

Analysis Of Distracted Driver Behaviour Using Self Organizing Maps

by

Matthew Immanuel Samson

A Thesis

presented to

The University of Guelph

In partial fulfilment of requirements

for the degree of

Master of Science

in

Computer Science

Guelph, Ontario, Canada

© Matthew Immanuel Samson, January, 2020

ABSTRACT

ANALYSIS OF DISTRACTED DRIVER BEHAVIOUR USING SELF ORGANIZING MAPS

Matthew Immanuel Samson

University of Guelph, 2020

Advisor:

Dr. David A. Calvert

Driving can be a complicated process, but with sufficient practice, it becomes surprisingly more easier. People tend to forget that even the smallest distractions can have great consequences. Nowadays, experienced drivers are skilled enough to perform multiple tasks like listening to music or texting while simultaneously concentrating on driving. This thesis studies driving under different distractions and how they affect different drivers. The behaviour of individual drivers are also studied to make conclusions on how distractions affect drivers.

To understand a driver's behaviour, their driving patterns are studied by constructing Self Organizing Maps and training them on the drivers' datasets. This results in a structure that maps each driver under a particular distraction to their behaviour. The map is then studied by developing labels based on the features of the datasets. These labels serve as test cases to examine different behaviour of each driver, from which conclusions regarding the disruptiveness of each distraction.

ACKNOWLEDGEMENTS

To my supervisor, Dr. David A. Calvert,

For his impeccable mentoring,

To the members of the advisory committee,

Dr. Hamilton-Wright, Dr. Wirth and Dr. Grewal,

For their insightful comments,

To my family, friends and colleagues,

for their understanding and support,

THANK YOU.

Contents

Abstract	ii
Acknowledgements	iii
1 Introduction	1
1.1 Application of the Self Organizing Map	2
1.2 Brief Overview on Experiments Conducted	2
1.3 Basic Structure of the Thesis	3
2 Literature Review	4
2.1 The Driving Simulator	4
2.1.1 Driving Simulators in Psychology	5
2.1.2 Drive Lab - Experimental Setup	6
2.1.3 Experiments Conducted and Results Obtained	6
2.2 Artificial Intelligence Techniques Applied	10
2.2.1 Machine Learning	11
2.2.2 Brief Intro to Artificial Neural Networks	14
2.2.3 Unsupervised Learning with Neural Networks	16
2.2.4 Self Organizing Maps	19
2.2.5 Other Literature behind the Proposed Model	21

3	Methodology	25
3.1	Description of Data Collected	25
3.2	Feature Engineering and Data Modifications	27
3.3	SOM Architecture and Parameters involved	29
3.3.1	Network Architecture and Training	29
3.3.2	Parameters Involved	30
3.4	Brief description of the SOM Algorithm	33
3.5	Labeling the SOM	36
3.5.1	Thresholding each neuron	37
3.5.2	Labels to be applied	37
3.6	Further Feature Augmentation	39
4	Experimental Results	41
4.1	Result analysis of a single driver	42
4.2	Result analysis of four drivers	57
4.3	Result analysis of 40 drivers:	63
4.4	Thresholding and label to classifier	73
4.5	Driver and Distraction results comparison	107
4.5.1	Individual class analysis	110
4.5.2	Results Analysis and Findings	113
5	Conclusions & Future Work	125
5.1	Conclusions	125
5.2	Future Work	130
	References	133

Appendices	138
A Splitting driver 1's dataset into segments	138
B Driver and distraction analysis by labeling	141
C Combined analysis under each class	151

List of Tables

4.1	Results from the Brake Analysis label	88
4.2	Results from Speed Limits label	90
4.3	Results from Linear Acceleration label	92
4.4	Results from Turning Acceleration label	93
4.5	Results from Altitude Acceleration label	95
4.6	Results from Gap between Lanes label	97
4.7	Results from Split two Segments label	98
4.8	Results from Split three Segments label	100
4.9	Results from Split four Segments label	103
4.10	Results from Split five Segments label	106
4.11	Results containing percentages and number of drivers distinctive under every class under each distraction	113

List of Figures

3.1	A miniature example of the regular dataset being converted into the windowed dataset, with a chosen window size of five	29
3.2	Examining the Decay rate of hyper-parameters. The learning rate (yellow squares) begins at 1.0 and the sigma (blue squares) starts at 0.5.	36
4.1	Examining the brake behaviour of driver 1	44
4.2	Examining the speeding behaviour of driver 1	45
4.3	Examining the acceleration behaviour of driver 1	47
4.4	Examining the acceleration behaviour of driver 1, turning to the left or right	49
4.5	Examining the acceleration behaviour of driver 1, going over lower or higher altitudes	51
4.6	Examining the lane gap behaviour of driver 1	53
4.7	Examining splitting the dataset of driver 1 into five segments	55
4.8	Examining how each of the distractions are visualized on the map	56
4.9	Examining the six feature labels of drivers 1 to 4 using datasets of all three distractions	59
4.10	Examining splitting the dataset of drivers 1 to 4 using datasets of all three distractions	61

4.11	Examining each distraction in the model trained on the first four drivers and on all the three distractions.	62
4.12	Examining each individual driver in the model trained on the first four drivers and on all the three distractions.	63
4.13	Examining the six feature labels of drivers 1 to 40 using datasets of all three distractions	66
4.14	Examining splitting the dataset of drivers 1 to 40 under all three distractions by dividing into	68
4.15	Examining each distraction in the model trained on data of all the 40 drivers and on all the three distractions.	69
4.16	Examining each individual driver in the model trained on data of all 40 drivers and on all the three distractions.	70
4.17	Examining each individual driver in the model trained on data of all 40 drivers, but augmented with the driver label.	71
4.18	Examining how the threshold over the set of activated neurons of driver 1 dataset under the hands-free distraction works	73
4.19	Examining how the threshold over the set of activated neurons of driver 1 dataset under the music distraction works	74
4.20	Examining how the threshold over the set of activated neurons of driver 1 dataset under the text distraction works	74
4.21	Examining applying the brake pressure label over the maximal points set of driver 1	76
4.22	Examining applying the speed limits label over the maximal points set of driver 1	77
4.23	Examining applying the linear acceleration label over the maximal points set of driver 1	78

4.24	Examining applying the turning acceleration label over the maximal points set of driver 1	79
4.25	Examining applying the altitude acceleration label over the maximal points set of driver 1	80
4.26	Examining applying the gap between lanes - label over the maximal points set of driver 1	81
4.27	Examining applying the split into two sections - label over the maximal points set of driver 1	82
4.28	Examining applying the split into three segments - label over the maximal points set of driver 1	83
4.29	Examining applying the split into four segments - label over the maximal points set of driver 1	84
4.30	Examining applying the split into fives sections - label over the maximal points set of driver 1	85
4.31	Examining the activity of drivers under all the distractions under the non-brake class label.	110
4.32	Examining the activity of drivers under all the distractions under the brake class label.	111
A.1	Examining splitting the dataset of driver 1 into two segments	138
A.2	Examining splitting the dataset of driver 1 into three segments	139
A.3	Examining splitting the dataset of driver 1 into four segments	140
B.1	Examining the braking behaviour by separating the classes within the brake pressure label based on percentages	141
B.2	Examining the speeding behaviour by separating the classes within the speed limits label based on percentages	142

B.3	Examining the acceleration along X axis behaviour by separating the classes within the linear acceleration label based on percentages . . .	143
B.4	Examining the acceleration along Y axis behaviour by separating the classes within the turning acceleration label based on percentages . .	144
B.5	Examining the acceleration along Z axis behaviour by separating the classes within the altitude acceleration label based on percentages . .	145
B.6	Examining the lane gap behaviour by separating the classes within the gap between lanes label based on percentages	146
B.7	Examining the activity of drivers under different section of their dataset by separating the classes within the split into two sections label based on percentages	147
B.8	Examining the activity of drivers under different section of their dataset by separating the classes within the split into three sections label based on percentages	148
B.9	Examining the activity of drivers under different section of their dataset by separating the classes within the split into three sections label based on percentages	149
B.10	Examining the activity of drivers under different section of their dataset by separating the classes within the split into three sections label based on percentages	150
C.1	Examining the activity of drivers under all the distractions under the below average speed class label.	151
C.2	Examining the activity of drivers under all the distractions under the above average speed class label.	151
C.3	Examining the activity of drivers under all the distractions under the backward deceleration class label.	152

C.4	Examining the activity of drivers under all the distractions under the forward acceleration class label.	152
C.5	Examining the activity of drivers under all the distractions under the left acceleration class label.	153
C.6	Examining the activity of drivers under all the distractions under the right acceleration class label.	153
C.8	Examining the activity of drivers under all the distractions under the higher altitude acceleration class label.	154
C.7	Examining the activity of drivers under all the distractions under the lower altitude acceleration class label.	154
C.9	Examining the activity of drivers under all the distractions under the left center lane class label.	155
C.10	Examining the activity of drivers under all the distractions under the right center lane class label.	155
C.11	Examining the activity of drivers under all the distractions under the split two 1st section class label.	156
C.12	Examining the activity of drivers under all the distractions under the split two 2nd section class label.	156
C.13	Examining the activity of drivers under all the distractions under the split three 1st section class label.	157
C.14	Examining the activity of drivers under all the distractions under the split three 2nd section class label.	157
C.15	Examining the activity of drivers under all the distractions under the split three 3rd section class label.	158
C.16	Examining the activity of drivers under all the distractions under the split four 1st section class label.	158

C.17 Examining the activity of drivers under all the distractions under the split four 2nd section class label.	159
C.18 Examining the activity of drivers under all the distractions under the split four 3rd section class label.	159
C.19 Examining the activity of drivers under all the distractions under the split four 4th section class label.	160
C.20 Examining the activity of drivers under all the distractions under the split five 1st section class label.	160
C.21 Examining the activity of drivers under all the distractions under the split five 2nd section class label.	161
C.22 Examining the activity of drivers under all the distractions under the split five 3rd section class label.	161
C.23 Examining the activity of drivers under all the distractions under the split five 4th section class label.	162
C.24 Examining the activity of drivers under all the distractions under the split five 5th section class label.	162

Chapter 1

Introduction

Driving a car, like most things, is a complicated process but it can be learnt with experience due to its repetitive and predictive nature. When inexperienced drivers (and even experienced drivers in some cases) become distracted, this leads into accidents that can end up in loss of life. But how do we determine which distractions would need the most attention and end up distracting from safely driving the vehicle? This can be determined by analysing data on drivers who are tested under different kinds of distractions. The goal is to understand specific differences between distractions and between drivers which in turn will lead into making conclusions on the way each distraction affects the behaviour of the drivers whose data is analysed.

By analysing these sequential datasets of different distracted drivers, knowledge can be obtained about their varying driving patterns. This can be used to build models that represent the driver's patterns and then point out patterns that are distinct to each driver and the associated distraction. These patterns can then be analysed to identify various characteristics of individual drivers which are further studied to identify distinctions between drivers under different distractions. The features of

the dataset are also studied as a way to identify differences between drivers and distractions. The differentiating factors are based on the driver behaviour. Some examples of this are when the brakes are applied, when speeds up, when the driver makes left or right turns and many others.

1.1 Application of the Self Organizing Map

A Self-Organizing Map (SOM) or Kohonen Map, which is an unsupervised neural network architecture is applied to all the experiments that are performed in this thesis. Its unique characteristics of dimensionality reduction, simple visualization and spatially distributed structuring are used to gauge patterns found in the dataset. More specifically the nature of the structure formed after training the SOM are used to identify unique driver patterns. The resulting trained model is topologically clustered wherein similar vectors in the datasets that the model is trained on, will be found closer to each other. This leads to the topologically clustered structure of the SOM and the results are represented over a 2D map. Each coordinate on the resulting map is a neuron and it contains a set of the closest matching vectors from the input dataset. When the trained map is tested, the patterns that appear are examined which leads to conclusions on regarding individual drivers and on drivers under the effect of distractions.

1.2 Brief Overview on Experiments Conducted

The initial experiments cover examine how the SOM represents the respective datasets that it is trained on. In this case, tests are conducted to uncover the trained topological structure which are then used to examine further test cases focussing

mainly on two points:

- Analysis between drivers under each of the three distractions to understand the similarities and dissimilarities between them and make conclusions on their behaviour.
- Analysis between distractions on either a single driver at a time or all the drivers in the datasets to make final conclusions on what distractions affect what specific behaviour and if such results are repeated among a majority of the drivers.

1.3 Basic Structure of the Thesis

Following the introduction, the literature review covers all the concepts that will be covered in this thesis including discussions on the dataset that are used, brief introduction on the Artificial Intelligence (AI) techniques that are used leading into explaining the SOM with specific techniques that are applied to the training algorithm. The next chapter covers the methodologies that are used in this research. In this chapter, the datasets are discussed in more detail followed by a description of how the SOM algorithm the respective datasets. The next chapter discusses all the results collected as part of this research, how they are obtained and the reasons for each experiment. The final chapter concludes the research by summarising the results and conclusions obtained followed by a discussion of future research.

Chapter 2

Literature Review

The Literature Review will cover papers on Driving Simulators, their value and the type of data collected along with the various experiments conducted with their respective findings. It will then focus on describing Artificial Intelligence (AI) techniques that could be applied to specific datasets, starting with brief descriptions on numerous Machine Learning (ML) and Artificial Neural Systems (ANS), followed by placing emphasis on Unsupervised Learning (UL) with Artificial Neural Networks (ANNs). This will then lead into the main model used in this thesis called the Self-Organizing Map (SOM) or the Kohonen Map.

2.1 The Driving Simulator

A driving simulator is a machine that utilizes both computer aided motion of an object, together with dynamic simulations. The driver is placed in an artificial environment as a substitute for most aspects of actual driving ^[1]. While in no way is it perfect because of its very nature of “consistent safety” being guaranteed causing drivers to adopt a comparatively care free attitude (even unconsciously), there is a

lot of research that supports the idea of new drivers being trained in a simulated kind of environment. With the growth in technology and the advantages of simulations, today's driving simulators help in studying vehicle design, intelligent highway design and human driving behaviour under different conditions ^[1].

However it is important to note that the very nature of the driving simulator causes a consequential bias in the data collected. When participants in the experiments realize the safe nature of the driving simulator, they tend to be more relaxed in driving. This can not only cause bias but if not given attention, it will lead into data that is far off from the actual data that could be collected in a real-time scenario.

2.1.1 Driving Simulators in Psychology

Driving simulators are valuable research tools which further enable the study of the seemingly simple yet highly complicated activity of driving. Testing out stereotypes was found to yield interesting results, for example, elderly drivers who read documents stating that they are bad drivers (elderly people in general which is a stereotype), found that their driving ability changed in a way that was positively related to following distance but negatively impacted their brake reaction time^[2]. Others include testing the Terror Management Theory where self-awareness leads to anxiety, looking for some social influence like peer influence or the self-concept of driving. In some cases participants who heard “pro-risky” comments (like “go faster!”) drove faster and had more accidents compared to those who heard anti-risky comments (like “drive slowly”). Many interesting correlations were found with age, driving experience, etc. ^[2]. A very interesting study was done with participants exposed to many accident scenes in movies and those who had numerous accidents in video games. It was found that such participants seem to accelerate very quickly up to 160

km/hr and completed their race course much faster.

2.1.2 Drive Lab - Experimental Setup

The driving simulator based on which the data is collected for this thesis is located at the laboratory named as the Drive Lab. It is set up by the Department of Psychology at the University of Guelph. The driving simulator used in the experiments is a model developed by Oktal which in turn constitutes technology (especially the SCANNER software) developed by pioneers from both the industrial and scientific fields of Simulation and Virtual Reality respectively. The model consists of a base full scale simulator connected to the body of a Pontiac G4 convertible. The body of the vehicle is then surrounded by 300° of wrap around projector screens. Almost all the controls that would be present in a vehicle are replicated in the simulator as well, including the steering, pedals and automatic transmission ^[3]. Numerous features are monitored under the Symptom Assessment Scale (SAS) that aim to identify a patient's distress relating to their physical symptoms such as difficulty sleeping, nausea, bowel problems, etc. A rich dataset consisting of numerous features (further explained in the next section 2.1.3) were collected by the simulator at a temporal resolution of 62.5 Hz.

2.1.3 Experiments Conducted and Results Obtained

This section discusses interesting results obtained on by using data from driving simulators. Various driver characteristics were analysed in different environments, to understand behavioural patterns of drivers in general. One of the features analyzed was the factor of experience in driving ^[7]. The study aims to further understand while driving, how different object size and shape (along the course) serve as qualifiers as

to what is relevant to the driver and what isn't. An experienced driver automatically allocates levels of relevancy to the information that enters the drivers' brain at a very high rate as compared to performing any other activity. To do the test, participants were divided into 2 groups of 12. The experiment was conducted by showing both the groups a pair of images. The first image followed by the second image which had a minor change from the first image. The goal for the 2 groups was to identify what the change was while the process was repeated for 300 times. The results are expressed in terms of error times and response times. Novice drivers had higher error rates. The key factor was experience which lead to experienced drivers giving better results on irrelevant objects. Whereas in response times experience was not the key factor. Results showed that safety relevance played a more important role as relevant objects were recognised faster than irrelevant objects. Age did not play a major role in response times but gender was shown to affect the results with males having a faster response time than females.

Another study was conducted on driver distracted data similar to the Dual Task dataset (mentioned in section 2.1.3) where drivers are made to listen to an audio book while driving through simple and complex courses ^[3]. Interestingly distractions were actually helpful in cases of dealing with boredom and fatigue. The experiment conducted uses a dual task methodology to determine whether there was interference between tracking and driving. It was conducted in a simulated environment where the driver was made to be attentive to multiple objects at any given time. The performance decreased significantly when participants were moving while simultaneously tracking objects. It was assessed based on the mean and standard deviation of the distance between the lead vehicle and the participants' vehicle. A total of 53 participants were tested and 28 were selected to do tracking alone and tracking

while driving portions of the study and the remaining 25 did only driving without tracking. The participants reported significant decrease in tracking performance as the number of vehicles increased. So the first group reported poorer results as the number of objects increased. Drivers in the first group displayed shorter fields of view as they had to keep track of many objects while conversely the other group had chosen to give the lead vehicle a greater distance for safety reasons.

