes). Using the adam optimiser, all the layers were fine-tuned (none were frozen), and I used a ReduceLROnPlateau scheduler with factor 0.5, patience 2. I set a maximum epoch of 15 due to computational constraints, as developing it using a MacBook does not offer GPU.

4.3 Live Drowsiness Detection System

4.3.1 System Architecture

I built the dashboard using PyQt5, with the modules communicating through my main class DrowsinessDetectionDashboard. Video processing occurs in a separate thread to maintain the UI's responsiveness. For face and eye detection, we use OpenCV's Haar cascades, followed by inferring the ResNet-18 Model for the classification element, which is done in real time. The main class does contain the functionality to periodically invoke the old, overfitted

CNN model for drowsiness probability estimation and the XGBoost model for pre-drive checks, so if we continue development and fix this, this is an easy switch to make.

4.3.2 Real-time Monitoring Features

- **Microsleep Detection:** Triggered when the ResNet-18 model detects 15 consecutive frames with closed eyes (approximately 0.5-1 second), generating visual and audio alerts.

This is triggered when the ResNet-18 model detects 15 consecutive frames with closed eyes (approximately 0.5-1 second). This produces both visual and audio alerts. I decided to switch to Audio files recorded by myself because text to speech was computationally demanding and buggy via an API. In practice and testing, I found it has an even better desired effect of startling the driver/subject.

- **Blink Rate Monitoring:** The system tracks rapid open-closed-open eye transitions (blink duration $\leq 0.5s$) and calculates blinks per minute over a continuous time window.

To calculate the blink rate, we track only open-closed-open eye transitions done in a rapid manner, and calculate the blinks per minute by dividing these sequences over a continuous time window.

- **Pre-Drive Safety Check:** This is manually triggered via the GUI, this feature captures a frame, extracts facial landmarks using Dlib, calculates the eight engineered features, and feeds them to the XGBoost model to determine if the driver is safe to proceed.

To run the Pre-Drive Safety Check, you must trigger this manually via the GUI. This feature captures a frame, and extracts the facial landmarks using Dlib, feeding the features specified earlier into the model and playing one of two pre-recorded messages conditionally.

4.3.3 Fallback Mechanisms

For camera dropouts and UI crashes, we do have heavy error handling, however false positives simply require the user to hear out the pre-recorded message as there is no interrupt implemented yet.

Chapter 5

Results

5.1 XGBoost Model

In this section, we look at the development and performance of the XGBoost model used for drowsiness detection. We look at it both as a standalone model, and then when it is stacked into an ensemble.

5.1.1 Without Stacking

Architecture and Features

In this model, we utilise facial landmarks to extract the key drowsiness indicators. These include the following: left and right eye aspect ratios (measuring eye openness), average eye aspect ratio, difference between left

and right eye ratios, mouth aspect ratio (measuring mouth openness), eye-mouth ratio, and normalised head pitch and yaw angles.

Dataset and Training

This model was trained to use person-based splitting, helping us prevent it from learning individual facial characteristics as opposed to drowsiness indicators. Our test set contained 257 samples. Some of the key hyperparameters included 1000 estimators, maximum depth of 5, and learning rate of 0.1.

Performance Metrics

The standalone XGBoost model achieved an accuracy of 77% with class-specific metrics as follows: Non-Drowsy Class (0) achieved precision 0.86, recall 0.75, and F1-score 0.81; Drowsy Class (1) achieved precision 0.67, recall 0.81, and F1-score 0.73. The model's ROC-AUC was 0.84.

Confusion Matrix

The confusion matrix for the standalone XGBoost model (Focused Model) is presented in Figure 5.1.

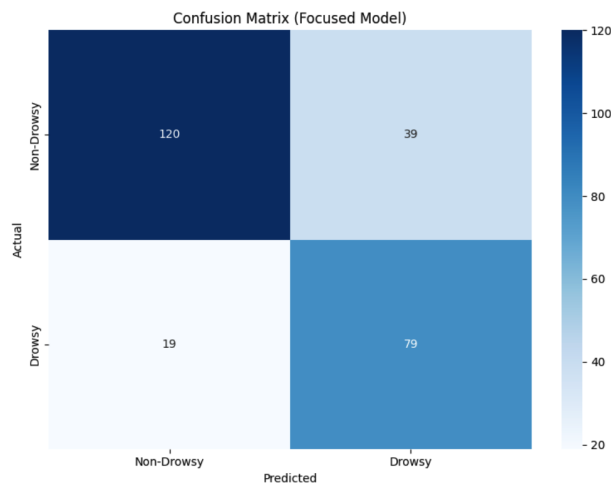


Figure 5.1: Confusion Matrix for the Standalone XGBoost Model (Focused Model).

5.1.2 With Stacking

Ensemble Architecture

The stacking ensemble combined three base learners (XGBoost, LightGBM, RandomForestClassifier) with a Logistic Regression meta-learner using 5-fold cross-validation.

Performance Comparison

The stacked model only showed minimal performance changes compared to the standalone XGBoost model with an accuracy of 77.4% (this is unchanged from standalone XGBoost). For the Non-Drowsy Class, the model achieved

precision 0.85, recall 0.77, and F1-score 0.81; for the Drowsy Class, it achieved precision 0.68, recall 0.79, and F1-score 0.73. The ROC-AUC was 0.843, representing a slight improvement of 0.51% over standalone XGBoost.

Confusion Matrix

The stacked model’s confusion matrix showed 123 true negatives, 36 false positives, 22 false negatives, and 76 true positives, as depicted in Figure 5.2.

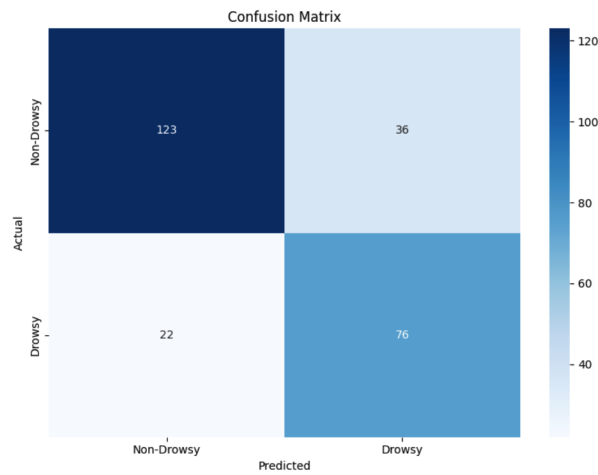


Figure 5.2: Confusion Matrix for the Stacked Ensemble Model.

5.1.3 Feature Importance

The most important features (normalised) for the models were mouth aspect ratio (~ 0.16), head pitch (~ 0.15), head yaw (~ 0.15), right eye aspect ratio (~ 0.13), and eye-mouth ratio (~ 0.12). Figure 5.3 illustrates the normalised feature importances across different models.

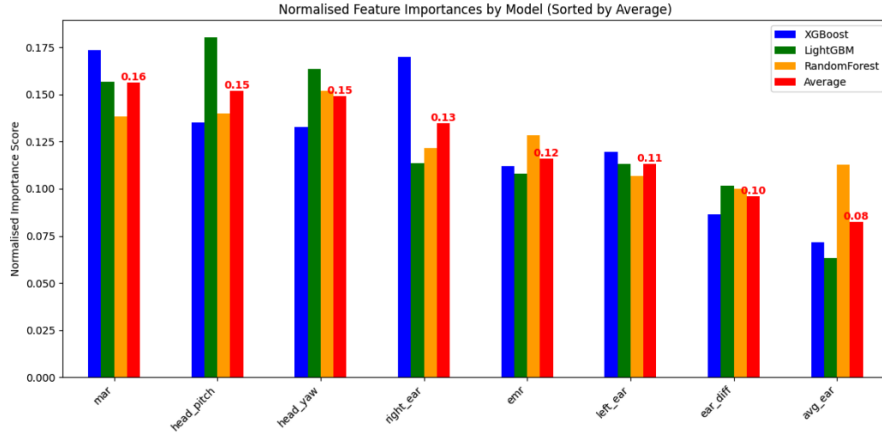


Figure 5.3: Normalised Feature Importances by Model (Sorted by Average).

By performing a SHAP analysis on the just the XGBoost element of the stacked ensemble, we can infer that the eye-mouth ratio, head yaw, and mouth aspect ratio had the most significant impact on predictions, as shown in Figure 5.4.

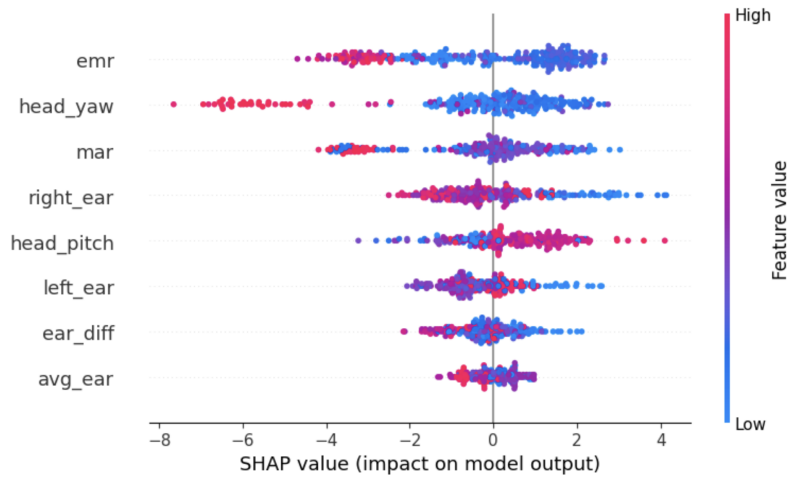


Figure 5.4: SHAP Value Summary Plot for the XGBoost Component, Showing Impact of Features on Model Output.

5.1.4 Notable Observations

The calibration analysis showed that the XGBoost model (standalone) outperformed the stacked ensemble slightly when it came to calibration, specifically at the higher end of the predicted probabilities. Additionally, the performance demonstrated significant variety across the different individuals, which suggests that person specific factors are driver behind variation in the model’s accuracy.

5.2 Custom CNN for Drowsiness Detection: Results

A custom Convolutional Neural Network (CNN), designated DrowsinessDetectionCNN, was developed and evaluated for the task of classifying facial images into 'Drowsy' (eyes closed) and 'Non-Drowsy' (eyes open) states.

5.2.1 Dataset and Preprocessing

The model was trained on grayscale images of size 224×224 pixels. The dataset was partitioned as follows:

- **Training set:** 19,040 images (Drowsy: 9,005, Non-Drowsy: 10,035)
- **Validation set:** 4,080 images (Drowsy: 1,930, Non-Drowsy: 2,150)
- **Test set:** 4,080 images (Drowsy: 1,929, Non-Drowsy: 2,151)

The training data underwent the following transformations:

- Resize to (224, 224)
- Convert to Grayscale
- HorizontalFlip (probability = 0.5)
- RandomBrightnessContrast (brightness_limit=0.15, contrast_limit=0.15, probability=0.5)
- GaussianBlur (blur_limit=(1.0, 3.0), probability=0.125)
- GaussNoise (variance_limit=(0.1, 0.3), probability=0.125)
- ToTensorV2 conversion

5.2.2 Model Architecture

The DrowsinessDetectionCNN architecture is defined as:

- **Input:** Grayscale images (1 channel, 224×224)
- **Convolutional Layers:**
 - Conv1: `nn.Conv2d(1, 8, kernel_size=3, padding=1)` → BN1
→ ReLU → MaxPool(2,2)
 - Conv2: `nn.Conv2d(8, 16, kernel_size=3, padding=1)` → BN2
→ ReLU → MaxPool(2,2)

- Conv3: `nn.Conv2d(16, 32, kernel_size=3, padding=1) → BN3`
`→ ReLU → MaxPool(2,2)`
- Resulting feature map size after pooling: $32 \times 28 \times 28$

- **Fully Connected Layers:**

- FC1: `nn.Linear(32 * 28 * 28, 64) → BN4 → ReLU → Dropout(0.4)`
- FC2: `nn.Linear(64, 32) → BN5 → ReLU → Dropout(0.4)`
- FC3: `nn.Linear(32, 2)` (Output layer for 2 classes)

- **Weight Initialisation:** Kaiming Normal for Conv2D/Linear layers, constant 0 for biases, and BatchNorm weights constant 1.