Another study tested the yet to be proven idea that driving automobiles required multi-object tracking which shows how many objects can be detected and kept track of while driving ^[4]. The first experiment to understand this employed using a dual task methodology to determine if there is interference in tracking and driving. The participants had to keep track of vehicles, both in their vicinity and other specific “tracking vehicles” either in or out of their vicinity. The impact of how drivers maintain consistent multi-object tracking is also investigated and performance is determined by how many vehicles are in the vicinity of the driver and can be tracked. This experiment shows that while it is possible to perform multi-object tracking while driving, it causes the driving performance is decrease. In another experiment, it was found that drivers are more attentive in detecting changes by about 250 milliseconds faster in the vehicle that they are tracking (for example any change in the vehicle immediately in front of them) but the question investigated was whether multi-object tracking is really beneficial to drivers. The study concluded that multi-object tracking is in many cases advantageous to drivers however it is inconclusive whether it is actually done in day-to-day driving. It was also found that when drivers tried to perform both driving and tracking simultaneously, the performance of both driving and tracking was greatly affected. It was found to reduce a drivers’ stable position on the road, thus compromising steering and speed control. Notably this study was

one of the first to show that multi-object tracking was at least possible while driving.

Yet another study was done on understanding mind wandering while driving, as to whether it would be caused due to longer periods of driving or by differences found in executive working memory that was determined by the Sustained Attention to Response Task (SART) ^[5]. Each participant completed a total of 3 drives each of which were 20 minutes long in the driving simulator and were periodically asked questions of whether they felt distracted while driving. Whenever they reported that they were not focussing while driving, it served as an index for their mind wandering while driving. Another rating that was considered was asking the participants how they felt it was to focus on driving. At four points during each 20 minute drive, the participant was asked if they were thinking about driving and they pushed either the “Yes” or “No” buttons. They reported the difficulty of focussing on driving increased from the 1st to the 3rd drives and there was a marginal increase in steering variability. Another analysis was done that found that a percentage of trials was correlated with reports of emotional state, while the most predictive correlation was found between the mind wandering during the drives and the participants’ respective number of hours of sleep on the night before the experiment.

Many studies were conducted based on accidents, more particularly on vehicle collision type of accidents. Most reasons for vehicle collisions was found to be because of tailgating which was in turn difficult to remediate due to the fact that drivers were poor at estimating distance between vehicles known as the headway. This led to a study done focussing on understanding how headway is maintained by novice and experienced drivers ^[6]. The first goal was to directly compare novice and experienced drivers. Investigation was done into whether to follow experienced drivers as to what

they perceived to be a minimum safe headway or to follow the “2 second rule” intervention. As expected a main effect of driver experience in driver headway was found where experienced drivers revealed that they tended to overestimate where they stop. The second goal was to develop an automated reward based approach to encourage longer headways in order to address tailgating. Current systems of adaptive cruise control and collision avoidance initiate the braking system when vehicles are too close to each other. This automated systems do not solve the problem that initiates tailgating but can also diminish the drivers’ capacity to thwart potential collisions. The study proposes to convey headway distance to the drivers using an in-vehicle display. The system provides the driver with an objective measure of headway in real time together with a long term representation of performance. Comparison with previous techniques showed that the in-vehicle display recorded better and more consistent headway. It was concluded that the introduced system would be used as a tool for early driver training or to help preserve situation awareness when using driver assistance systems.

2.2 Artificial Intelligence Techniques Applied

Artificial Intelligence (AI) is a deeply diverse field encompassing numerous disciplines with the goal of making computers to perform tasks that are only humanly possible. The sections that follow will discuss the common literature in Machine Learning (ML) and Artificial Neural Systems (ANS). This will then lead into the ANS model that would be applied to the driver dataset – the Self-Organizing Map (SOM) and its characteristic features. Then we look into further literature behind the way the dataset is organized, the way in which the model is trained, how the models’ memory works in understanding the data and finally other models that employ

similar mechanisms.

2.2.1 Machine Learning

Most AI algorithms and models are trained, having their foundation in the numerous ways that humans learn from each other and our environment ^[8]. Most methods of learning are divided into 3 types on the basis of how the model or agent perceives the environment and the constraints placed on them. They are:

1. **Supervised Learning:** Learning with the help of a teacher is what is known as supervised learning ^[8]. In this case there is a “Teacher” that gives the model an answer to a question and the model is tasked with understanding how that answer is obtained. The idea is that given enough such “questions and answers”, the model would attain enough understanding to make generalised decisions in that particular field. The “Teacher” can be considered in conceptual terms as one having knowledge on the environment and this knowledge is represented as a set of input – output examples. The questions will be in the form of an input vector of data and the answer will be the desired output given by the teacher. The idea is that the model would make a prediction with the input vector, compare the prediction with the desired output and then send the error (difference between them) back to the model in such a way that the model is able to understand and realize the correct answer (the desired output) by making respective changes to its parameters. Commonly used algorithms include Support Vector Machines, Linear Regression, Decision Trees and Neural Networks.

Supervised learning is generally used into two types of problems:

- **Classification problems:** In classification the input set is divided into multiple classes (usually in discrete fashion) from which the model is made to predict the right class (output) for the right input. Some examples of applications are Spam Detection, Optical Character Recognition (OCR) and any other problems that require knowledge of distinguishing between classes.
- **Regression problems:** In regression the input set is divided into multiple classes (can be both discrete and continuous) in which the model makes a prediction of a quantity that is continuous and not restricted to a set of class labels. Some examples of applications include making predictions on cost values such as house prices, making predictions with respect to time series datasets and basically any problem requiring predictions on specific values as opposed to grouping input data into classes.

2. **Unsupervised Learning:** Learning without the help of a teacher is what is known as unsupervised learning ^[8]. In the absence of a “Teacher”, there is no desired answer (output) that is available for any question. Rather such models focus on understanding the underlying nature by grouping them in numerous ways. Instead of making predictions based on the questions, they are grouped together in such a way that common questions are more close to each other than uncommon questions. It is a kind of self-organized Hebbian Learning ^[9] which helps locate patterns in the data without previously defined labels. Common unsupervised algorithms include clustering – hierarchical clustering, k-means, even neural networks such as Self-Organizing Maps (SOMs), Auto encoders, etc. and other latent variable models like Principal Component Analysis (PCA)

and Singular Value Decomposition (SVD). Further details will be provided in sections 2.2.3 and 2.2.4 as this is one of main components of this thesis.

Two general methods in Unsupervised Learning are:

- **Clustering and Dimensionality Reduction:** The goal of clustering is to group together similar inputs and is used to deal with data containing a large number of dimensions ^[10]. The distance measure usually used to compare objects to determine their group is the Euclidean distance. However two objects may be similar despite differences in position, orientation and scale. Some common applications of clustering are Medical Imaging, Market Research, Social Network Analysis and Software Evolution.
- **Principal Components:** This is a more statistical procedure that aims to convert a set of correlated variables into a set of linearly uncorrelated principal components. It is even a standard technique used to perform dimensional reduction ^[10]. It relies on the concept of eigenvalue decomposition that can be seen as fitting a p-dimensional ellipsoid to the data where each axis of the ellipsoid is a principal component. Common applications include Quantitative Finance and Neuroscience.

3. **Reinforcement Learning:** In both supervised and unsupervised learning, the model is not permitted to look at its environment. This restriction is lifted with reinforcement learning. While there is still no “Teacher” to provide the right answer, the input-output mapping is performed by continuous interaction with the environment in order to minimize a scalar index of performance ^[9]. In place of a teacher, the system is built around a “Critic” that works on a primary reinforcement signal called the heuristic reinforcement signal where both of

which are scalar values. The goal is to reduce a function on the expectation on the total cost of actions taken over a sequence of steps instead of the immediate cost. In many cases, it turns out that actions taken earlier found better results. Such actions should be learned by the system and then taken and feed back to the environment. The environment is usually represented as a Markov decision process (MDP) and many such reinforcement techniques utilize dynamic programming to be solved. It has attained remarkable results in the fields of AI in gaming, robotics and evolution programming.

2.2.2 Brief Intro to Artificial Neural Networks

In the 1940's with the McCulloch-Pitts ^[23] model which was the first mathematical model that used the all-or-none output mechanism implemented by a step threshold function. None of these models had the ability to learn. In 1949, Donald Hebb's studies ^[9] on neurons led him to formulate a learning rule that states that the efficacy of a synapse increases if there is presynaptic activity followed closely with a postsynaptic activity called the Hebbian learning. In 1957, Rosenblatt applied learning rules to the McCulloch-Pitts to develop the "Perceptron" which was shown to learn to separate between two classes ^[25]. This forms the basis for today's general neural network called the "Multi-Layer Perceptron" (MLP). However at the time, criticism by Minsky & Papert together with the Credit Assignment Problem (CAP) encountered in training the MLP caused the research in neural networks to diminish greatly. Interest in neural networks resurged with the Back-Propagation algorithm that was developed individually by Rumelhart & Hinton ^[14] and LeCun ^[24] towards the end of the 1980s, which while not being the perfect solution, widely addresses the CAP.

An Artificial Neural Network (ANN) is a massively parallel distributed processor made up of simple processing units called neurons that have a natural capability to manipulate and store experiential knowledge and making it available for use. It resembles the brain in such a way that knowledge is obtained from the environment by a learning process and memory is stored and represented in the form synaptic weights, which are basically the connections between the neurons ^[8].

A basic ANN model consists of computational units called neurons which are arranged layer-wise. The first layer accepts the input vector and contains as many neurons as the lengths of each input vector. The last layer is called the output layer that contains as many neurons that are sufficient to obtain some understanding regarding the input data (for example, a decision to be made). The problem here is that not every problem can be solved with just an input and output layer and so as and when required, the usage of hidden layers of variable number of neurons has become common place in constructing ANN models ^[11].

ANNs usually store and operate on two types of patterns – spatial patterns which can be represented as a single static image and spatiotemporal patterns which is a sequence of spatial patterns similar to a sequence of static images. The way the memory in an ANN is represented can be understood as either Content Addressable Memory (CAM) where data is mapped to addresses using a matrix that stores all the values of the synaptic weight connections (aptly named the weight matrix), or as associative memories where data is mapped to other data that works by providing output responses based on respective input stimuli ^[12]. The model can either be trained using auto associative memory where all the patterns are stored one after another or by hetero associative memory which stores patterns in pairs. The hetero

associative mapping was later extended from pairs to a kind of window applied and slid over the pattern set ^[13]. The way the patterns (vectors) are stored represents the way that they will be used to train the network.

One of the greatest achievement of the ANN is its capability to satisfy the Universal Approximation Theorem (in theory at least), according to which it is possible to create an ANN model that would generalize to any function, dataset or problem provided that there are enough neurons in its hidden layer ^[11]. Basically no matter what function we wanted to compute, there is an ANN that could do the job. However there's always the problem of having too many neurons that could either result in overfitting the model or add unnecessary complexity (as the theorem does not place any limit on the number of neurons). The beginning of the "Backpropagation" (BP) algorithm ^{[14],[15]} solved the complexities of training multiple layers in a neural network by means of applying feedback accross the network. Feedback in BP is understood as the gradient (or the slope) that is first calculated as the error at the output layer and then transferred as the gradient obtained at the successive layer. The gradient of the error is then applied to each neuron that causes the weights to change in the direction of the optimal result (preferably). If the change leads to a better result, it is kept under certain conditions to counter the local minima problem, the change will either be maintained or discarded thus leading into the basic idea behind the Gradient Descent (GD) algorithm.

2.2.3 Unsupervised Learning with Neural Networks

Unsupervised learning with ANNs has its foundation in the Hebbian learning rule ^[9] where the connections are reinforced irrespective of an error and is specifically a function on the potential between the two neurons of the connection. While the most

commonly used unsupervised learning neural network techniques are the Adaptive Resonance Technique (ART) and the Self-Organizing Maps (SOM), numerous other valuable neural networks will be discussed in this section with SOMs being discussed in the next section.

Grossberg Techniques:

In 1968, Grossberg developed an ANS model that had a single layer of neurons and works as an auto associative, nearest-neighbour classifier to store analog spatial patterns. It learns using either Hebbian or competitive learning and the model is called the Additive Grossberg (AG) model ^[12]. Once trained, it results in a one-layer feedback structure where the neurons correspond to the input features. The neuron activations work by either self-exciting (positive) or neighbour inhibiting (negative) lateral feedbacks. The firing of neurons work as a kind of Short Term Memory (STM) and the final nearest-neighbour classification works dynamically with respect to each input by making the closest resembling neuron to be maximally activated and the least resembling ones to be nullified. At the end of training, the optimal result is that only one neuron would be activated to the maximum at 1 while all others would be inhibited to the minimum which is 0 for each respective input pattern.

Adaptive Resonance Theory (ART):

The Adaptive Resonance Theory was developed by both Grossberg & Carpenter ^[16] continuing from previous research on the AG models and is used to address problems of pattern recognition and prediction. The basic ART is an unsupervised neural network that consists of one layer of neurons called the comparison layer and another layer called the recognition layer. The model accepts a vector of input in the first layer and transfers it to the best match in the recognition layer i.e. the neuron that has

the weight vector that is closest to the input vector. A negative signal proportional to the distance between each respective weight vector and the input vector is sent to all the other neurons thus inhibiting their output and training the network. What's interesting about the competitive learning here is that after the inhibition is done, the model compares the winner with a “vigilance” parameter to make sure that the input vector is in normal ranges as the input seen before. If it is within ranges, the winner activations are updated such that they come closer to the respective input vector, but if the difference is out of the normal ranges, even the winning neuron will be inhibited. This is the reset function that keeps reducing neuron activations one by one that do not overcome the vigilance parameter. The value of the vigilance parameter inherently serves as a control variable for how the memories are to be modified ^[12].

Hopfield Networks:

The Hopfield networks is a variation on the Recurrent Neural Networks (RNN) which feeds its existing output back into the neuron with the aim that the model would understand contextual information. It is a single layer, auto associative, nearest-neighbour encoder similar to the AG model, which works in continuous time and stores analog spatial patterns. The model is trained using a thresholding function. It was in 1985 that it was first applied to an ANS network architecture by Hopfield & Tank ^[19] and it was shown that highly interconnected networks using non-linear analog neurons are very effective in computing and have 3 major forms of parallelization in the input, output and the network interconnectivity between the neurons. The neurons are modelled as amplifiers which use sigmoid monotonic activation functions with their synaptic weights altering between 0 and 1 for normal outputs (excitatory response) and between 0 and -1 for inverted outputs (inhibited outputs). Results were used to solve difficult optimization problems like the Traveling

Salesman problem where each input was the distance between cities. In order to approximate more complex problems, a larger number of neurons would need to be used.

2.2.4 Self Organizing Maps

Tuevo Kohonen began his ANS research working on randomly connected paradigms but then quickly shifted to focus on associative memories. Later improvements in 1973 resulted in the optimal linear associative memory to find an optimal mapping of vectors between associative memory and linear vectors. Kohonen's further research led to the development of a competitive learning algorithm called the Linear Vector Quantization (LVQ) which automatically determines the best reference vectors for a large set of n-dimension data points. This has also been called the self-organizing feature map because of its early success in organizing sounds into a phonotopic map which is a part of the brain responsible for understanding sounds, also called the auditory cortex ^[17].

The idea of self-organized feature maps in a topological manner was first published by Kohonen in 1982 ^[17]. The main discovery was the self-organizing capabilities of a simple network containing adaptive physical elements that received signals from an input space and automatically mapped them onto a set of output responses in a way that the responses seen in the output acquire a somewhat topological design. The discovery was followed by finding out that topologically correct maps of structured distribution were formed from an initial map where no such structure existed (called retino-tectal mapping). The main objective was to demonstrate that external signal activity, assuming a proper structural and functional description of system behavior is sufficient for enforcing such mappings into the system. The first experiment was

done using an array of units, a neighbourhood detecting function and an adaptive process applied on the parameters. The topology of the array is determined by the neighbours to each unit. The mapping is ordered when the neighbours are found to be similar to the matching unit. Through these initial simulations, multiple different results are obtained all having the same meaning but differently distributed. To prevent this, “seeds” i.e. pre-defined input weights should be used. Two phases are described, phase 1 attempts to define and understand the clustering activity and this is fine as long as it attains the proper form. In phase 2, the adaption of the input weights is described. The paper then discusses reasons for choosing a neural network model and explains some problems encountered such as the “pinch” phenomenon which occurs when data vectors do not spread out in a planar form but in the shape of a ring and the “collapse” phenomenon where all the weight vectors obtain the same value and is observed when the range of lateral interaction between vectors and neurons was too large.

Further improvement to the organizing maps leading into the current SOMs (or the Kohonen map) was once again done by Kohonen in 1990 ^[18]. The paper lays emphasis on the interesting spatially organized “internal representations” of the various features of the input signals and their abstractions and this is unique among all architectures and algorithms suggested for neural networks. The topographical organization formed is similar to the cortices of animal brains. After fine tuning the weight vectors, the map can even detect patterns in noisy signals. The spatial segregation of different responses and their organization into topologically related subsets results in a high degree of efficiency in typical neural network operations. The author makes an interesting comparison with studies on “brain maps” that show evidence that internal representations of information in the brain are generally

organized spatially, in theory at least. A single vector of data is given as input, and then a search begins for a match with the right weight vector, usually by Euclidean distance and the winner is the shortest one. Next the weights are updated (never independently as this inhibits competitive learning) and tend to attain values that are ordered along the axes of the network. To enforce lateral interaction in a general form, a neighbourhood set is defined around the winning neuron and its width could be a time-variable i.e. wide at the beginning and slowly narrowing down with its decay following that of a general bell curve function. Some experiments are conducted to show hierarchical representations in data which show that if the input has well defined probability density function, the weight vectors would tend to imitate it. Other experiments include LVQ in the sense of classifying (labeling) each weight vector using a kind of majority voting. Further results include that fine tuning via linear vector quantization is the best approach to classification tasks in the SOM.