5.2.3 Training Details

The model was trained on a CPU with the following parameters:

- **Loss Function:** `nn.CrossEntropyLoss` with label smoothing of 0.1.
- **Optimiser:** `torch.optim.Adam` with a learning rate of 0.001.
- **Scheduler:** `ReduceLROnPlateau` (mode="min", factor=0.1, patience=5).
- **Epochs:** 100, with early stopping (patience of 40 epochs).

Performance Evaluation

On the independent test set, the model achieved the following performance metrics:

- **Accuracy:** 0.9831
- **Precision:** 0.9753
- **Recall:** 0.9930
- **F1 Score:** 0.9841

The confusion matrix can be found in Figure 5.5, quantifying the model's performance.

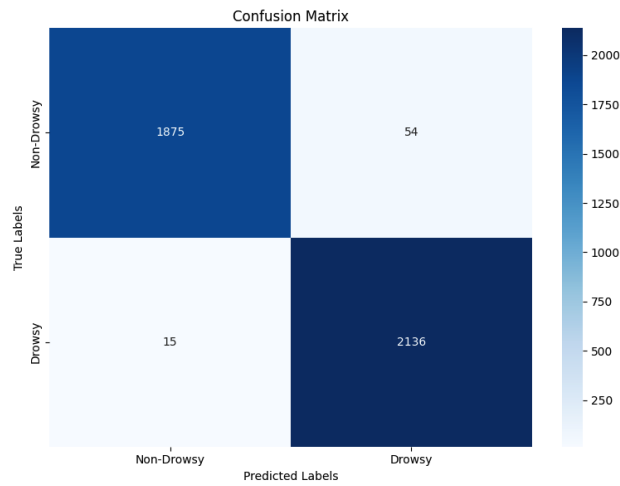


Figure 5.5: Confusion matrix of the DrowsinessDetectionCNN on the test set.

The training and validation loss and accuracy curves over epochs are shown in Figure 5.6.

Model Interpretability

In order for us to get an idea of the model's decision making process, I utilised Grad-CAM (Gradient-weighted Class Activation Mapping). Figure 5.7 shows

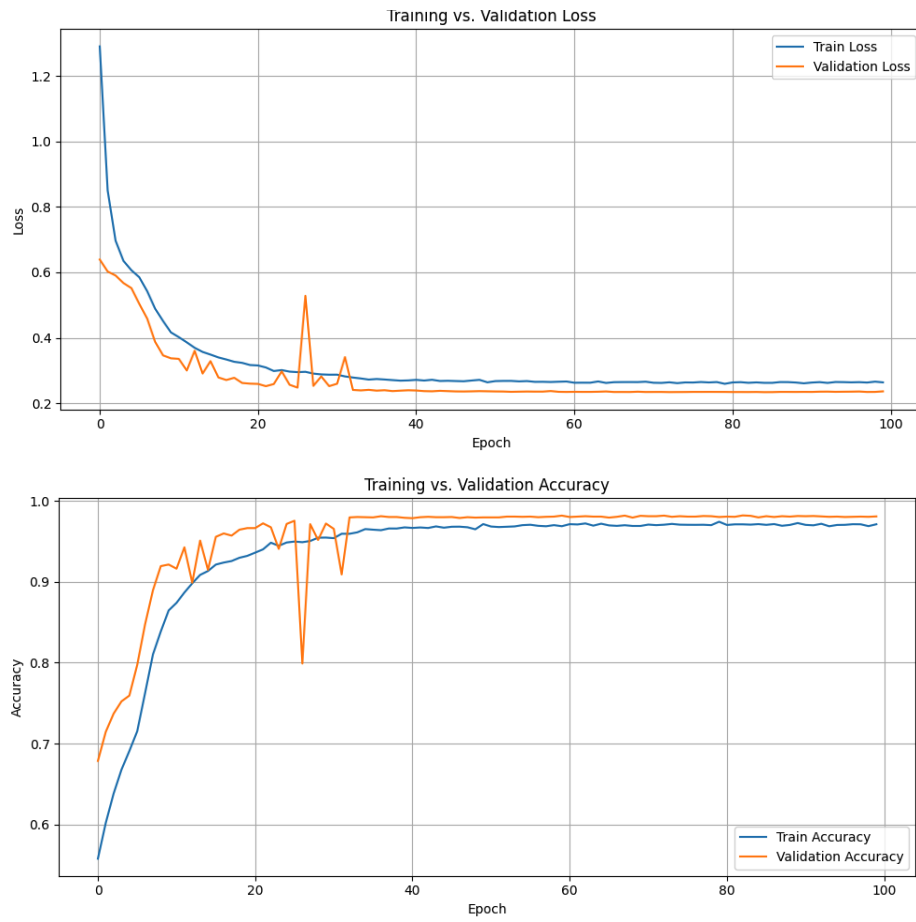


Figure 5.6: Training vs. Validation Loss (top) and Accuracy (bottom) throughout the training process.

an example Grad-CAM overlay, visualising the importance of different features in the image.

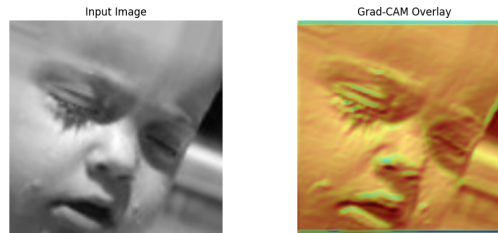


Figure 5.7: Example Grad-CAM overlay on a test image, indicating regions influencing the CNN’s prediction.

5.3 ResNet Model (Eyes Open vs. Closed)

In this section, we discuss the results from the ResNet-based model, which classifies eye states as either open or closed. This model takes advantage of transfer learning from a pre-trained ResNet-18 architecture.