2.2.5 Other Literature behind the Proposed Model

Sammon Mapping:

Dimensional reduction is a problem encountered in the field of AI where in the data collected is of many dimensions (features) which makes it very hard to visualize or understand. However in numerous attempts to perform dimensional reduction, starting from using the PCA to extract only principal components from the data, they all result in losing the dataset's underlying structure when being reduced to 3 or lower dimensions. Another technique that works much better is called the Sammon mapping ^[22].

In Sammon mapping, the goal is to preserve the distance as much as possible between each pair of points in a multi-dimensional space when reducing it to 2

dimensions. A good way to start this method is by first picking out the principal components and finding the respective 2D counterpart. Next the error is calculated using a function that finds the distance between each pair of points in both the multi-dimensional space and the 2D space. This error is then used to show how much the next 2D map should be altered such that the inter-point distances in the multi-dimension will be preserved. It should be noted that the final result is not optimal but as close as possible to the optimal solution. Resulting dimension reduction from numerous experiments were found to be much better than PCA [22]. The ML like structure of the algorithm even encouraged an ANN variant of Sammon mapping being developed to take advantage of the numerous parallel computations involved.

SOM on Time Series:

The dataset analysed and applied in this thesis is in essence a Time Series. A time series is a series of data points indexed with respect to time which can be defined as a sequence taken of a successive equally spaced points in time. Time series are used in all sorts of fields like pattern recognition, signal processing, weather forecasting and more recently, in the field of Data Science to build models that make predictions with respect to time. In dealing with Time Series data, SOMs are usually not the first choice. This is because of numerous other supervised architectures such as the regular Multi-Layer Perceptron (MLP) and Radial Basis Functions (RBF) which are chosen thanks to their generalization ability [20]. The dynamics of the time series can be described by means of a nonlinear regressive model such as an autoregressive model based on the idea of making predictions.

The paper even points out interesting reasons for using SOMs such as the local nature and growing architectures of SOMs in general [20]. One of the algorithms

mentioned basically talks about performing vector quantization on vectors together with their associated one step-ahead observations (a kind of windowing and hetero associative mapping) which is the idea behind a VQTAM and this can be used to learn static (memoryless) mapping. The input vectors are arranged in such a way that each vector is a combination of an input and the output vectors (current and the one step ahead). However during updating the weights, two separate updates are done on the respective input and output weight vectors. The VQTAM model is improved through the use of geometric interpolation and smooth output values can also be done using a VQTAM model in the form of RBF-like networks. This method is also applied to system identification and adaptive filtering. Once trained, the results can be interpreted by using rule extraction procedures. Rules are then applied to each winning neuron based on the input values that activate them.

Windowed input to Neural Networks:

One of the keys to understanding time series is maintaining context between subsequent events to so as to form a kind of spatio-temporal structure in order to comprehend the dataset and make predictions. The first neural network model that applied this was the NetTalk to solve speech recognition problems in 1986 ^[13]. The intent behind the model was to try and understand the complexity of learning simply human cognitive tasks, more specifically focussed on converting text to speech, which after training on a corpus of informal words was found to have good performance and generalisation capabilities.

The model used the basic structure behind an MLP network with 3 layers and applied the sigmoid activation function. The input layer consists of 7 “groups” of neurons each taking 7 characters as input leading to the “windowed” input layer

format where each neuron accepts 7 inputs (7 characters). The expected output containing 7 neurons is the center character of each window applied on each group of input neurons. The window is then moved over the sequence of characters by one step and the overall process is repeated. The results showed that a relatively small network could capture most irregularities of the dataset. This is very similar to the discoveries found by the Time Delay neural network in 1989 ^[21]. Its one of the first times the BP algorithm was applied to train neural networks and also when windowed input was applied and slid through the dataset.

Chapter 3

Methodology

The methodology behind this thesis begins by describing the features selected from the dataset in more detail followed by various modifications done to the data. The structure of the network used and the way data is fed into the network is then explained. Then the SOM algorithm is described in detail followed by how the training algorithm is applied on the SOM. Once trained, the SOM mapping is created between each map coordinate from the resultant maps and the data points from the respective dataset. This final trained map is then labeled using numerous existing labels based on data collected from the participants of the driving simulator experiments and other interesting labels obtained from analysing the features of the dataset.

3.1 Description of Data Collected

The dataset that this thesis will be based on is called the Dual Task dataset. It is one of the many datasets collected in the Drive Lab by Dr. Lana M. Trick and her students from the Department of Psychology in the University of Guelph. The

dataset is collected in the form of a Time Series of each driver. It is comprised of 40 young adults who are within the ages of 17 to 22 years old, of whom nine are male and 31 are female. Each participant made to drive under different distractions to examine how they perform multiple tasks whilst driving. The three distractions placed on the drivers were 'Hands Free' (driving while having a conversation via a 'hands free' device), 'Music' (driving while listening to music) and 'Text' (driving while texting). The data is collected separately for each condition and for each driver, so a total of 120 Time Series' (40 drivers x three distractions). Values for all of the features is obtained from the sensors on the driving simulator sampled at 62.5 Hz. The following are the features that are present in the datasets:

1. **Time:** The value of time is being recorded at 62.5 Hz (0.016 s) until the end of the experiment. It is measured in seconds (s).
2. **Brake Pressure:** A sensor at the brake sends signals indicating how much pressure was applied to the brake when it is pressed. When not pressed the values would be very small (approximately zero), however when pressed the values vary from 0.1 to ranges of 50 to 70. It is measured in Bar (1 Bar = 100000 Pascal).
3. **Tangential Speed:** The instantaneous speed measured by the simulator at every time-step. It is measured in m/s .
4. **Tangential Acceleration X:** The change in tangential speed measured along the x-axis with respect to the motion recorded by the simulator in a forward or backward direction (driving on a straight road). It is measured in m/s^2 .
5. **Tangential Acceleration Y:** The change in tangential speed measured along the y-axis with respect to the motion recorded by the simulator in a left or

right direction (driving around a turn). It is measured in m/s^2 .

6. **Tangential Acceleration Z:** The change in tangential speed measured along the z-axis with respect to the motion recorded by the simulator in an up and down direction (driving over varying altitudes). It is measured in m/s^2 .
7. **Lane Gap:** The distance between the center of the simulator and the driving track. It is measured in meters (m).

Other features were collected in the dataset but they were not examined in this work. The “Time” feature is removed since the sequential nature of the data collected is sufficient to understand when each event takes place.

3.2 Feature Engineering and Data Modifications

Resampling the datasets:

The dataset is sampled at 62.5 Hz which means data is collected at every 0.016 seconds. Although this is valuable, in regards to human nature and the respective actions taken, it would be hard to imagine any significant action done in 0.016 seconds in regards to driving a car. For this reason, the dataset is resampled at the rate 10 Hz. However this was not reduced further due to possible risk of loss of data, thus end up in leaving out valuable context between events (each data sample).

General ML related modifications:

The dataset is normalized or scaled within each feature such that every value in each respective feature is between 0 and 1. The normalization is based on the below formula:

$$NormX_i = (X_i - X_{min}) / (X_{max} - X_{min}) \quad (3.1)$$

Where:

X_i = data sample 'i' of feature X

X_{min} = The minimum value in feature X

X_{max} = The maximum value in feature X

Finally the datasets are placed under a window of size five meaning that the network is made to be trained by sliding windows of width five with an overlap of four samples. This means that instead of showing the network one data point at a time (which is resampled at $1/10^{th}$ of a second), we show the network five data points at any given time (so five times $1/10^{th}$ of a second is half a second) and this is described in fig. 3.1. Similarly the labels to be tested over the resulting trained SOM are placed under the same window of size five and the center of the window is chosen to be the main label for every window of input data points to the network. However consideration was put into choosing the window size with the aim of preventing too much context (large window) that causes loss of previously learned information, while also preventing the network from being exposed to too little context such that it never realizes there exists any context in the first place. **Final dataset arrangements:**

The final datasets are arranged in the four following sets before being used to train numerous SOM models, which will be discussed in the next chapter:

1. **Hands-Free dataset:** Includes the data from all the drivers while placed under the distraction of using a hands-free device while driving.
2. **Music dataset:** Includes the data from all the drivers while placed under the distraction of listening to music while driving.
3. **Text dataset:** Includes the data from all the drivers while placed under the distraction of texting on a hand-held device while driving.

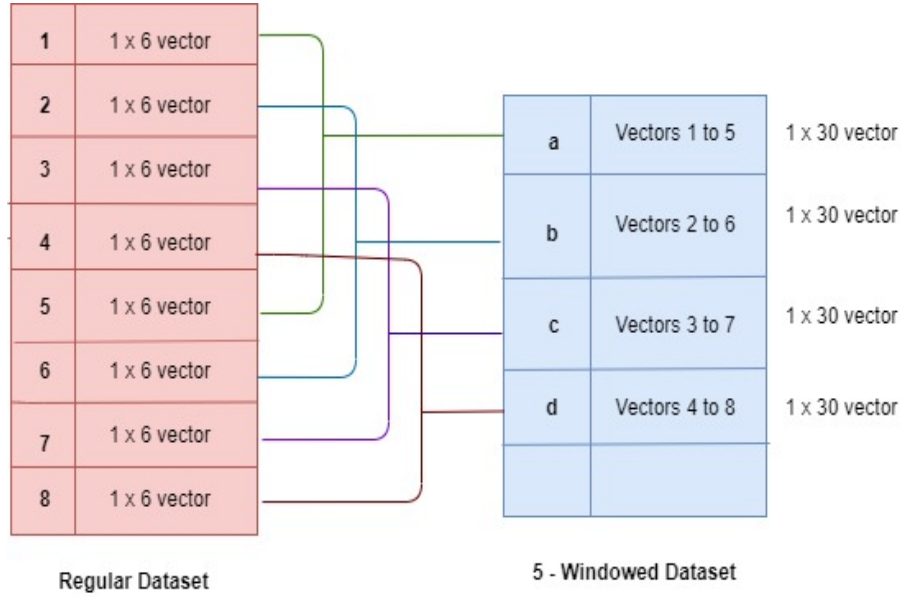


Figure 3.1: A miniature example of the regular dataset being converted into the windowed dataset, with a chosen window size of five

4. **Full dataset:** Includes the data from all the drivers while placed under the 3 distractions mentioned above. It is a combination of all the datasets - the hands-free dataset, the music dataset and the text dataset.

3.3 SOM Architecture and Parameters involved

In this section, the architecture of the SOM model is described along with the way the model is trained. This is then followed by describing the parameters involved and the reasons why their respective values are decided.

3.3.1 Network Architecture and Training

The architecture of the SOM consists of 2 layers. The input layer contains as many neurons as the features in the dataset. In our case the dataset has 6 features, however

with the window size of 5 meaning that when the network is trained, it uses 5 data points at the same time with each having 6 features, thus leading to a total of 30 features which in turn leads to having 30 neurons in the input layer. The output layer is the resulting self-organized mapping where each coordinate on the 2D map corresponds to a neuron. The size of the maps (depending on the number of neurons in the output layer) are varied based on the size of the dataset with larger maps (more neurons) being given for the model which trained on the larger dataset.

Before training the networks, the weights to each neuron in the output layer are first randomly set, where each weight is a randomly generated vector of length 30 (number of features after applying the window). Once initialised, the weights are saved and repeated for maps of similar size, so as to have a uniform start to training each model. The training of the network is carried out using the sequential batch training of the windowed dataset. The process starts at the first windowed vector and goes all the way to the end and are then repeated for as many iterations as required, in the case of the experiments conducted the number of iterations used is 10.

3.3.2 Parameters Involved

The control variables used in the SOM algorithm are pre-set after much experimentation and are reasoned as follows:

Input layer neurons: The number of neurons in the input layer is determined by the number of features in the each vector that would be used to train the network. So in the case of the original dataset of 6 features, after applying the window of size 5, each vector contains 30 values and thus each input layer would contain 30 neurons. In the next case, of generating the individual driver as a new feature, there are 7 features and after applying the window of size 5, each vector contains 35 values and

thus each input layer would contain 35 neurons.

Output layer neurons: The number of neurons in the output layer is represented by size x size, where size can be any number. So for example when size = 20, a 20x20 map (which is the output layer) will be generated which is a total of 400 neurons with each neuron being represented as an (x, y) coordinate on the 20x20 map.

Some main cautions are against using too large or too small maps. Large maps tend to cause different patterns (comparatively recessive patterns) to become more visible, however too large of a map causes the same patterns that ought to appear together to become separated and thus end up missing or misinterpreting such patterns. Smaller maps tend to cause different patterns to come together (resulting in comparatively dominant patterns), however maps too small cause the patterns to simply overlap each other to a point where there wouldn't be any discernable patterns.

Sigma: The value of sigma indicates the distance of the neighbourhood of each neuron when being used in weight updates. This neighbourhood function determines how much each weight vector within the neighbourhood is updated and this update decreases from the neurons closer to the BMU to the neurons within the neighbourhood but furthest away from the BMU. Further description is given in section 3.3. After experimentation, it was determined that the sigma value in all further experiments would be set to 1.0 as this creates a neighbourhood that starts off by covering the entire map and then slowly becomes decayed at the end of each epoch.

Learning Rate: The value of the learning rate determines how big of a change would be applied to the weights during their respective updates to prevent over or

under fitting. However, in an unsupervised learning environment, it is hard to define when models over fit or under fit the dataset. The learning rates are defined between 0 to 1 which works as a control over how much the weights within a neighbourhood are updated. In this case, once again after much experimentation it was found that the highest learning rate (1.0) caused the network weights to be trained much faster however causing changes so fast that important patterns tend to overlap each other. Learning rate that is too small (lesser than 0.01) requires too much training time and even causes important features to be more scattered as opposed to coming together. To prevent both a learning rate that is too fast and one that is too slow, 0.5 is chosen as an intermediate between the two. Once again further description is provided in section 3.3.

Number of Epochs: The number of epochs determine how many number of iterations that the model is trained on. A single iteration is usually not sufficient for a model to learn as the changes would apply to the whole dataset but the model would not retain enough "knowledge" (changes to weights) to adapt to the dataset. So this process has to be repeated numerous times. Too few cycles or epochs of training results in models that do not learn enough about the dataset, however too many cycles would usually cause the network to over train, but in unsupervised learning, these would be wasted iterations and time as the control variables would be decayed to an extent that the network would not learn anymore. For example, a when sigma is too small, the resulting gaussian function would be too small to make enough of a change over the neighbourhood matrix. Hence after much experimentation, seven epochs was determined to be the number of epochs used to train all the models that will be discussed in the next chapter.

3.4 Brief description of the SOM Algorithm

The same SOM algorithm given below is applied on all the models trained as part of this thesis. The only changes are noted in the way the input vectors are shown to the network (after applying the window), either in batch – one vector after another in a sequential order or in random – one vector is picked at random from the whole dataset and this is repeated as many times as required.

A brief description of the SOM algorithm is as follows:

Step 1: Initialisation

The randomly generated weights are saved. The parameters –

1. **Sigma:** Used to determine the neighbourhood of the winning neurons whose weight vectors will need to be updated. It is initialised to 1.0. Sigma is applied into the neighbourhood function as a value between 0 (indicates that the neighbourhood distance is only restricted to each single neuron) and 1 (indicated that the neighbourhood distance covers the whole distance of the map).
2. **Learning rate:** Used to determine how much each neuron's respective weight vector would be updated. It is initialised to 0.5. Learning rate is between 0 (no change in weight updates) and 1 (maximum change in weight updates).

At the end of every epoch of training, the hyper-parameters will be decayed with the aim of learning smaller patterns than those learned in the preceding epoch.

Step 2: Best Matching Unit - Competitive Learning

Training begins by passing the first input vector into the input layer. Based on

this input vector, the map will be activated. This means that the weight vector that most closely matches the input vector is selected and the neuron in the output layer to which the weight vector is connected is called the Best Matching Unit (BMU). This returns the winning neurons respective coordinate to next be used in updating the weights once the neighbourhood of the BMU is acquired in the next step.

Step 3: Neighbourhood Function - Cooperative Learning

Next the neighbourhood of the BMU will be found using a neighbourhood function. The function so chosen is the Gaussian function that will be applied to the (and is centered around the) BMU coordinate and the sigma hyper-parameter such that:

$$G = \text{outerproduct}(e^{(x_{range}-x_{win})^2/d}, e^{(y_{range}-y_{win})^2/d}) \quad (3.2)$$

Where:

G = neighbourhood matrix $d = 2 * \pi * \sigma^2$

(x_{win}, y_{win}) = BMU coordinate

x_{range} = the set of values along the X-axis (0 to map size)

y_{range} = the set of values along the Y-axis (0 to map size)

This is then followed by applying the calculated product of the neighborhood and learning rate over the weight matrix. This product is then summed with the current weights to result in the new and updated weights. The weights within the neighbourhood will be updated with the strength of the update being the highest at the BMU and lowest around the edges of the neighborhood. The updated weights are got by:

$$Weights_{Updated} = Weights_{Previous} + (G * Lr) \quad (3.3)$$

Where:

G = neighbourhood matrix

Lr = learning rate

Step 4: Repetition

The above steps 2 and 3 are repeated for every 30 featured input vector and once the last vector is complete, the end of one epoch is reached.

Step 5: Epochs and Decay

The total number of epochs selected to train all of the models is seven. At the end of each epoch, the hyper-parameters sigma and learning rate are decayed using an asymptotic decay function. They both use the same function given below:

$$decay_{new} = (decay_{prev}) / (1 + ((iterations_{current}) / ((iterations_{total}) / 2))) \quad (3.4)$$

Where:

$iterations_{current}$ = current iteration number

$iterations_{total}$ = total number of iterations

Both the hyper-parameters are decayed, one after the other, after which the above steps 2, 3 and 4 are repeated for as many epochs as required. An example of the way both sigma and the learning rate are decayed is shown below in fig. 3.2.

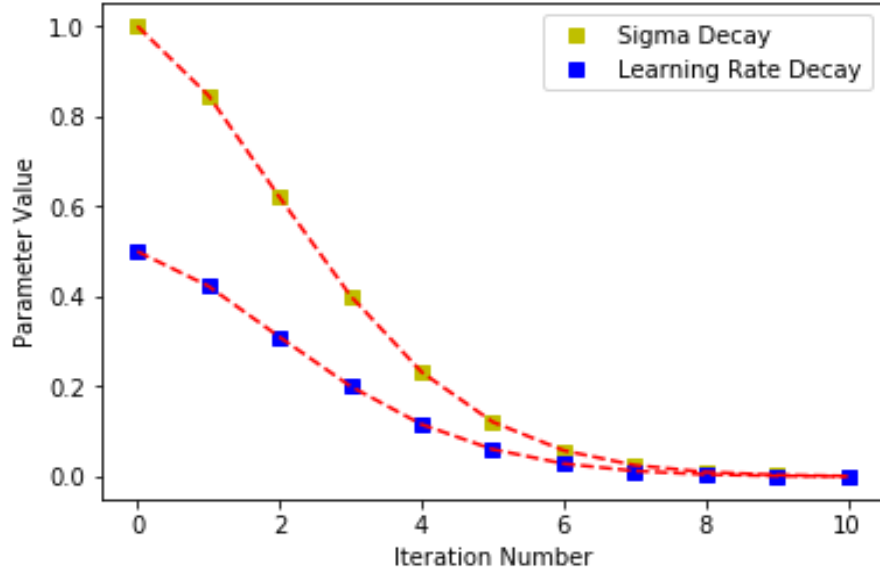


Figure 3.2: Examining the Decay rate of hyper-parameters. The learning rate (yellow squares) begins at 1.0 and the sigma (blue squares) starts at 0.5.

3.5 Labeling the SOM

Once the SOM is trained, the mapping between the neurons (map coordinates) and the data points in the respective datasets is generated in a one to many kind of matching where one coordinate has a list of data points – further showing similarities in data points mapped to the same neuron. This mapping will then be saved and reloaded so as to reuse the same mapping testing many different kinds of labels.

The labels are selected by understanding and analysing each feature of individual participants' respective driving datasets. Labels generated from feature analysis are then applied on all the drivers in the respective map. They are applied with the aim of identifying certain specialized driving patterns of each driver in order to:

1. Compare individual driver patterns (obtained from the same label) with every other driver in the dataset to make conclusions about either trends found

among drivers or patterns that are unique to particular drivers.

2. Compare individual driver patterns (obtained from the same label), not with other drivers but among themselves under each of the distractions discussed to observe how much each driver is affected by the different distractions.

3.5.1 Thresholding each neuron

Once the mapping is obtained, the map is plotted based on each respective label. However when there are more than two labels, most of the coordinates have lists of points which map to multiple labels and this in turn results in each coordinate having multiple labels making it hard to gain any insight into how the labels and data are organised. To overcome this, each point is passed into a maximum threshold that returns only the label with the maximum number of instances such that each point is now indicated by only a single label. This is applied to all the maps once their respective mappings have been completed and saved.

3.5.2 Labels to be applied

The different feature based labels that are applied to all the trained SOM models, each of which are discrete in nature and their descriptions are as follows:

1. **Brake analysis:** Locations where the driver applies the brakes are labeled versus when the brakes are not applied. This is done by examining the brake pressure feature.
2. **Speed limits:** Locations where the driver speeds up (goes above the average speed) versus where the driver slows down (goes below the average speed). This is done by examining the tangential speed feature.

3. **Linear acceleration:** Labels to identify increasing speed rates along the x-axis, so between forward and backward directions to understand when the driver accelerates versus when the driver decelerates are labeled. This is done by examining the acceleration along the x-axis feature.
4. **Turning acceleration:** Labels to locate when the driver makes a turn are obtained using the y-axis to indicate when they turn left versus when they turn right. This is done by examining the acceleration along the y-axis feature.
5. **Altitude acceleration:** Labels to find when the driver goes over varying altitudes using the z-axis, so as to differentiate between acceleration over a higher altitude versus the acceleration down a lower altitude. This is done by examining the acceleration along the z-axis feature.
6. **Gap between lanes:** Labels to identify when the driver comes away from the center of the lane towards the left versus when they drive closer to the right of the center of the lane. This is done by examining the lane gap feature.
7. **Partition the datasets:** Label different parts of the each dataset to visually analyze if certain sections of a driver's are either similar to or different from the others.
 - (a) **Two Partitions:** Dividing the dataset into two equal sections and labeling them as the first section versus the last section.
 - (b) **Three Partitions:** Dividing the dataset into three equal sections and labeling them as the first, second and then the last section.
 - (c) **Four Partitions:** Dividing the dataset into four equal sections and labeling them as the first, second, third and then the last section.

- (d) **Five Partitions:** Dividing the dataset into five equal sections and labeling them as the first, second, third, fourth and then the last section.
- 8. **Distraction Type:** Each participant is labeled based on the type of distraction they are under to understand driver behavior under different distractions.
- 9. **Driver ID:** Each participant has an individual unique label to be distinguishable from each other.

3.6 Further Feature Augmentation

The SOM's topological structure ensures that similar data points are closer to each other while dissimilar ones are further away from each other as seen especially in the feature analysis labels. It also shows sub-clusters (where each coordinate on the map or neuron is mapped to a list of data points) found in each cluster which display a kind of hierarchical arrangement. However ironically, it is this interesting property that also makes the picking out of specific patterns to be particularly difficult, especially in the context of this thesis and the corresponding time-series datasets involved.

As shown in all of the feature analysis labels, the structure formed by the SOM is limited by the features that are used as input. To overcome this more features are introduced into the SOM so as to obtain the required patterns. Overall this can be a complicated feature engineering problem, so we make it the main aim of this thesis to distinguish between individual participants in the experiments. While there are numerous ways to investigate the datasets, so to simplify this process, the main focus is placed on observing the structure based on analysing the differences between each respective driver.

This results in augmented datasets where the only change is that instead of 6 features in the dataset, there are now 7 features courtesy of the segmented feature. The main change here is that once the window of size 5 is applied, each vector will now contain 35 features (previously 30 features). The models are now trained and follow the same approach where the mapping of each model is obtained and the testing is conducted using the same set of labels.

In the next chapter, all the results will be discussed along with all the parameters mentioned above, with only changing parameter being the number of neurons in the output layer (map size) with respect to the dataset that will be used.

Chapter 4

Experimental Results

This chapter is a compilation of all the significant results obtained as part of this research. The results are divided by sections based on the models that are trained. The first section covers results from models trained on the datasets of a single driver over which thresholds are applied followed by applying the labels which are completely described in section 3.4. The next section covers results similar to the previous section except that the models are trained on multiple drivers' datasets to understand specific differences between individual drivers.

The final section of models includes those that are trained on all the drivers' respective datasets. Once the main labels are discussed, the following steps are taken:

1. Thresholds are then applied, similar to the previous sections except that each dataset is looked at one at a time versus all the other datasets mapped to the model. This results in a filtered set of points (the set of most activated points) that represent an individual dataset as opposed to every single point to a vector in the dataset. These maximal points denote unique patterns of each individual dataset in the model.
2. The same labels are once again applied, but this time only to each maximal set

of points. This will result in a simpler and more efficient way to understand individual driver behaviour in relation to all the drivers in the dataset as opposed to looking at the complete mappings of individual datasets. Each label is then applied as a classifier that then calculates the percentage of how much each maximal set of points belong each class within each label.

3. After the percentages of each class within each label are calculated and compared with each other, each class is then compared based on the three distractions to understand and make conclusions on individual driver behaviour under different distractions.

Before discussing all the results found, the values assigned to the hyper-parameters are assigned as follows:

1. Sigma is set to 1.0 which encompasses the whole map.
2. Learning rate is set to 0.5 to prevent very fast and very slow convergence.
3. Input layer contains 30 neurons which is six features times a window of five.
4. Output layer contains a varying number of neurons according to each model.
5. Number of iterations that every model is trained on is seven.

4.1 Result analysis of a single driver

The models used in this section are trained only on a single driver (the first driver in the dataset). The trained models are not modified in any way and are only labeled with different labels. The models trained are as follows:

1. **Hands-Free distraction driver 1:** This model is trained on the “Hands-Free” distraction of driver 1 and the number of output neurons is 900 neurons over a 30x30 size map.
2. **Music distraction driver 1:** This model is trained on the “Music” distraction of driver 1 and the number of output neurons is 900 neurons over a 30x30 size map.
3. **Text distraction driver 1:** This model is trained on the “Text” distraction of driver 1 and the number of output neurons is 900 neurons over a 30x30 size map.
4. **Combination datasets of driver 1:** This model is trained on the combination of datasets of all the distraction of driver 1 and the number of output neurons is 900 neurons over a 30x30 size map.

Label 1 Brake Analysis: Labeling the models mentioned above by when the driver applies the brakes (green) versus when the brakes are not applied (red) as shown in figure 4.1. This based on the brake pressure feature in the dataset.

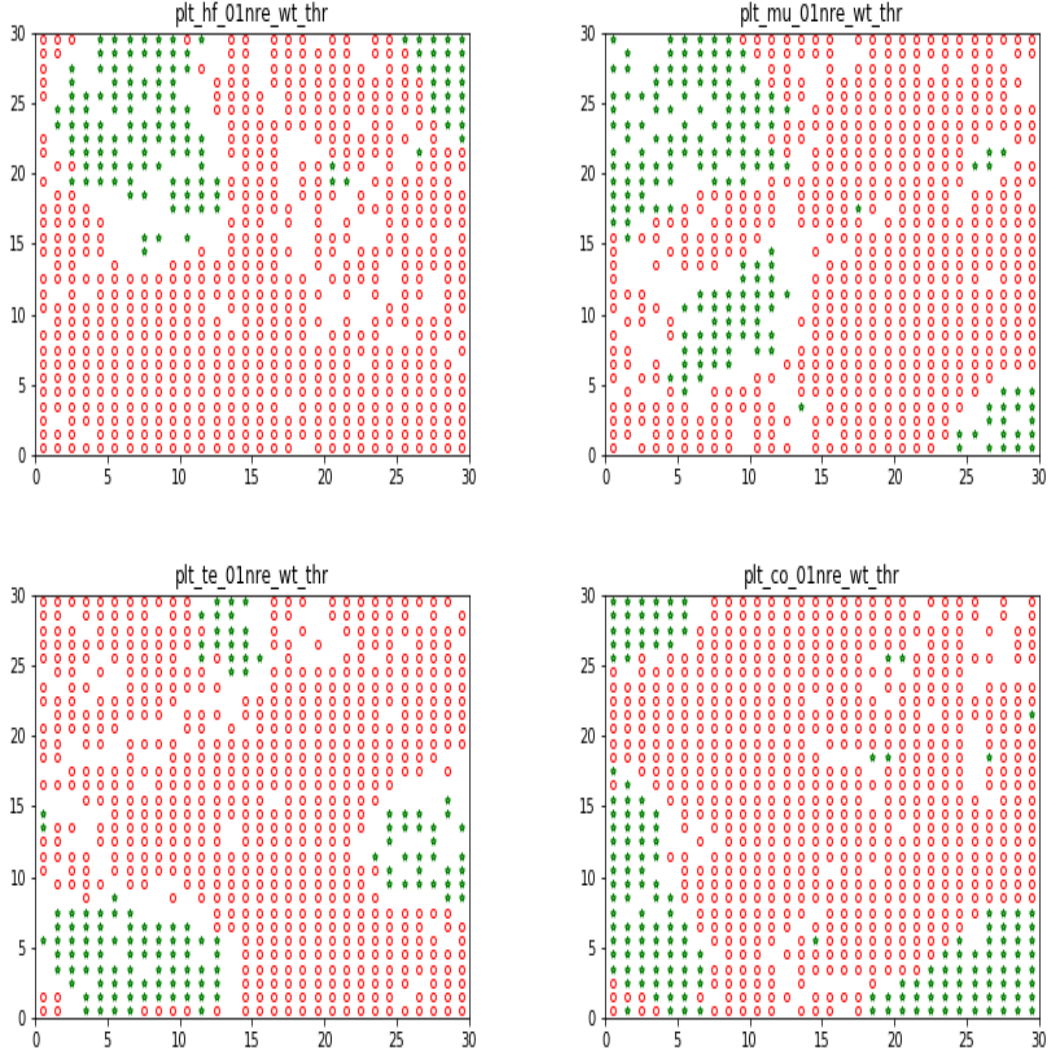


Figure 4.1: Examining the brake behaviour of driver 1
(a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
(c) Text distraction (bottom-left) and (d) combination of datasets
(bottom-right).

There is a clear divide between vectors (the red and green points) when the driver applies the brakes versus when the brakes are not applied. This is further strengthened by observing the space between the two labels. This denotes that there are neurons which do not map to any vectors thus proving the map's strong

distinction between “brake” and “non-brake” vectors.

Label 2 Speed Limits: The average speed found in all the 120 datasets was found to be 77 km/hr. Labeling the models mentioned above when the driver speeds up and goes above 77 km/hr which is above the average speed (green) versus when the driver slows down going below the same average velocity (red) as shown in figure 4.2. This is based on the tangential speed feature in the dataset.

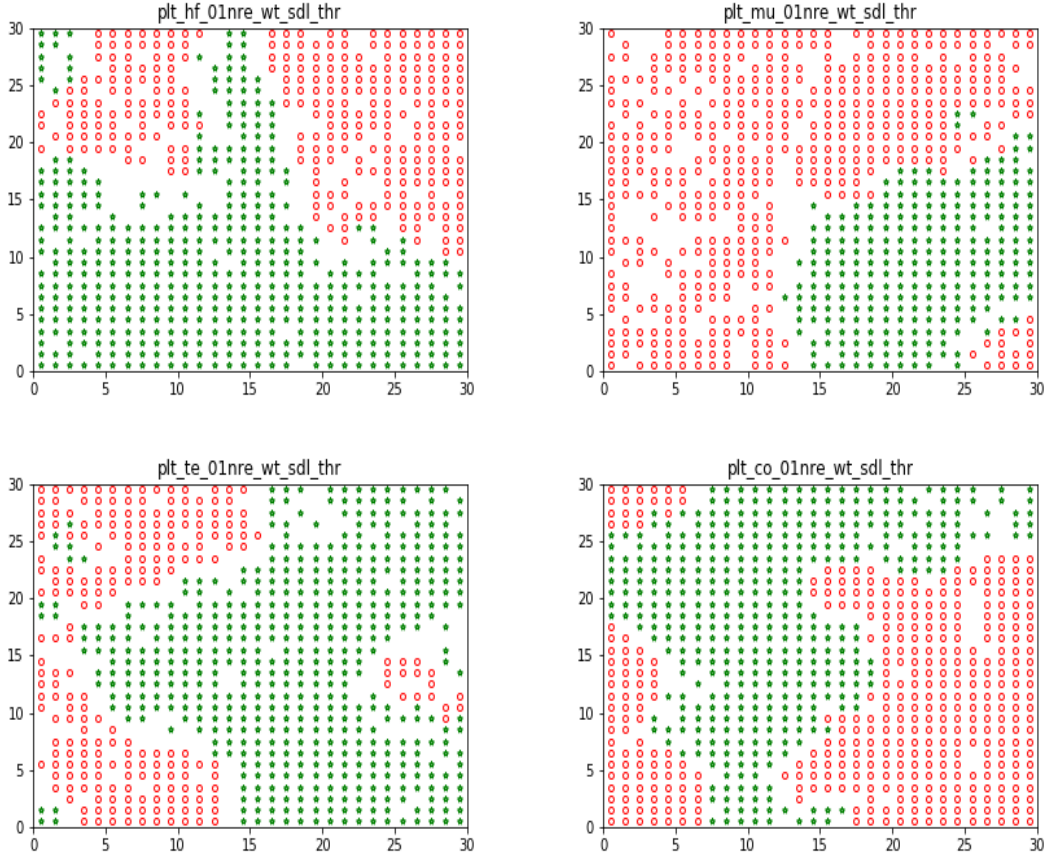


Figure 4.2: Examining the speeding behaviour of driver 1 while under (a) Hands-Free distraction (top-left), (b) Music distraction (top-right), (c) Text distraction (bottom-left) and (d) combination of datasets (bottom-right).

Similar to the previous label, it is clear there is a divide between when the driver speed up versus when they slow down. Another indication is that the labels in this map are almost opposite to the labels under the previous braking behaviour. This illustrates that speed and brake pressure are inversely related to a certain extent such that lower the speed, higher is the pressure being applied on the brakes so indicating that the driver is trying to slow down.

Label 3 Linear Acceleration: Labeling the models mentioned above based on when the driver accelerates (positive values) (green) versus when the driver decelerates (negative values) (red) and this is shown in figure 4.3. This is based on the acceleration along the X axis feature based on the forward motion observed in the dataset.