5.3.1 Architecture details (ResNet variant, input size):

Variant: ResNet-18. The pre-trained ResNet-18 model from `torchvision.models` was used as the backbone. Input Size: Images were resized to 224x224 pixels.

5.3.2 Preprocessing (normalisation, augmentation):

Normalisation: Images were normalised using the standard ImageNet mean [0.485, 0.456, 0.406] and standard deviation [0.229, 0.224, 0.225]. Training

Augmentation: Random Horizontal Flip, Random Rotation (10 degrees), and Colour Jitter (brightness=0.2, contrast=0.2). Validation/Test Preprocessing:

Only resizing and normalisation were applied. All images were converted to RGB format and then to PyTorch tensors.

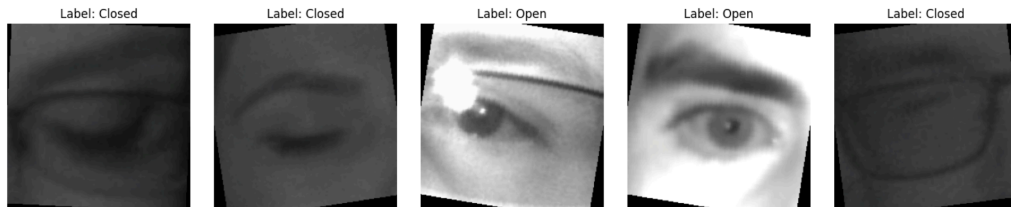


Figure 5.8: Sample Training Images with Labels, Illustrating Data Augmentation Effects.

5.3.3 Dataset Sizes

Training set size: 65,340 images. Validation set size: 16,335 images. Test set size: 3,223 images.

5.3.4 Hyperparameters:

Optimiser: Adam. Loss Function: Cross-Entropy Loss. Initial Learning Rate: 0.001. Learning rate schedule: ReduceLROnPlateau; learning rate was reduced by a factor of 0.5 if validation loss did not improve for 2 epochs. Batch size: 32. Number of epochs: 15 (with early stopping). Early Stopping Patience: 5 epochs. Note that the training completed all 15 epochs, as LR was reduced and validation loss improved again before patience ran out. Training time: 11127.53 seconds (185.46 minutes) for the logged run.

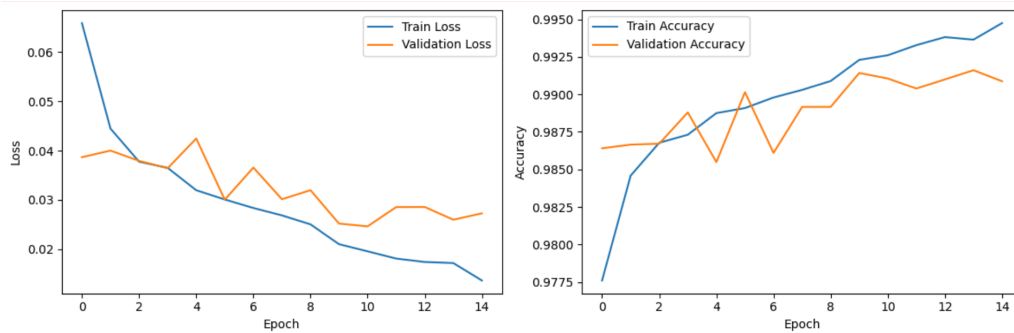


Figure 5.9: Training and Validation Accuracy and Loss per Epoch for the ResNet Model.

5.3.5 Performance metrics:

Final Test Accuracy: 0.9054 (or 90.54%). Final Training Accuracy (Epoch 15): 0.9948. Final Validation Accuracy (Epoch 15): 0.9909.

Precision, Recall, F1-score (open vs. closed):

From the confusion matrix (Closed=0, Open=1): True Positives (Open correctly predicted): 1355; False Positives (Closed predicted as Open): 3; True Negatives (Closed correctly predicted): 1563; False Negatives (Open predicted as Closed): 302.

For "Open" (Class 1): Precision = 0.9978, Recall = 0.8177, F1-score = 0.8989.

For "Closed" (Class 0): Precision = 0.8381, Recall = 0.9981, F1-score = 0.9112. ROC-AUC: 1.00 (as indicated visually on the ROC curve plot).

5.3.6 Confusion matrix:

The confusion matrix, presented in Figure 5.10, summarises the classification performance on the test set.

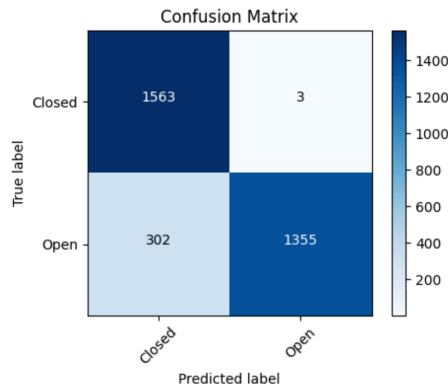


Figure 5.10: Confusion Matrix for the ResNet Model on the Test Set.

The Receiver Operating Characteristic (ROC) curve, shown in Figure 5.11,

further brings to light how the model can differentiate between the 'open' and 'closed' eye states across various threshold settings.

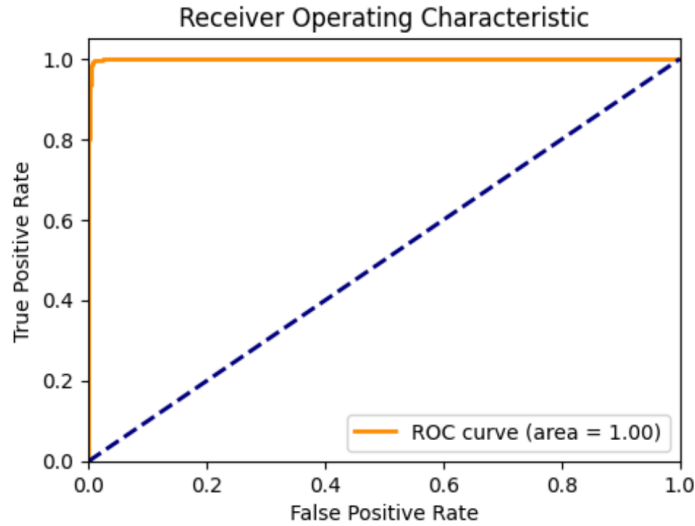


Figure 5.11: Receiver Operating Characteristic (ROC) Curve for the ResNet Model on the Test Set.

5.3.7 Notable observations:

The model achieved very high training and validation accuracy (above 99%). High training and validation accuracy is observed, however the lower test accuracy at ~90.54% indicates some drop in generalisation to unseen test data. When reviewing the confusion matrix (Figure 5.10), we can see that the model has a higher likelihood of misclassifying an "open" eye as "closed". This indicates we may have some bias towards closed, or maybe the subtly open eyes cause the model some difficulty.

The Grad-CAM visualisation of an open eye (illustrated in Figure 5.12) being

predicted as closed shows the model is focusing in the right area, but there may be a lack of strong activation where it expects there to be.

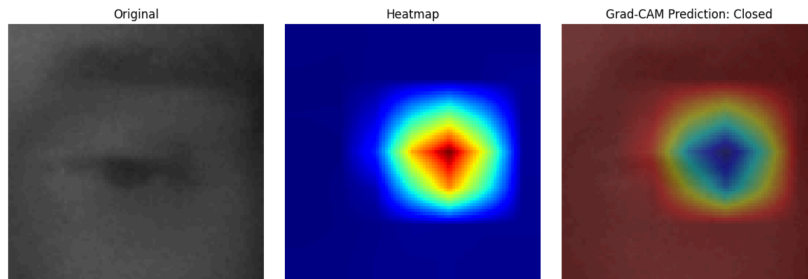


Figure 5.12: Grad-CAM Visualisation of an Incorrect Prediction by the ResNet Model.

The learning rate was reduced during training (e.g., Epoch 9/15 in one log: Learning Rate: 0.000500 (reduced)), indicating the scheduler was active. The final run (15 epochs) also had LR reductions.

5.4 Real-World Deployment Results

Here we discuss the results, feedback and insights we got from deploying and evaluating the drowsiness detection dashboard in a real-world driving scenario.

5.4.1 On-Road Test Setup

The dashboard was evaluated in a real-world scenario under the following controlled conditions to ensure safety and consistency: **Subject:** Jay Lowe

(brother of Toby Lowe); **Vehicle:** Audi A3; **Conditions:** Daytime, low-speed driving on safe road conditions; **Duration:** Approximately 2 minutes of continuous driving for each test iteration focusing on specific functionalities.

5.4.2 **Functionality Performance**

The performance of key features was assessed both with a standard webcam in a static environment and during the on-road tests. Table 5.1 summarises these findings.

Table 5.1: Functionality Performance Summary: Webcam
vs. On-Road

Feature	Home	Road	Notes
Real-time eye status (ResNet-18)	✓	✓	Accurately tracked eye open/closed states and blink rate without drop-outs. Ran at full frame rate, with no missed frames or misclassifications.
Pre-drive drowsiness check (XG-Boost)	✓	✓	Consistently classified alert vs. drowsy states across varying levels of fatigue. Produced stable confidence scores before departure.

(continued on next page)

(Table 5.1 continued)

Feature	Home	Road	Notes
Rolling drowsi- ness check (CNN)	X	X	Generalisation issues caused poor performance. On-road vibration and lighting caused constant "drowsy" outputs—nearly 100% false-positive under motion.
Audio alerts	✓	✓	All bugs resolved; alerts were clear and delivered with the intended urgency. Playback consistently sounded within 0.2 s of the detection event.
Microsleep de- tection	✓	✓	Reliable for lighter skin tones; degraded on darker skin. Correctly flagged eye closures [0.5 s for light skin, but missed equivalent events on dark skin—suggests more diverse training data.
UI responsive- ness	✓	✓	Occasional, non-blocking crashes; did not interrupt the 2 min test. Averaged one freeze per test (downtime $\approx$ 3 s), self-recovered without restart.

5.4.3 Qualitative Feedback & Observations

In addition to quantitative performance metrics, qualitative feedback was gathered from the test subject, and several observational notes were made during the deployment.

Subject's Comments

The test subject provided the following feedback: "The alert sound was jarring but exactly what I'd want if I were falling asleep." "The interface felt intuitive, but would be more practical if embedded directly into a smartphone camera rather than run on a laptop."

Observed Issues

During the on-road testing, the following issues were observed: **False positives:** Only the rolling-CNN check produced spurious "drowsy" readings on-road; **System glitches:** Screen freezes were noted during the pre-drive routine, though they did not affect the core functionality during the main driving test.

5.4.4 Prototype Demonstration

Video demonstrations of the prototype being tested in various scenarios are available online: Microsleep and Pre-Drive Test (Static and In-Car): <https://>

youtube.com/shorts/1xpchZcw7V8?feature=share; Microsleep Test (On-Road Driving): <https://youtube.com/shorts/kuYC7JP8y-s?feature=share>; Full Feature Test (In Office): <https://youtu.be/B7uF1UFbE1Y>.

Figure 5.13 shows a snapshot of the driver drowsiness monitoring system’s dashboard interface as used during testing.