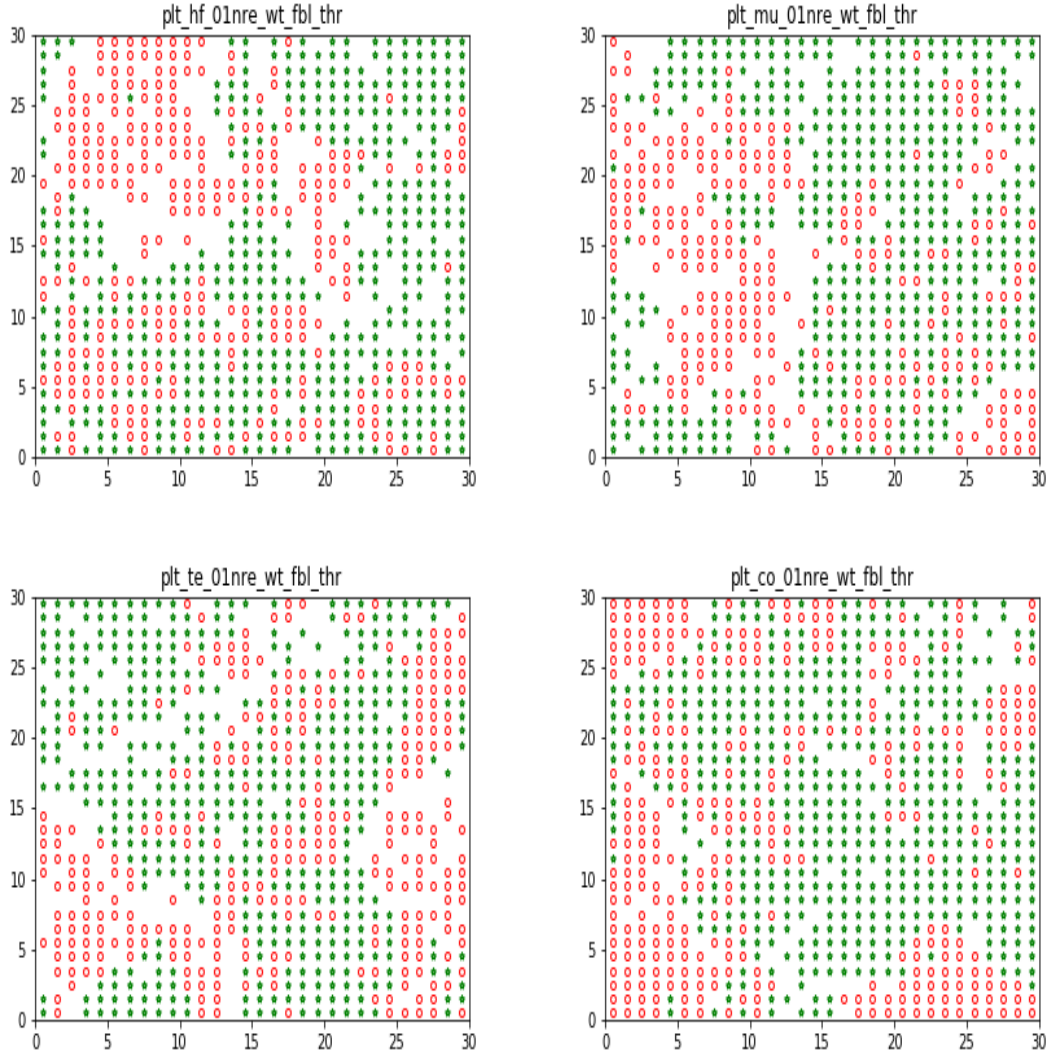


Figure 4.3: Examining the acceleration behaviour of driver 1
 (a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
 (c) Text distraction (bottom-left) and (d) combination of datasets
 (bottom-right).

In this case, the label appears a bit more scattered as compared to the previous two labels. It appears that the strength of the values of the acceleration along X feature is comparatively lesser than that of the previous labels and hence the labeling shows that there are no large sections having a common label. This is in comparison

with the results from label 1 where there is a divide between labels, however there does not seem to be such a divide in this case. Once again it appears that the acceleration is related to the other features, for example, when the driver accelerates, there is no pressure on the brakes and there would also be a possible change in label from below average speed to going above the average speed.

Label 4 Turning Acceleration: Labeling the models mentioned above based on when the driver accelerates either to the right (positive values) (green) or to the left (negative values) (red) and this is shown in figure 4.4. This is based on the acceleration along Y axis feature and can help identify when the driver turns to the left or right.

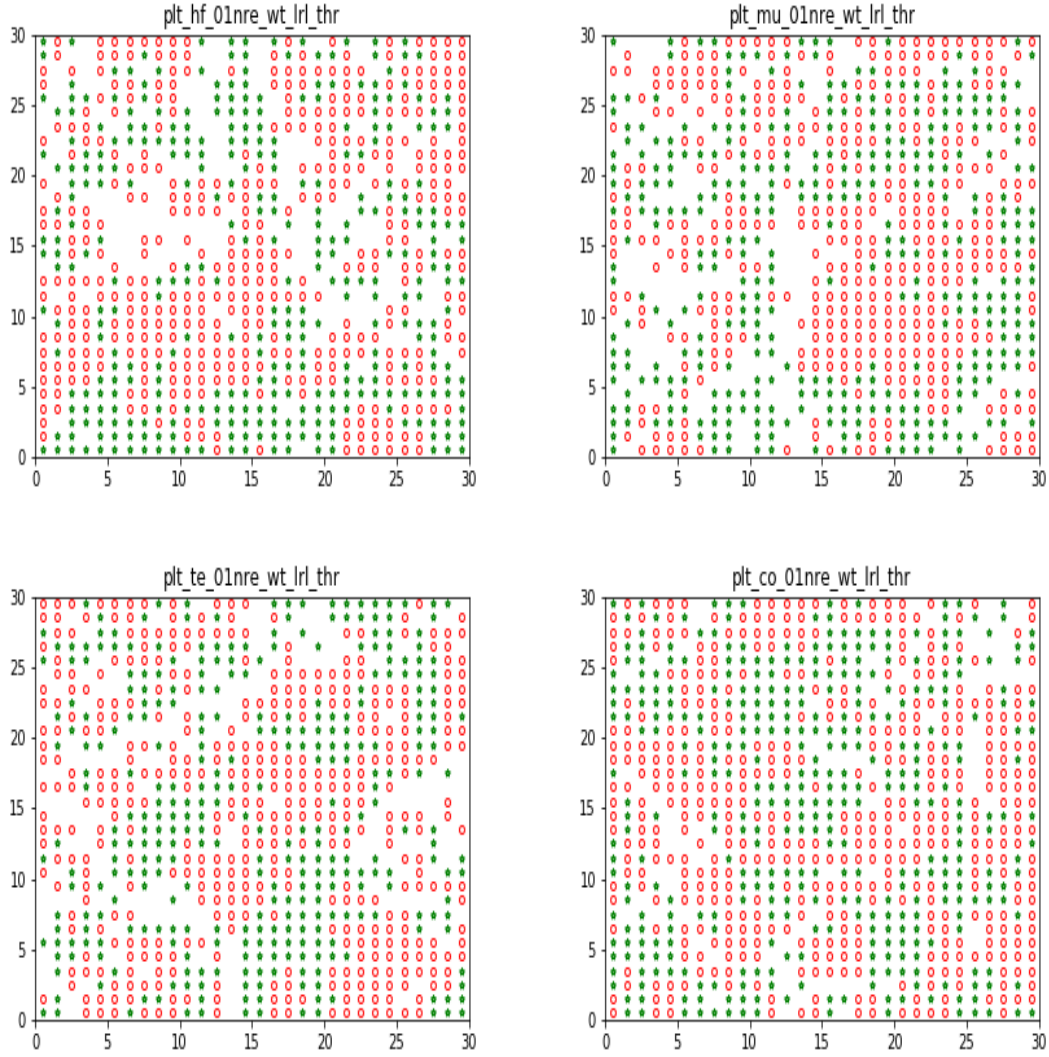


Figure 4.4: Examining the acceleration behaviour of driver 1, turning to the left or right
 (a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
 (c) Text distraction (bottom-left) and (d) combination of datasets (bottom-right).

When analysing the acceleration behaviour in turning left or right, the labeling shows that although sections are divided similar to the previous acceleration along X axis label, the vectors are closer to each other which causes the sections to be larger.

For example, the driver's acceleration along the Y axis would be more towards the left (red) when turning left and even the brakes would be pressed either before or during this action. Also the speed would change while turning.

Label 5 Altitude Acceleration: Labeling the models mentioned above based on when the driver accelerates over a higher altitude (positive values) (green) versus when the driver accelerates over a comparatively lower altitude (negative values) (red) as shown in figure 4.5. This is based on the acceleration along Z axis feature in the dataset.

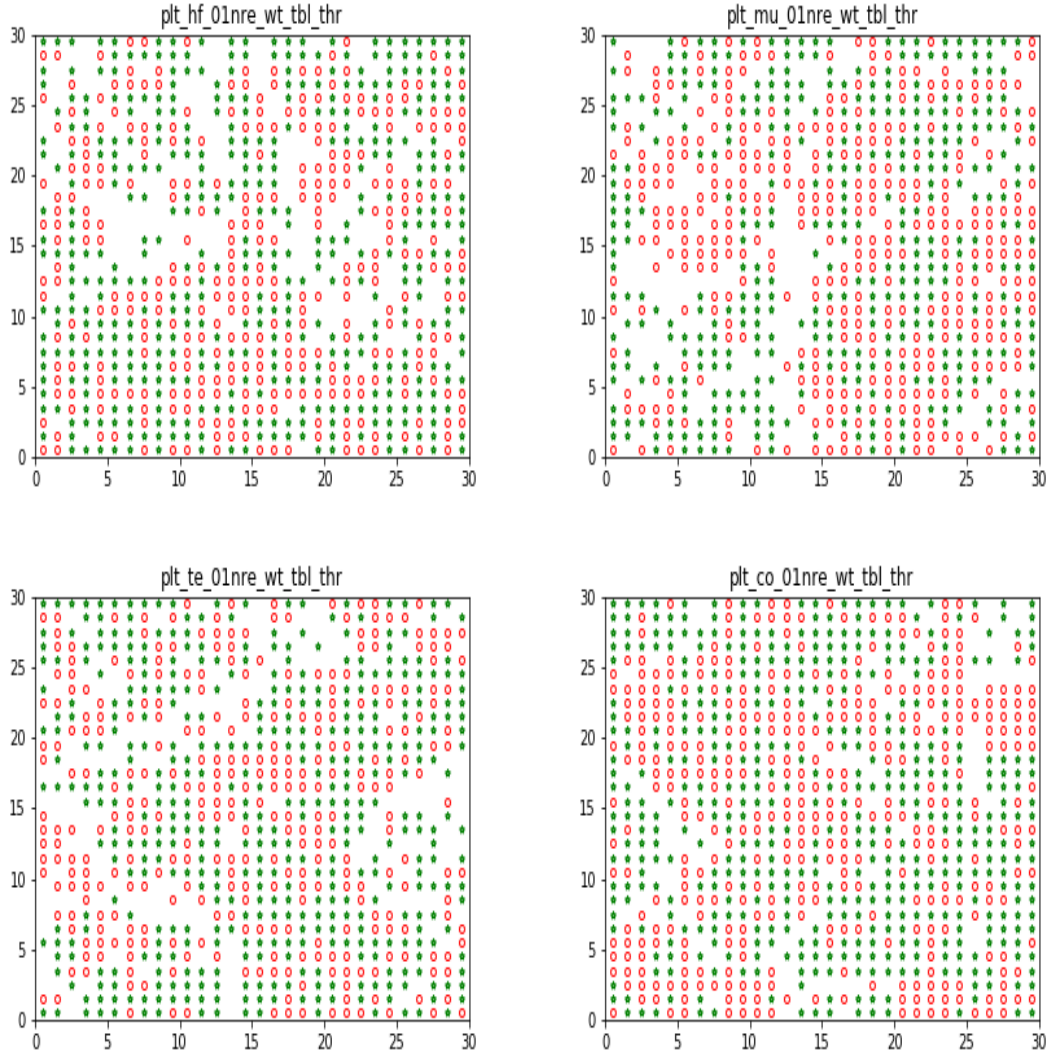


Figure 4.5: Examining the acceleration behaviour of driver 1, going over lower or higher altitudes
 (a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
 (c) Text distraction (bottom-left) and (d) combination of datasets (bottom-right).

In this labeling, the sections are much more scattered indicating that either there is variation in the dataset but not over long enough ranges to cause large sections to appear on the map or that the other features are more prevalent such that the

vectors tend to be arranged more in such a direction.

Label 6 Gap between lanes: Labeling models mentioned above based on when the driver tends to drive to the right of the center of the lane (positive values) (green) versus when driving to the left side of the lane (negative values) (red) as shown in figure 4.6. This is based on the lane gap feature in the dataset.

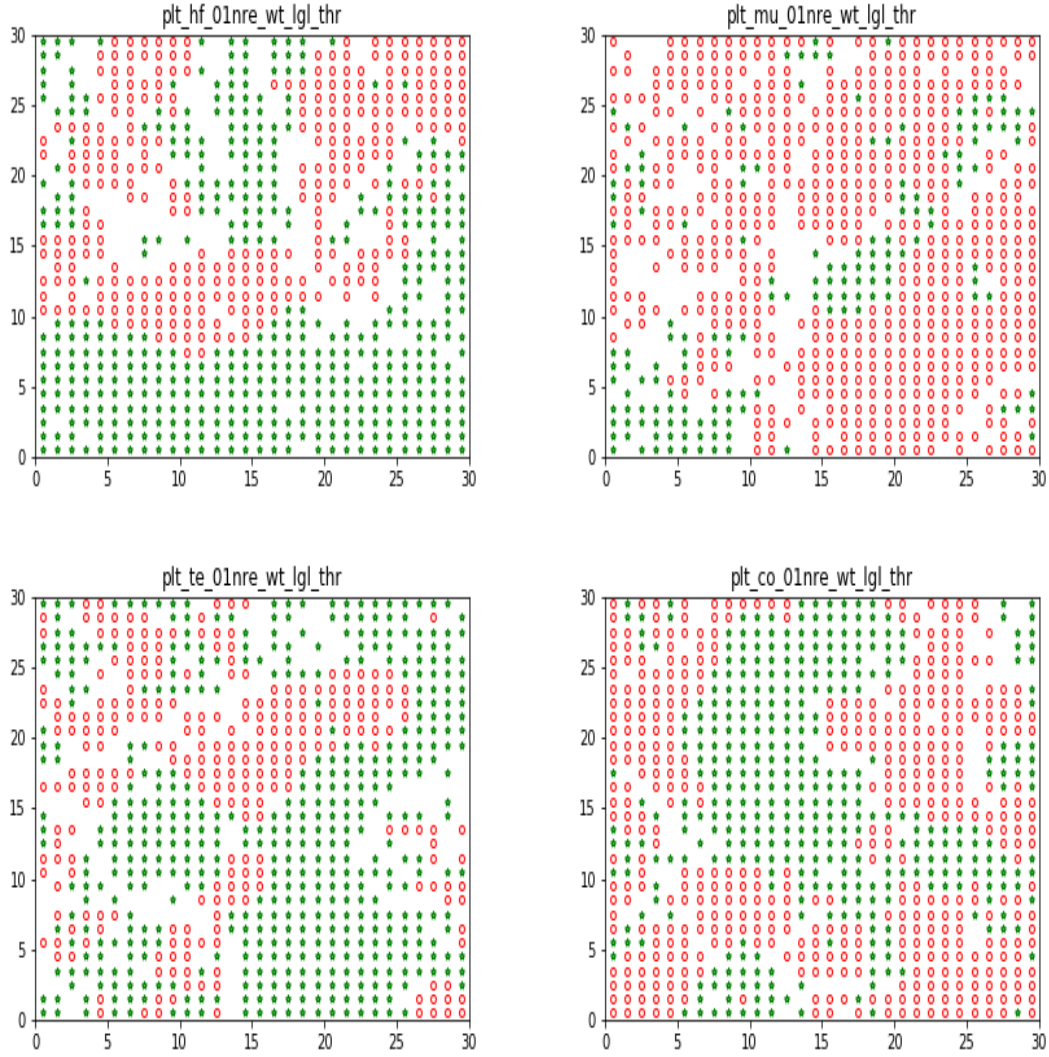


Figure 4.6: Examining the lane gap behaviour of driver 1
(a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
(c) Text distraction (bottom-left) and (d) combination of datasets
(bottom-right).

This label shows more clear partitions in the results between the red and green portions as compared to other acceleration labels and so it more strongly activates the SOM.

Label 7 Split Sections: Labeling the models mentioned above to represent the sequential flow of dataset. This is done by dividing the dataset into sections and labeling each section represent when most actions are performed.

Each dataset is tested by dividing them into segments. Experiments include dividing each dataset into two to five segments, however the results shown in figure 4.7 are where each dataset is divided into five segments – the first section (red), the second section (green), the third section (blue), the fourth section (yellow) and the last section (black).