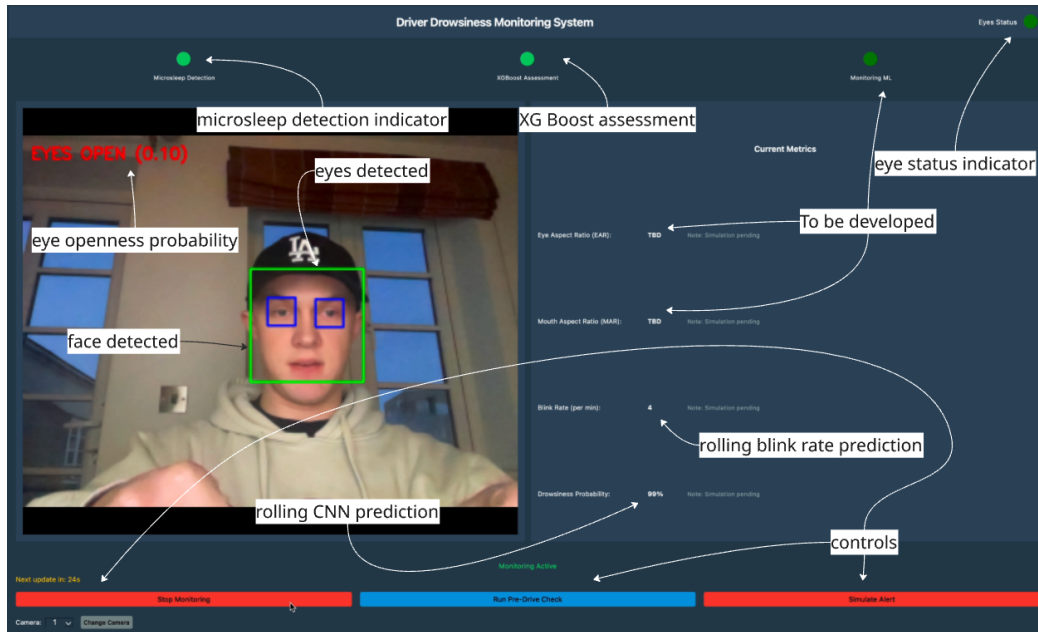


Figure 5.13: Driver Drowsiness Monitoring System Dashboard Interface

5.4.5 Summary of Real-World Deployment

Overall, the real-world deployment of the dashboard was a success, demonstrating consistently high performance for the real-time eye state classification (ResNet-18), pre-drive drowsiness testing (XGBoost), and audio alerting components under realistic driving conditions. Immediate areas identified for

improvement include the generalisation capabilities of the rolling drowsiness CNN inference, which exhibited a high rate of false positives on-road. Furthermore, the generalisability of the microsleep detection component across diverse skin tones needs to be addressed, likely through the incorporation of more skin-tone-inclusive training data. While minor UI glitches were observed, they did not critically impact the core functionality during the brief test periods.

Chapter 6

Discussion

Through this project, we have successfully engineered and tested a comprehensive system that can combat driver drowsiness in practice. It integrates both traditional machine learning and deep machine learning, showcasing strength in both domains. Our traditional XGBoost model which leverages custom-engineered geometric features, demonstrated strong performance for such a small dataset and difficult task (general drowsiness detection). Perhaps its biggest delivery was the understanding that it gave us through the feature importance analysis. The person-based data splitting was crucial in ensuring that the performance we received was accurate.

The deep learning approaches delivered higher accuracy for eye-state classification. While operating on full-face images, it achieved an impressively high $\sim 98\%$ test accuracy, which highlights the potential for robust detection. While it achieved a slightly lower (but still solid) $\sim 90\%$ test accuracy. Note: In the abstract on page 2 we targeted 90% accuracy for ResNet, while our

actual results on page 81 and 83 achieved about 90.54%. I use the more general 90% from abstract for consistency with the text above. The ResNet-18 model proved its effectiveness as a lighter weight choice that was perfect for its integration into the dashboard.

The core functionalities of the dashboard were proven to be successfully built when the microsleep detection, audio alerts, pre-drive assessment and blink rate were successfully tested in real-time. There were of course limitations; the initially developed rolling CNN exhibited poor generalisation on the road, and was later found to be suffering from data-leakage in the late stages in the project. This perhaps could have been avoided with more rigorous evaluation, and was a certain learning point. Furthermore, the microsleep detection with ResNet-18 showed that it struggled on darker skin tones, suggesting potential for more diverse training data. While the stacking gave us knowledge on feature importance, the marginal improvement it gave us in its objective (to classify drowsiness) indicates to me that the complexity was not warranted.

We can be reassured that these discoveries align with the research and literature. Deep learning illustrated itself as the strongest choice for classification, however it also underlined that real-world generalisation and diverse datasets are pitfalls you must be careful to avoid. Our combination of using deep learning for tasks requiring guaranteed performance, and an interpretable model which we must be able to explain for assessments, give us the best of both worlds, and means we take a balanced, practical approach.

Future work should put priority on improving the abilities of this system and others for consistency across skin-tone and combating bias. This could

potentially be done by retraining the CNN to improve its ability to adapt to dynamic conditions, or take a step further and explore spatio-temporal models (e.g., 3D CNNs) as discussed in the literature. Optimising our dashboard further, as suggested by the subject during on-road testing, is another equally opportunity-lain route.

Chapter 7

Conclusion

In this dissertation, we successfully address the challenges we set out to conquer. In all, we have developed, implemented, and evaluated a multi-model machine learning-based drowsiness detection system.

The primary objectives were met, including:

- The development of a geometric feature-based XGBoost model achieving $\sim 77\%$ accuracy for drowsiness detection.
- The design of a custom CNN achieving $\sim 98\%$ accuracy for full-face eye-state classification.
- The fine-tuning of a ResNet-18 model reaching $\sim 90\%$ accuracy for cropped-eye classification, suitable for real-time application.
- The integration of these models into a functional real-time Python dashboard featuring blink-rate metrics, drowsiness scoring, and audio

alerts.

- The successful deployment and evaluation of the system using a live camera in on-road conditions, demonstrating practical viability for key features.

The solution we have is strong and adaptable, with room for growth in future work. The potential is already there for the system to have a tangible impact in saving lives, and it is my hope that I can continue to build on this.