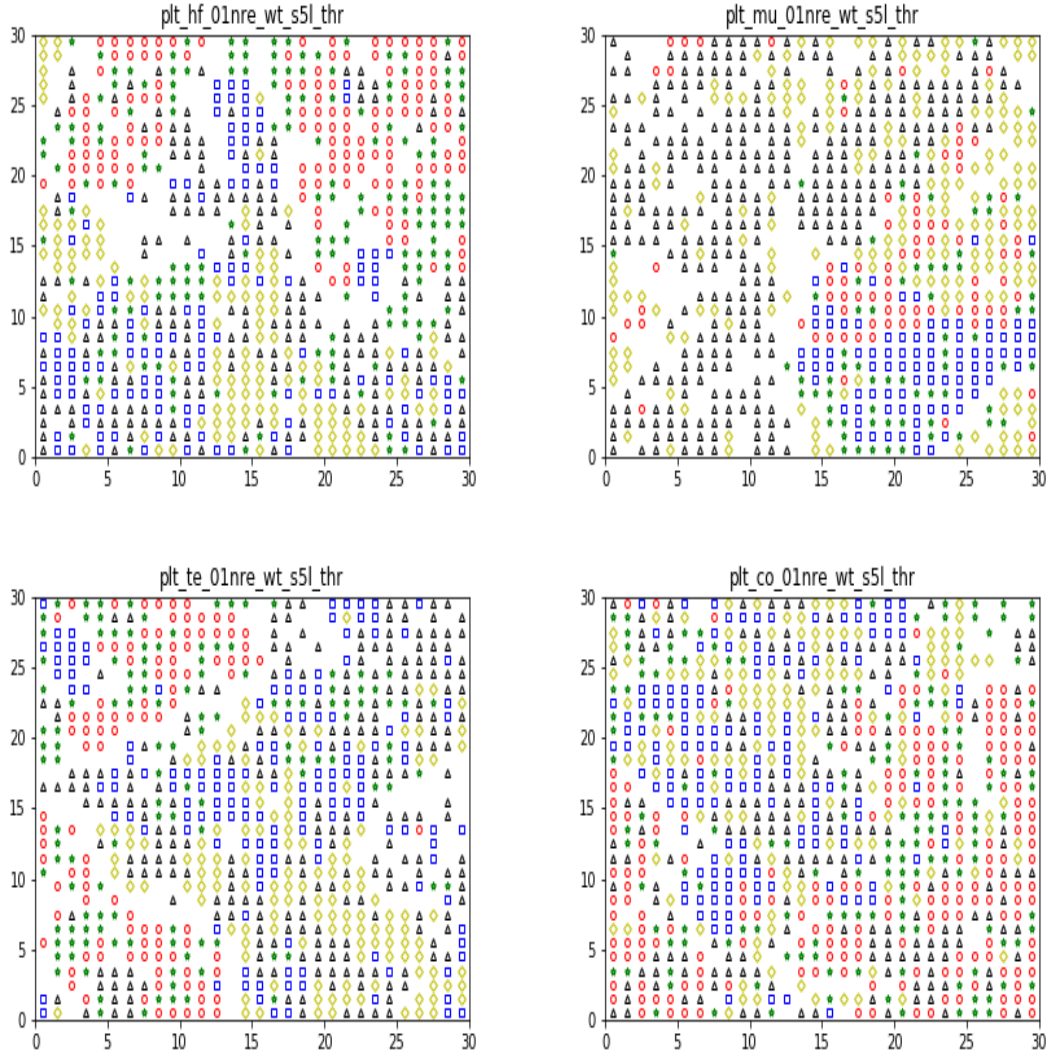


Figure 4.7: Examining splitting the dataset of driver 1 into five segments
 (a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
 (c) Text distraction (bottom-left) and (d) combination of datasets
 (bottom-right).

The continuing trend while labeling the map after dividing the data into sections is that the music distraction is more active towards the end of the drive (the fourth and fifth sections) while the other distractions are more active at the starting sections

(the first and second sections). Results by experiments on dividing each dataset into two to four segments can be seen in Appendix A.

Label 8 Distraction: Labeling only the model containing all the three distractions so as to differentiate between them. The map is labeled by the “Hands-Free” distraction (red), the “Music” distraction (green) and the “Text” distraction (blue) as shown in figure 4.8.

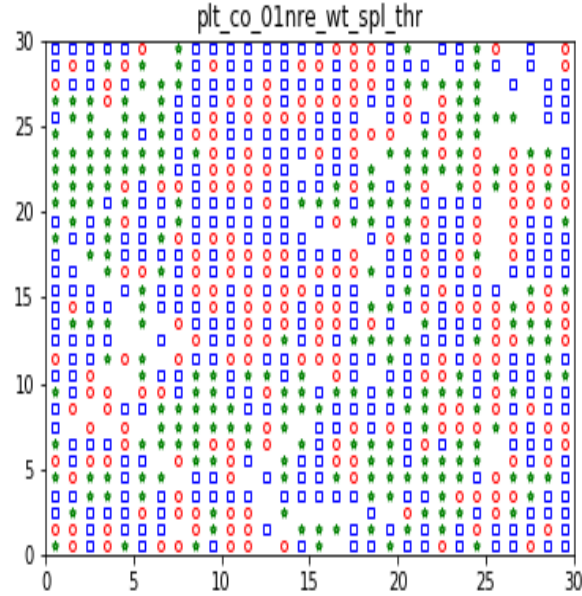


Figure 4.8: Examining how each of the distractions are visualized on the map

.

From labeling the each distraction, it is found that the music distraction contains larger sections of activated portions on the map, showing that driving activity while listening to music is more consistent than the same activity while being distracted by the hands-free and text distractions.

Results Summary:

- It was found that the six features based on which the models are trained, have been learned by the model. However this is completed with differing strengths which indicates that some features affect the resulting topological structure more than others which is a result of the overall nature of the set of actions that comprise the process of driving. For example the tangential speed feature takes up larger sections as opposed to acceleration along Z axis.
- Certain features are shown to differ in terms of detecting different distractions. For example the lane gap feature and the sections label which divides the dataset into parts was found to differ to a larger extent under music distractions versus the other two distractions.
- As this study is aimed towards understanding differences between distractions and most of the features do not highlight such differences, further tests are performed using multiple drivers to distinguish between both drivers and the distractions.

4.2 Result analysis of four drivers

The models used in this section are trained using four drivers (the first four drivers in the dataset) as a test to understand how a model would be trained on multiple drivers before finally applying and analysing all the drivers in the dataset. The models trained in this section are as follows:

1. **Hands-Free distraction drivers 1 to 4:** This model is trained on the “Hands-Free” distraction of drivers 1 to 4 and the number of output neurons is 1600 neurons over a 40x40 size map.

2. **Music distraction drivers 1 to 4:** This model is trained on the “Music” distraction of drivers 1 to 4 and the number of output neurons is 1600 neurons over a 40x40 size map.
3. **Text distraction drivers 1 to 4:** This model is trained on the “Text” distraction of drivers 1 to 4 and the number of output neurons is 1600 neurons over a 40x40 size map.
4. **Combination datasets of drivers 1 to 4:** This model is trained on the “Hands-Free” distraction of drivers 1 to 4 and the number of output neurons is 2500 neurons over a 50x50 size map.

Most of the results obtained were found to be similar to those obtained for the single driver in section 4.1. Hence a brief overview of the results using only the model containing all the distractions of drivers one to four is given below followed by labeling each driver before finally leading into models that are trained on all the datasets.

Feature based labels (labels 1 to 6 in section 4.1):

The very same labels in section 4.1 which are – brake pressure, speed limits, linear acceleration, turning acceleration, altitude acceleration and finally the gap between lanes as shown in figure 4.9.

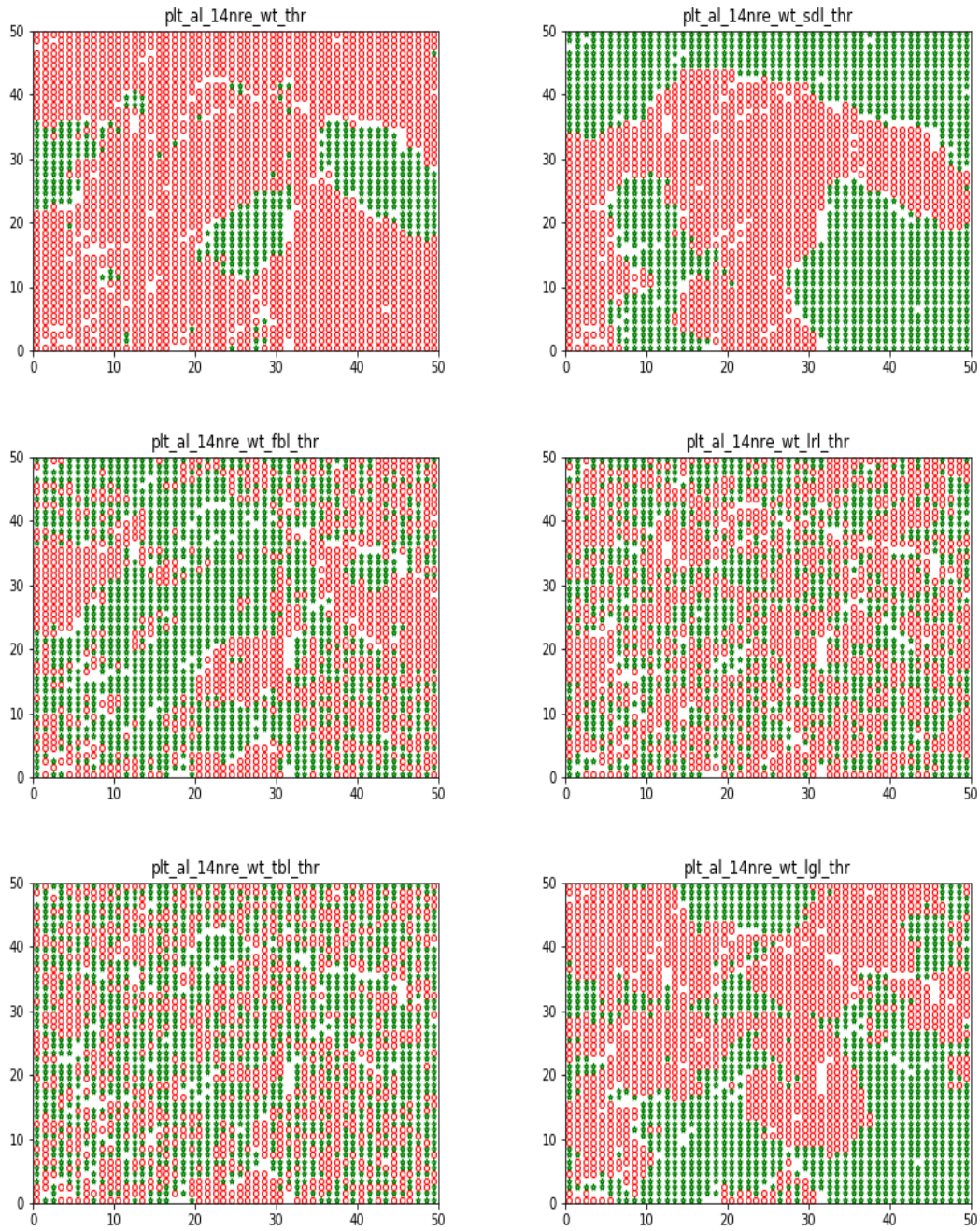


Figure 4.9: Examining the six feature labels of drivers 1 to 4 using datasets of all three distractions
(a) Brake pressure (top-left), (b) Speed limits (top-right), (c) Linear acceleration (middle-left), (d) Turning acceleration (middle-right), (e) Altitude acceleration (bottom-left) and (f) Gap between lanes (bottom-right).

As stated previously the results of drivers 1 to 4 including datasets of all distractions, after applying the feature labels is found to be similar to the results discussed from dealing with a single driver in section 4.2.

Section split labels (label 7 in section 4.1):

The same labels used to separate the dataset into sections is applied on the model trained on the four drivers with datasets of all three distractions with labels similar to those mentioned in section 4.1, where part one is labeled red, part two is labeled green, part three is labeled blue, part four is labeled yellow and part five is labeled black as shown in figure 4.10.

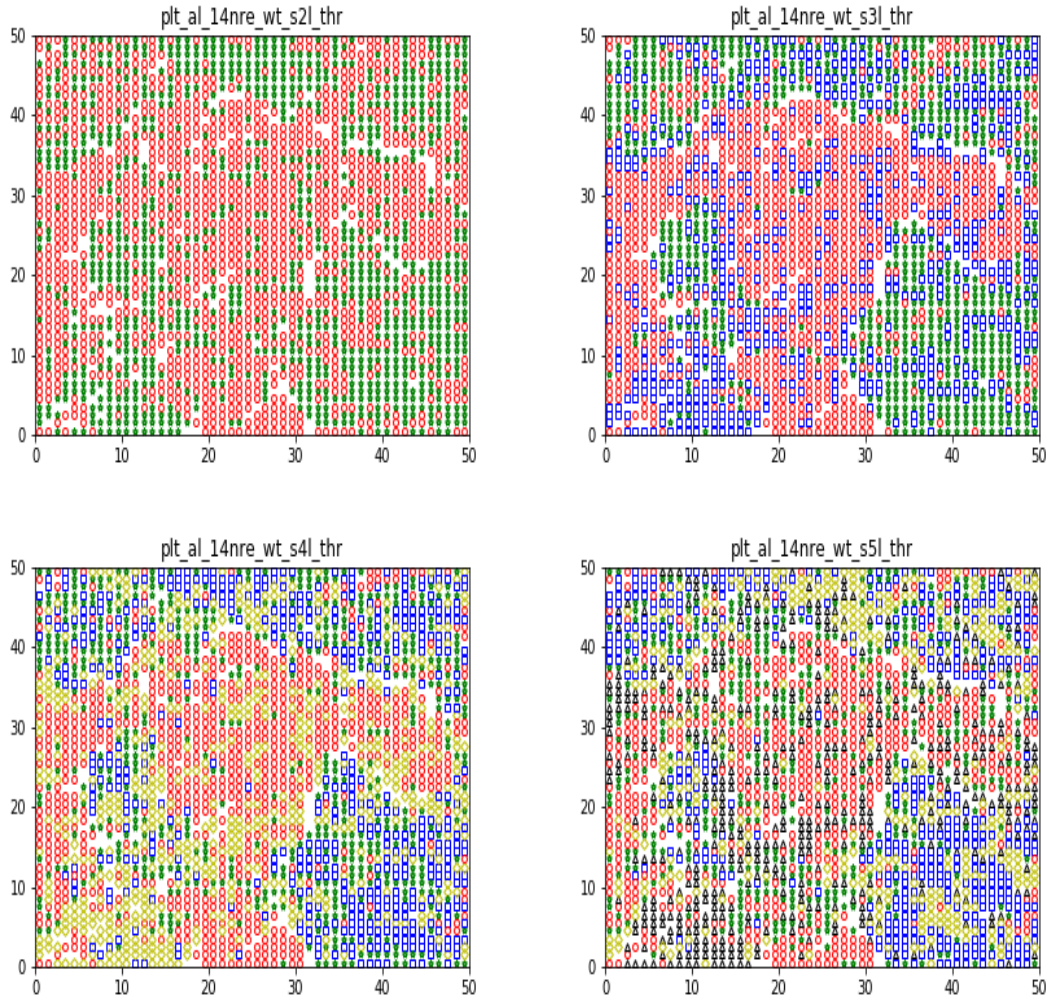


Figure 4.10: Examining splitting the dataset of drivers 1 to 4 using datasets of all three distractions by dividing into (a) Two sections (top-left), (b) Three sections (top-right), (c) Four sections (bottom-left) and (d) five sections (bottom-right).

Once again the results obtained here are very similar to those found in section 4.1. The common trend is maintained in that the music distraction is more active and prevalent towards the last sections as opposed to the hands-free and text distractions which are more prevalent in the first few sections.

Labeling each distraction:

The model that contains all the distractions and datasets of the four drivers is labeled to understand each distraction. The labels used are similar to those used in section 4.1 where the hands-free datasets are labeled red, music datasets are labeled green and finally the text datasets are labeled blue as seen in figure 4.11.

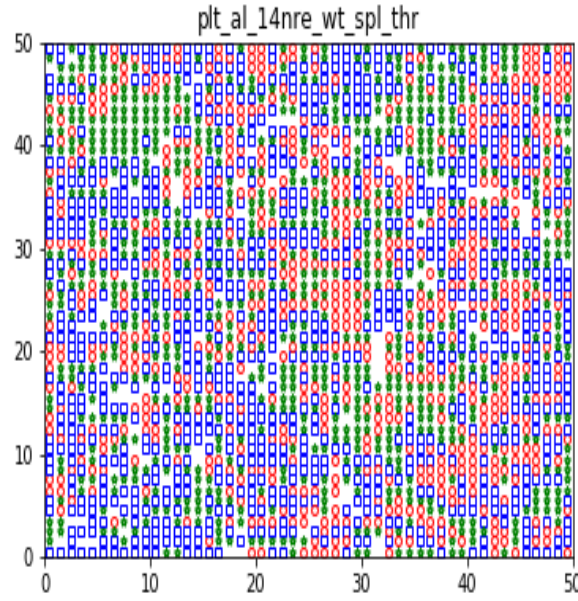


Figure 4.11: Examining each distraction in the model trained on the first four drivers and on all the three distractions.

Similar to the results in section 4.1, the music distraction has larger sections on the map, but in this case, it appears that listening to music and texting create the most unique behaviours in the map.

Labeling each driver:

The four drivers in the dataset that the model is trained on are labeled such that driver one is labeled with red, driver two is labeled with green, driver three is labeled

with blue and driver four is labeled with yellow as shown in figure 4.12.

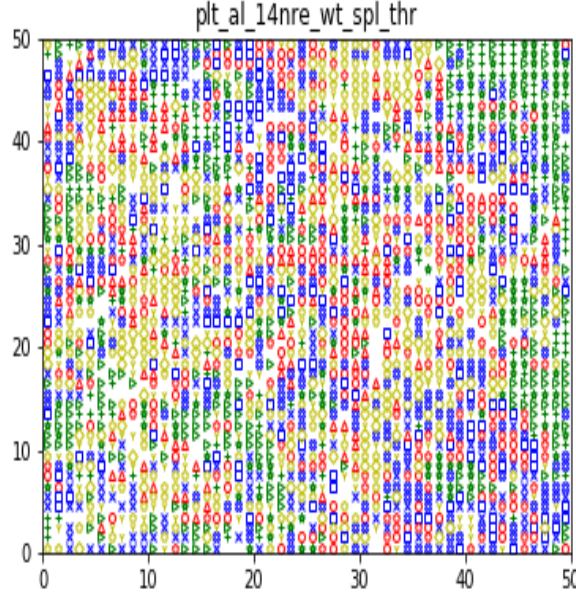


Figure 4.12: Examining each individual driver in the model trained on the first four drivers and on all the three distractions.

There are patterns that show each drivers' individuality from each other, such as the patterns of driver two (green) and driver four (yellow), while the others appear too mixed up to make any conclusion. A more drivers are added, some show particular patterns but others do not. How this is made use of will be described in detail in the next section where all the models used will be trained on the datasets of all the 40 drivers.

4.3 Result analysis of 40 drivers:

Finally the models trained in this section incorporate all the datasets of the 40 drivers including all the separate datasets on the three distractions for a total of 120

datasets. Individual driver patterns are labeled after applying the maximal set of points of each driver in comparison to all the other datasets. These maximal points are labeled using the existing labels but now each individual label is broken down into its specific classes. After this, percentages of how many points in the maximal set belong to each class are calculated and they serve as measures to distinguish between each unique driver pattern. The models trained on all 40 drivers are as follows:

1. **Hands-Free distraction drivers 1 to 40:** This model is trained on the “Hands-Free” distraction of drivers 1 to 40 and the number of output neurons is 3600 neurons over a 60x60 size map.
2. **Music distraction drivers 1 to 40:** This model is trained on the “Music” distraction of drivers 1 to 40 and the number of output neurons is 3600 neurons over a 60x60 size map.
3. **Text distraction drivers 1 to 40:** This model is trained on the “Text” distraction of drivers 1 to 40 and the number of output neurons is 3600 neurons over a 60x60 size map.
4. **Combination dataset of drivers 1 to 40:** This model is trained on the datasets of all three distractions of drivers 1 to 40 and the number of output neurons is 6400 neurons over a 60x60 size map.

While studying the models trained under individual distractions can lead to insight into unique driver patterns, it is difficult to compare the distractions that are trained on separate maps. When all distractions were trained on the same map, it is possible to make conclusions about both individual drivers and differences in distractions. For this reason only the results of the fourth model, trained on all the 40 drivers

using data of all the three distractions is examined. The same procedure in section 4.2 is repeated where the labeled results are discussed before getting into the problems encountered in the newly trained model on all 40 drivers

Feature based labels (labels 1 to 6 in section 4.1):

The very same labels in section 4.1 which are – brake pressure, speed limits, linear acceleration, turning acceleration, altitude acceleration and finally the gap between lanes as shown in figure 4.13.

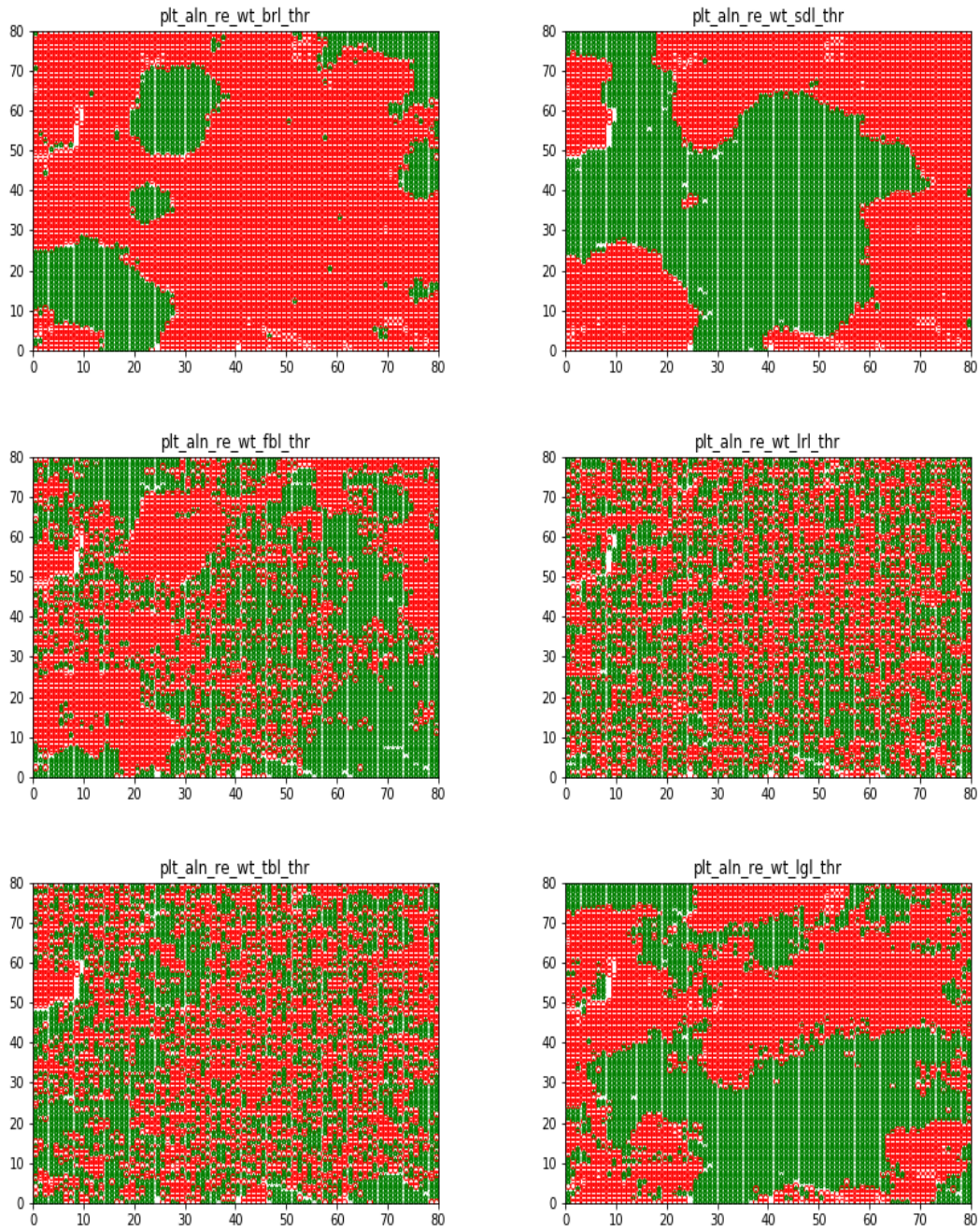


Figure 4.13: Examining the six feature labels of drivers 1 to 40 using datasets of all three distractions
 (a) Brake pressure (top-left), (b) Speed limits (top-right), (c) Linear acceleration (middle-left), (d) Turning acceleration (middle-right), (e) Altitude acceleration (bottom-left) and (f) Gap between lanes (bottom-right).

Some interesting points to note are that certain features such as the three accelerations along the X, Y and Z axes, slowly start to spread out when compared with the previous maps on a single driver and the four drivers, thus resulting in smaller sections on the map. This is once again due to the differences between each driver in all the datasets which will be further examined while labeling the drivers and distractions.

Section split labels (label 7 in section 4.1):

Once again the very same labels used to separate the dataset into sections is applied on the model trained on all the forty drivers with datasets of all three distractions with labels similar to those mentioned in section 4.1, where part one is labeled red, part two is labeled green, part three is labeled blue, part four is labeled yellow and part five is labeled black as seen in figure 4.14.

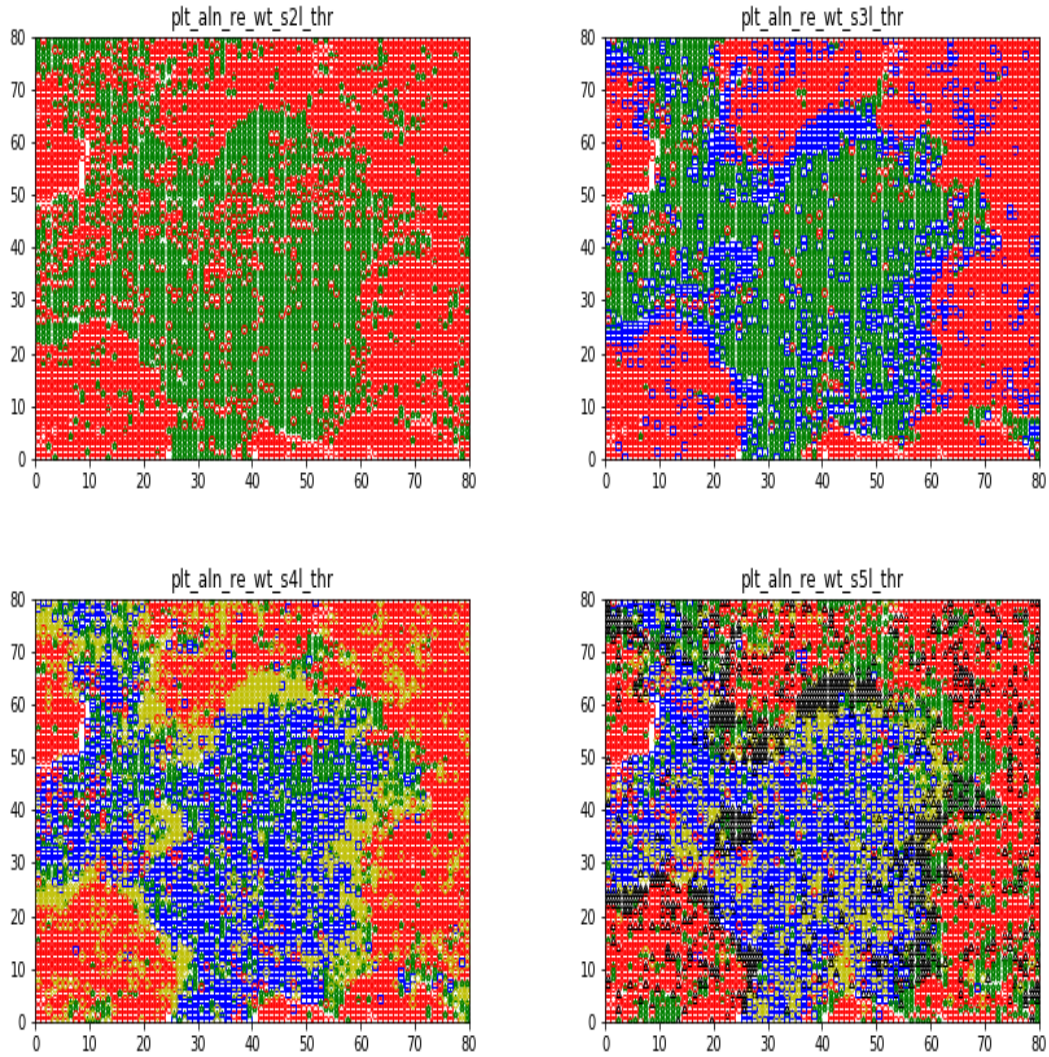


Figure 4.14: Examining splitting the dataset of drivers 1 to 40 under all three distractions by dividing into
 (a) Two sections (top-left), (b) Three sections (top-right), (c) Four sections (bottom-left) and (d) five sections (bottom-right).

Labeling each distraction:

The model that contains all the distractions and datasets of all the 40 drivers is labeled to understand each distraction. The labels used are similar to those used in section 4.1 where the hands-free datasets are labeled red, music datasets are labeled

green and finally the text datasets are labeled blue as seen in figure 4.15.

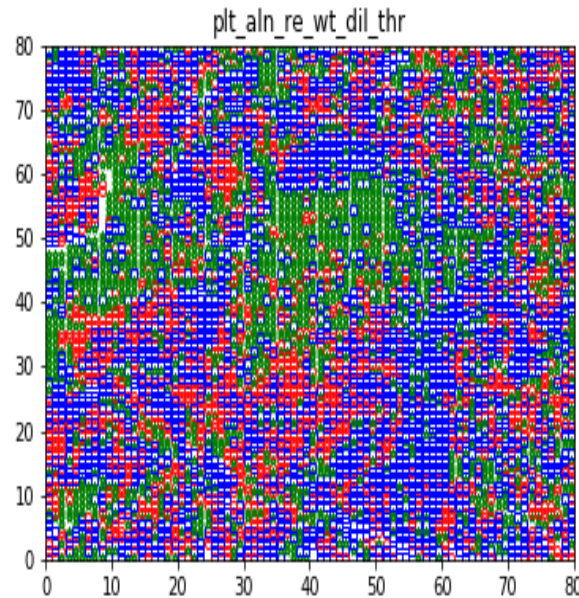


Figure 4.15: Examining each distraction in the model trained on data of all the 40 drivers and on all the three distractions.

In the final case of comparing distractions of all the drivers on a single model, the results show that the music distraction remains almost constant by having its own section but there are also scattered smaller sections. Interestingly, the patterns from the text distraction were also found to gain larger sections of the map, denoting common characteristics with drivers listening to music and drivers texting while driving.

Labeling each driver:

The 40 drivers in the dataset that the model is trained on, is labeled such that driver one is labeled with 40 unique colors to differentiate between them as seen in figure 4.16.

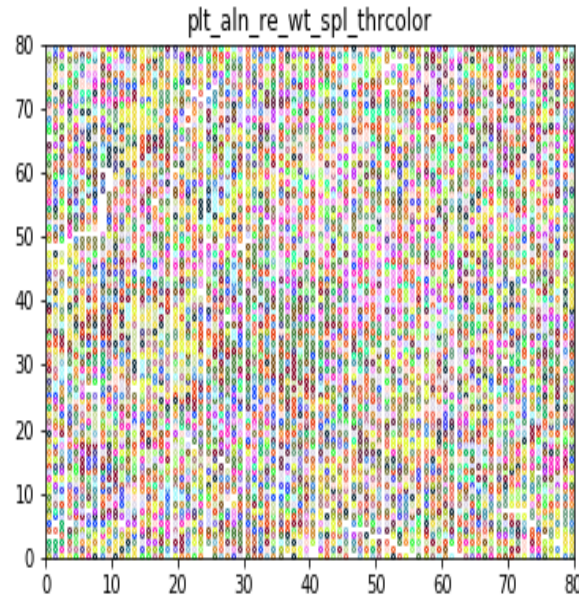


Figure 4.16: Examining each individual driver in the model trained on data of all 40 drivers and on all the three distractions.

No particular driver has their data clearly grouped together on the map. Although each point is a cluster of common characteristics, this does not seem to extend to other points nearby and so there is no clustering of individual drivers. The topological structure of the SOM forces patterns to be different from each other, so when patterns are found near each other as they indicate similarities between drivers.

Training an Augmented Model:

Numerous efforts were taken to include labels for individual drivers during training to gain better understanding of how each driver differs from each other. An example of this is mentioned in section 3.5 where the label of each driver is added to the six features of the dataset during training as the new seventh feature. This is done with the aim of bringing together clusters of drivers. These features were then trained

on a new model (all the parameters were maintained) where each driver was then labeled with each feature. The model had seven input neurons and the output layer had 900 neurons which is a 30x30 map. Data was randomly ordered during training which resulted in drivers, to an extent, being grouped together. This in turn led to certain other problems which will be described after viewing the results in figure 4.17.

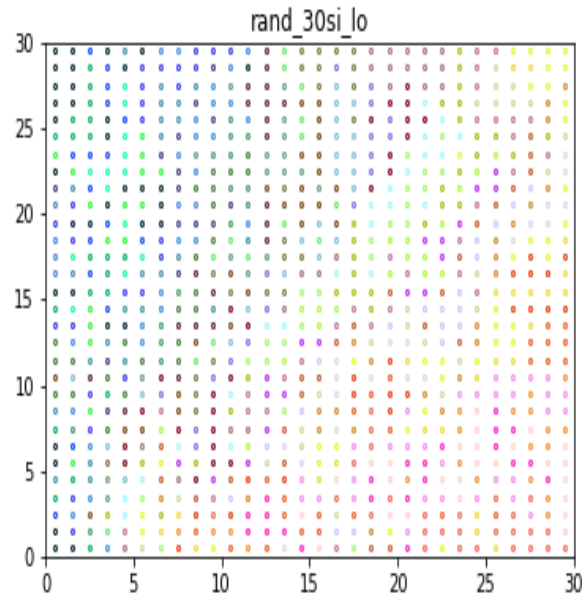


Figure 4.17: Examining each individual driver in the model trained on data of all 40 drivers, but augmented with the driver label.

Looking at the above result, the first assumption would be that it divides drivers in the map and causes individual drivers to be grouped. However this contains certain inherent issues that must be addressed and is the reason why the model is discarded from future analysis.

The first issue is in the way of labeling which adds the new seventh feature, the

driver ID, in turn changes the behaviour of the map. When labeled between 0 and 39 where each driver gets a label, driver 1 is labeled as 0 and driver 40 is labeled as 39. This makes the model learn that according to the new feature driver 1 is distinct from driver 40 which creates the structure that can be seen in the map where the labels move across the map in a diagonal pattern. Another issue is that this new feature overwhelms the patterns found by using the initial six features and replaces existing patterns that could lead to potentially interesting results. Due to these reasons this model has not been used and the idea of grouping drivers in this manner has been discarded in favor of the results from the other features in the data.

Results Summary:

- Results of labeling with the feature analysis and section split labels are very similar to those previously discussed in sections 4.1 and 4.2. The main difference is that as the number of datasets increases, some features in the models become more scattered over the map while others remaining constant, thus showing the changing patterns when more drivers are added to create the model.
- Attempts to make individual driver groups to form on the map despite each point already being a cluster of vectors of that particular driver were not successful because adding the driver label as a new feature resulted in patterns that were too heavily influenced by the newly added driver label. This in turn caused the loss of the other patterns formed by the initial six features and hence lost the individuality of each driver.

4.4 Thresholding and label to classifier

Due to the complexity of the previous maps, thresholding of the nodes was used to reduce the number of active nodes on the map to make it easier to understand. This map is trained on all 40 drivers with all the distractions. A threshold is applied to each node which selects the driver which has the largest number of activations in the node. The resulting set of neurons is the distinct points for each driver under the respective distractions.