References

- Albadawi, Y., Takruri, M., & Awad, M. (2022). A review of recent developments in driver drowsiness detection systems. *Sensors*, 22(5), 2069. doi: 10.3390/s22052069
- Arif, S., Munawar, S., & Ali, H. (2023). Driving drowsiness detection using spectral signatures of eeg-based neurophysiology. *Frontiers in Physiology*, 14, 1153268. doi: 10.3389/fphys.2023.1153268
- Bamidele, A. A., Kamardin, K., Abd Aziz, N. S. N., Sam, S. M., Ahmed, I. S., Azizan, A., ... Kaidi, H. M. (2019). Non-intrusive driver drowsiness detection based on face and eye tracking. *International Journal of Advanced Computer Science and Applications*, 10(7). doi: 10.14569/IJACSA.2019.0100775
- Chai, M., Li, S.-W., Sun, W.-C., Guo, M.-Z., & Huang, M.-Y. (2019). Drowsiness monitoring based on steering wheel status. *Transportation Research Part D: Transport and Environment*, 66, 95–103. Retrieved from <https://www.sciencedirect.com/science/article/pii/S1361920917306582> doi: 10.1016/j.trd.2018.07.007
- Chiew, L. X., & Hong, J. (2025, February 19). *iphone 16e kicks off apple's ai drive in 2025*. Canalys Insights. Retrieved from <https://canalys.com/>

REFERENCES

- [insights/apple-iphone16e-ai-2025](#)
- Cui, J., Lan, Z., Liu, Y., Li, R., Li, F., Sourina, O., & Müller-Wittig, W. (2022, June). A compact and interpretable convolutional neural network for cross-subject driver drowsiness detection from single-channel eeg. *Methods*, 202, 173–184. Retrieved from <http://dx.doi.org/10.1016/j.ymeth.2021.04.017> doi: 10.1016/j.ymeth.2021.04.017
- Department for Transport. (2011). *Fatigue and road safety: A critical analysis of recent evidence*. Retrieved from <https://webarchive.nationalarchives.gov.uk/ukgwa/20121103213009/http://www.dft.gov.uk/publications/rsrr-theme3-fatigue-road-safety-analysis/>
- Department for Transport. (2023). *Reported road casualties great britain, annual report: 2022*. Retrieved from <https://www.gov.uk/government/statistics/reported-road-casualties-great-britain-annual-report-2022/reported-road-casualties-great-britain-annual-report-2022#:~:text=In%202022%2C%20there%20were%20328,of%203%25%20compared%20to%202019> (Source of Crash Statistics Banding for Table 1.1)
- Dewi, C., Chen, R.-C., Chang, C.-W., Wu, S.-H., Jiang, X., & Yu, H. (2022). Eye aspect ratio for real-time drowsiness detection to improve driver safety. *Electronics*, 11(19), 3183. doi: 10.3390/electronics11193183
- Fujiwara, K., Iwamoto, H., Hori, K., & Kano, M. (2024). Driver drowsiness detection using r-r interval of electrocardiogram and self-attention autoencoder. *IEEE Transactions on Intelligent Vehicles*, 9(1), 2956–2965. doi: 10.1109/TIV.2023.3308575

REFERENCES

- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.
(<http://www.deeplearningbook.org>)
- Junaedi, S., & Akbar, H. (2018). Driver drowsiness detection using iris-based perclos estimation. In *Journal of physics: Conference series* (Vol. 1090, p. 012037). doi: 10.1088/1742-6596/1090/1/012037
- Kielty, P., Dilmaghani, M. S., Ryan, C., Lemley, J., & Corcoran, P. (2023). Neuromorphic sensing for yawn detection in driver drowsiness. In *Fifteenth international conference on machine vision (icmv 2022)* (Vol. 12701, p. 127010Z). doi: 10.1117/12.2680327
- Knaak, C., von Eßen, J., Kröger, M., Schulze, F., Abels, P., & Gillner, A. (2021). A spatio-temporal ensemble deep learning architecture for real-time defect detection during laser welding on low power embedded computing boards. *Sensors*, 21(12), 4205. Retrieved from <https://doi.org/10.3390/s21124205> doi: 10.3390/s21124205
- Lipton, Z. C., Elkan, C., & Naryanaswamy, B. (2014). Optimal thresholding of classifiers to maximize f1 measure. In T. Calders, F. Esposito, E. Hüllermeier, & R. Meo (Eds.), *Machine learning and knowledge discovery in databases* (Vol. 8725, pp. 225–239). Berlin, Heidelberg: Springer. Retrieved from https://doi.org/10.1007/978-3-662-44851-9_15 doi: 10.1007/978-3-662-44851-9_15
- MathWorks. (n.d.). *resnet18 (deep learning toolbox)*. Retrieved from <https://uk.mathworks.com/help/deeplearning/ref/resnet18.html>
- Murel, J. (n.d.). *What is transfer learning?* Retrieved from <https://www.ibm.com/think/topics/transfer-learning>
- Nasri, I. (2022). *Driver drowsiness dataset (ddd)*. Kaggle. Retrieved

REFERENCES

- from <https://www.kaggle.com/datasets/ismailnasri20/driver-drowsiness-dataset-ddd> (Accessed 2025)
- Olivas, E. S., Guerrero, J. D. M., Sober, M. M., Benedito, J. R. M., & Lopez, A. J. S. (2009). *Handbook of research on machine learning applications and trends: Algorithms, methods, and techniques*. Hershey, PA: Information Science Reference.
- Pan, S. J., & Yang, Q. (2010). A survey on transfer learning. *IEEE Transactions on Knowledge and Data Engineering*, 22(10), 1345–1359. doi: 10.1109/TKDE.2009.191
- RAC. (n.d). *Hgv driving hours explained*. Retrieved from <https://www.rac.co.uk/business/news-advice/fleet-management/hgv-driving-hours-explained> (Accessed: 2025-04-17)
- Ramzan, M., Khan, H., Awan, S., Ismail, A., Ilyas, M., & Mahmood, A. (2019). A survey on state-of-the-art drowsiness detection techniques. *IEEE Access*. doi: 10.1109/ACCESS.2019.2914373
- Road Safety Scotland. (n.d.). *Driver fatigue campaign*. Retrieved from <https://roadsafety.scot/road-user-advice/driver-fatigue/fatigue-campaign/>
- Sayed, E. A. (2024). *Driver drowsiness detection (cnn — mobilenetv2)*. Kaggle notebook. Retrieved from <https://www.kaggle.com/code/esraameslamsayed/driver-drowsiness-detection-cnn-mobilenetv2/notebook> (Accessed 2025)
- scikit-learn. (n.d.). *StandardScaler documentation*. <https://scikit-learn.org/stable/modules/generated/sklearn.preprocessing.StandardScaler.html>. (Accessed: April 2025)

REFERENCES

- Sharma, R. (2024). *Transfer learning and the mathematics behind it*. Retrieved from <https://medium.com/@rajat01221/transfer-learning-and-the-mathematics-behind-it-0988153178aa>
- Souto Maior, C. B., Moura, M. J. d. C., Santana, J. M. M., & Lins, I. D. (2020). Real-time classification for autonomous drowsiness detection using eye aspect ratio. *Expert Systems with Applications*, 158, 113505. doi: 10.1016/j.eswa.2020.113505
- Stone, M. (1974). Cross-validatory choice and assessment of statistical predictions. *Journal of the Royal Statistical Society: Series B (Methodological)*, 36(2), 111–133. Retrieved from <https://doi.org/10.1111/j.2517-6161.1974.tb00994.x> doi: 10.1111/j.2517-6161.1974.tb00994.x
- Ting, P.-H., Hwang, J.-R., Doong, J.-L., & Jeng, M.-C. (2008). Driver fatigue and highway driving: A simulator study. *Physiology & Behavior*, 94(3), 448–453. Retrieved from <https://doi.org/10.1016/j.physbeh.2008.02.015> doi: 10.1016/j.physbeh.2008.02.015
- Varun Shiva Krishna Rupani, V. V. S. T., & Tejith, K. (2024). *Real-time drowsiness detection using eye aspect ratio and facial landmark detection*. Retrieved from <https://arxiv.org/abs/2408.05836>
- Wiatowski, T., & Bölcskei, H. (2018). A mathematical theory of deep convolutional neural networks for feature extraction. *IEEE Transactions on Information Theory*, 64(3), 1845–1866. doi: 10.1109/TIT.2017.2776228
- Wijnands, J., Thompson, J., Nice, K., et al. (2020). Real-time monitoring of driver drowsiness on mobile platforms using 3d neural networks. *Neural Computing and Applications*, 32, 9731–9743. Retrieved from

REFERENCES

<https://doi.org/10.1007/s00521-019-04506-0> doi: 10.1007/s00521-019-04506-0

Xu, C., Huang, W., Liu, J., & Li, L. (2025). Detecting driver drowsiness using hybrid facial features and ensemble learning. *Information*, 16(4), 294. doi: 10.3390/info16040294

Yuan, H., Zhao, K., Yan, Y., Wan, L., Tian, Z., & Chen, X. (2025). Long tunnel group driving fatigue detection model based on xgboost algorithm. *Journal of Traffic and Transportation Engineering (English Edition)*, 12(1), 167–179. doi: 10.1016/j.jtte.2023.02.008

Zhao, Z., Zhou, N., Zhang, L., Yan, H., Xu, Y., & Zhang, Z. (2020). Driver fatigue detection based on convolutional neural networks using em-cnn. *Computational Intelligence and Neuroscience*, 2020, 7251280. Retrieved from <https://doi.org/10.1155/2020/7251280> doi: 10.1155/2020/7251280

.1 Code Listings

```
class DrowsinessDetectionCNN(nn.Module):
    def __init__(self, cfg: Optional[Dict[str, Any]] =
        ↪ None):
        super(DrowsinessDetectionCNN, self).__init__()

        # Default configuration
        if cfg is None:
            cfg = {
                'num_classes': 2,
```

```
        'dropout_rate': 0.4,
        'input_channels': 1,
        'input_height': 224,
        'input_width': 224
    }

    self.in_ch = cfg['input_channels']
    # Convolutional feature extractor
    self.conv1 = nn.Conv2d(self.in_ch, 8,
        ↪ kernel_size=3, padding=1)
    self.bn1 = nn.BatchNorm2d(8)
    self.conv2 = nn.Conv2d(8, 16, kernel_size=3,
        ↪ padding=1)
    self.bn2 = nn.BatchNorm2d(16)
    self.conv3 = nn.Conv2d(16, 32, kernel_size=3,
        ↪ padding=1)
    self.bn3 = nn.BatchNorm2d(32)
    self.pool = nn.MaxPool2d(2, 2)

    # Calculate flattened feature size
    feature_h = cfg['input_height'] // 8
    feature_w = cfg['input_width'] // 8
    self.feat_size = 32 * feature_h * feature_w

    # MLP head
    self.fc1 = nn.Linear(self.feat_size, 64)
    self.bn4 = nn.BatchNorm1d(64)
```

```
self.fc2      = nn.Linear(64, 32)
self.bn5      = nn.BatchNorm1d(32)
self.fc3      = nn.Linear(32, cfg['num_classes'])
self.dropout  = nn.Dropout(cfg['dropout_rate'])

def forward(self, x: torch.Tensor) -> torch.Tensor:
    x = self.pool(F.relu(self.bn1(self.conv1(x))))
    x = self.pool(F.relu(self.bn2(self.conv2(x))))
    x = self.pool(F.relu(self.bn3(self.conv3(x))))
    x = x.view(x.size(0), -1)
    x = self.dropout(F.relu(self.bn4(self.fc1(x))))
    x = self.dropout(F.relu(self.bn5(self.fc2(x))))
    return self.fc3(x)

def init_weights(self):
    for m in self.modules():
        if isinstance(m, nn.Conv2d):
            nn.init.kaiming_normal_(m.weight,
                ↪ mode='fan_out',
                ↪ nonlinearity='relu')
            if m.bias is not None:
                nn.init.constant_(m.bias, 0)
        elif isinstance(m, nn.BatchNorm2d):
            nn.init.constant_(m.weight, 1)
            nn.init.constant_(m.bias, 0)
        elif isinstance(m, nn.Linear):
```

REFERENCES

```
nn.init.kaiming_normal_(m.weight ,  
    ↪ mode='fan_out' ,  
    ↪ nonlinearity='relu')  
nn.init.constant_(m.bias, 0)
```

Listing 1: PyTorch implementation of the DrowsinessDetectionCNN with three Conv–BN–ReLU–Pool blocks and a 3-layer MLP head