Examples of applying the threshold:

Examples of this threshold are shown with respect to driver 1 under hands-free distraction in figure 4.18, under music distraction in figure 4.19 and under text distraction in figure 4.20 which are all shown below:

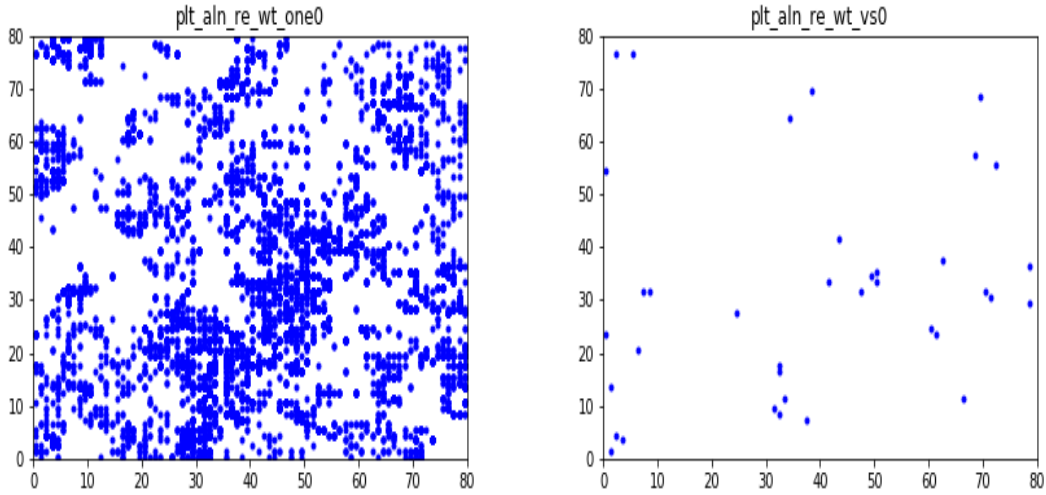


Figure 4.18: Examining how the threshold over the set of activated neurons of driver 1 dataset under the hands-free distraction works (a) All activated neurons versus (left) (b) maximally activated neurons after applying the threshold (right).

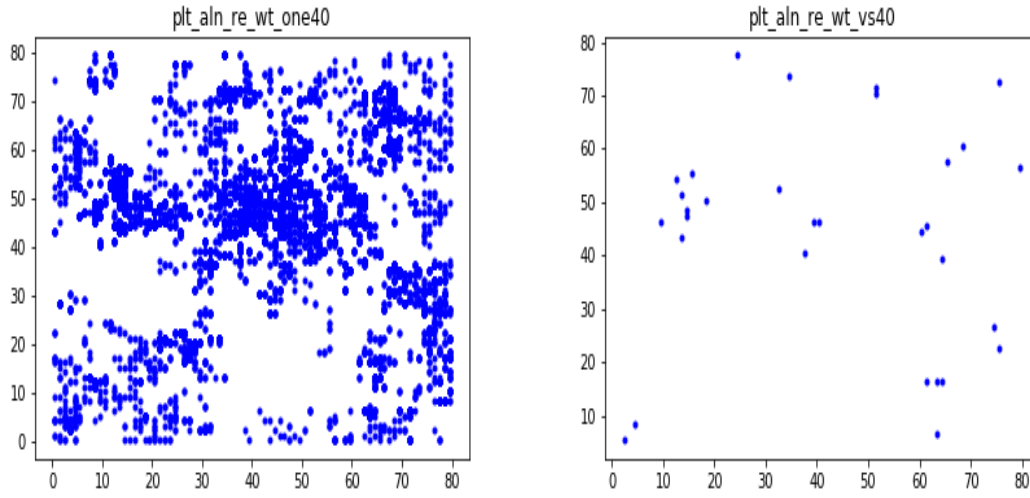


Figure 4.19: Examining how the threshold over the set of activated neurons of driver 1 dataset under the music distraction works
(a) All activated neurons versus (left) (b) maximally activated neurons after applying the threshold (right).

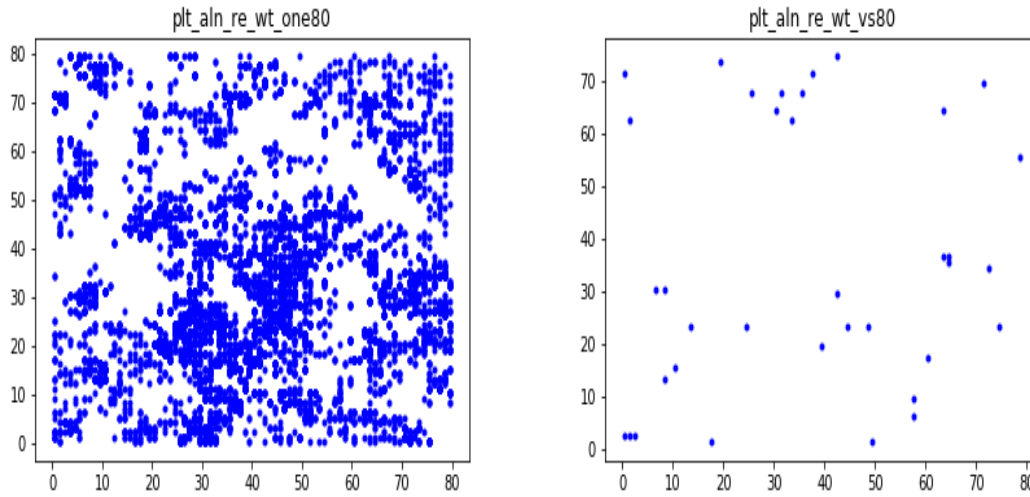


Figure 4.20: Examining how the threshold over the set of activated neurons of driver 1 dataset under the text distraction works
(a) All activated neurons versus (b) maximally activated neurons after applying the threshold.

In this way two factors are achieved, where the first being a very significant decrease in data points leading to results that uniquely distinguish each driver. The second and more important factor is that this reduced set of points means that they are only the most distinct points for their respective dataset and this leads to the individual driver patterns being explicitly understood by the model.

An important point is that the maximal points obtained do not show all the characteristics of that particular driver but only the most distinctive traits of each driver in comparison to other drivers. It was also found that certain sets of maximal points contain points that are activated for different drivers (sets overlap), meaning that the same neurons tend to be activated by different drivers. Although the threshold selects the driver who has the most activations of a particular node, it is possible for there to be two drivers with the same maximum. This is used as a way of representing drivers that share common traits.

Feature label analysis after thresholding:

The same labels based on the features that were used in sections 4.1, 4.2 and 4.3 are now applied to the maximal set for points of each driver obtained after applying the threshold. In this case the label will now be used as a classifier. Each maximal set is labeled, after which percentages are calculated that measure how much each driver (under one distraction at a time) belongs to one of the classes within that particular label.

The labels applied remain the same and the only difference is that after thresholding, there remain fewer points in the set. For example, consider the brake pressure label. It contains two types of labels – when the brakes are applied and when they are not applied. These two will now serve as classes and the percentages of how many maximal points fall under each category is calculated.

Individual feature label analysis

Label 1 Brake Analysis: Labeling the models mentioned above by when the driver applies the brakes (green) versus when the brakes are not applied (red) as shown in figure 4.21. This is based on the brake pressure feature in the dataset.

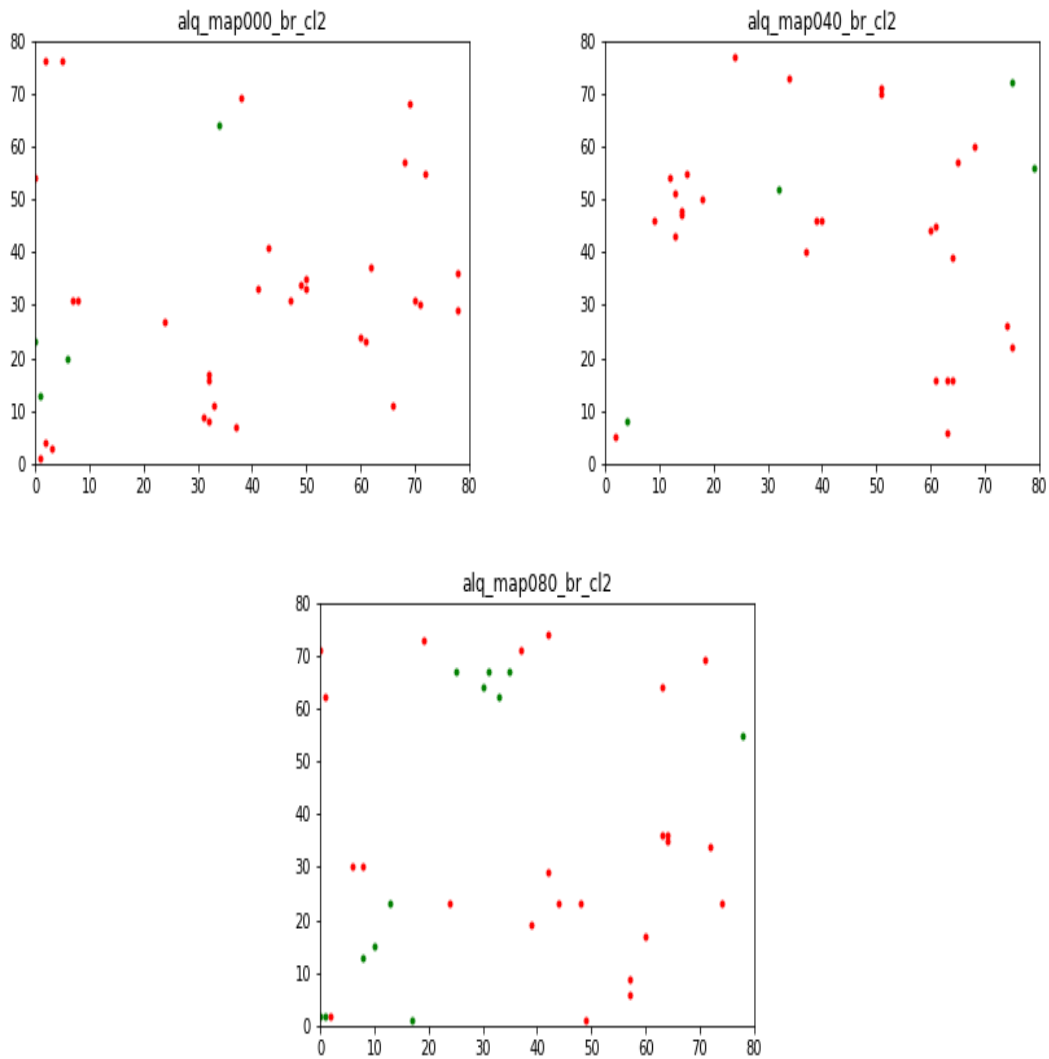


Figure 4.21: Examining applying the brake pressure label over the maximal points set of driver 1
(a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 2 Speed Limits: Labeling the models mentioned above when the driver speeds up and goes above 77 km/hr which is above the average speed (green) versus when the driver slows down going below the same average velocity (red) as shown in figure 4.22. This is based on the tangential speed feature in the dataset.

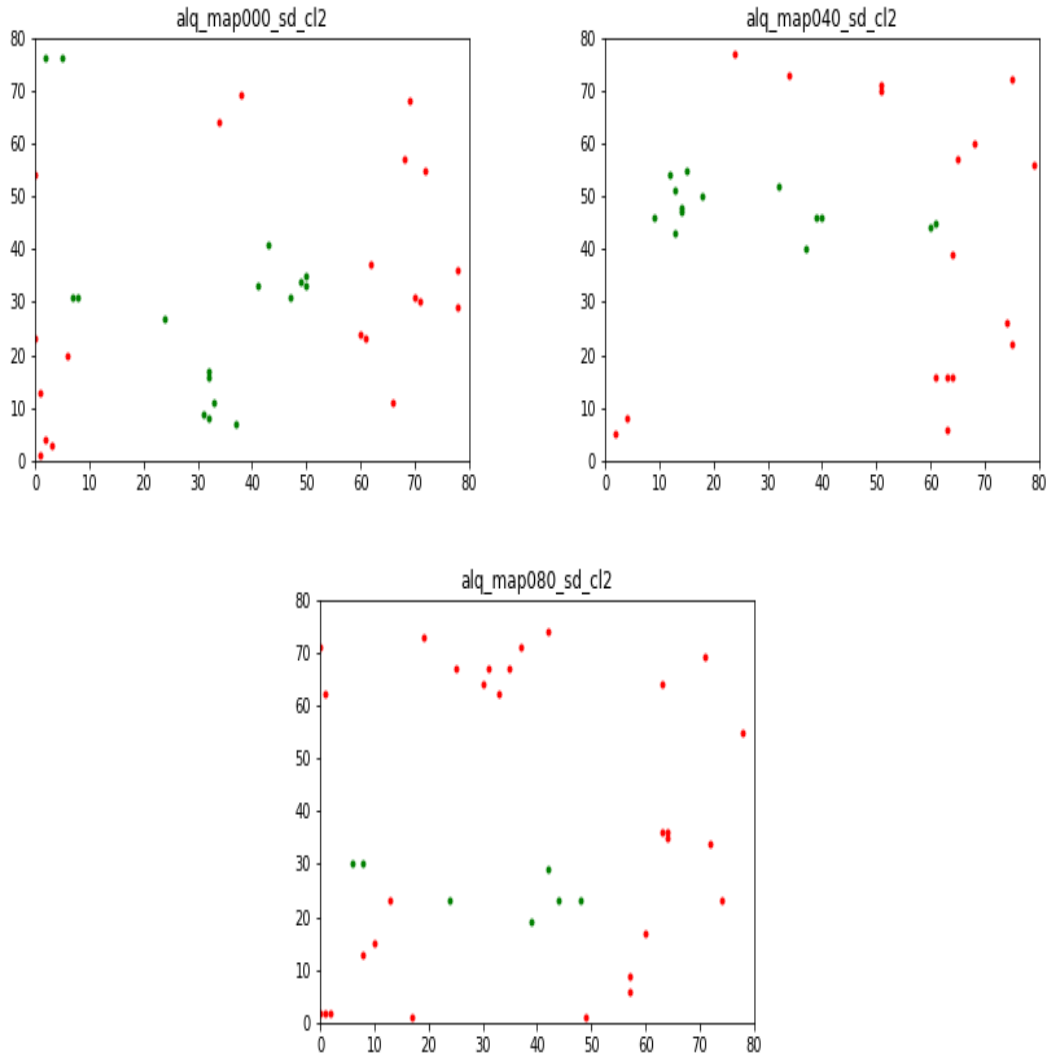


Figure 4.22: Examining applying the speed limits label over the maximal points set of driver 1
 (a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 3 Linear Acceleration: Labeling the models mentioned above based on when the driver accelerates (positive values) (green) versus when the driver decelerates (negative values) (red) and this is shown in figure 4.23. This is based on the acceleration along the X axis feature in the dataset.

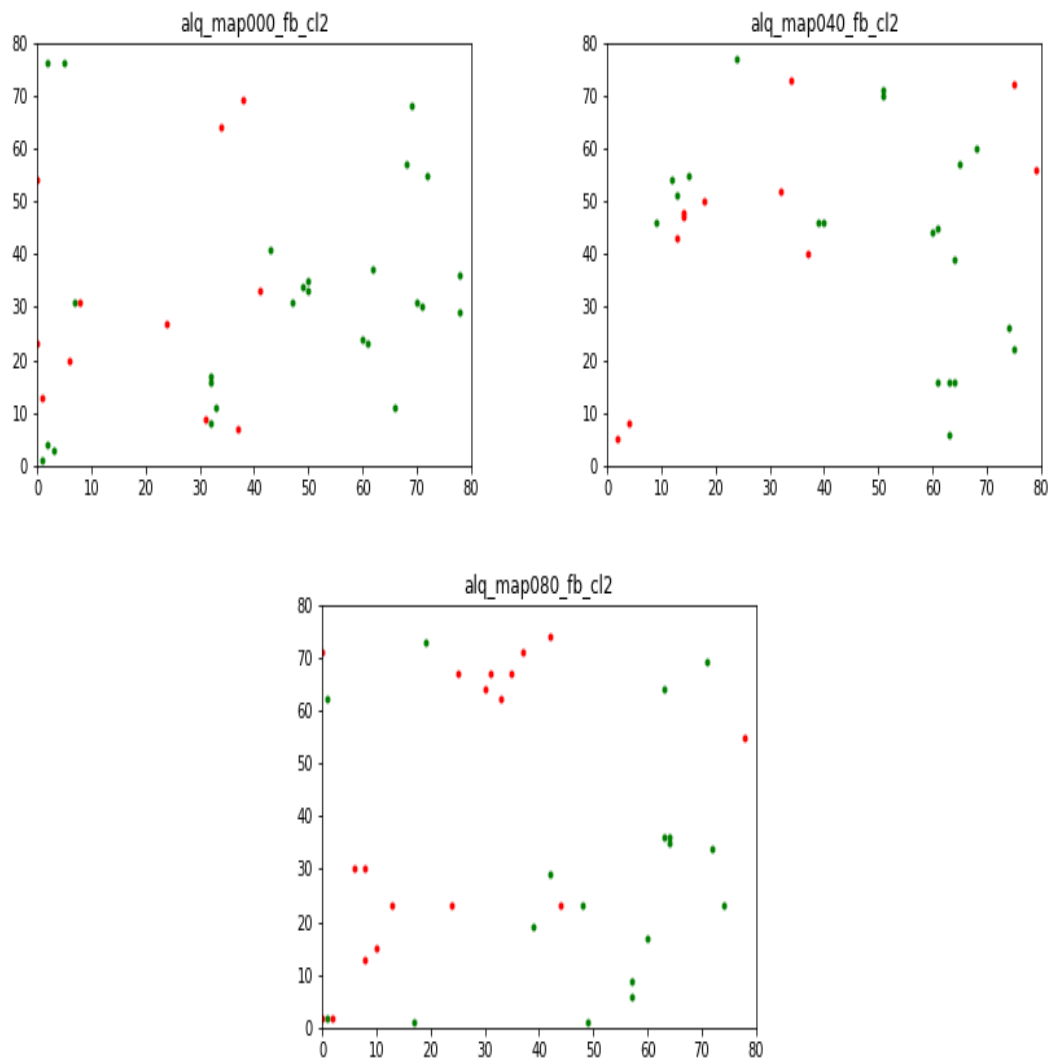


Figure 4.23: Examining applying the linear acceleration label over the maximal points set of driver 1
 (a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 4 Turning Acceleration: Labeling the models mentioned above based on when the driver accelerates either to the right (positive values) (green) or to the left (negative values) (red) and this is shown in figure 4.24. This is based on the acceleration along Y axis feature and can even help analyse when the driver turns to the left or right.

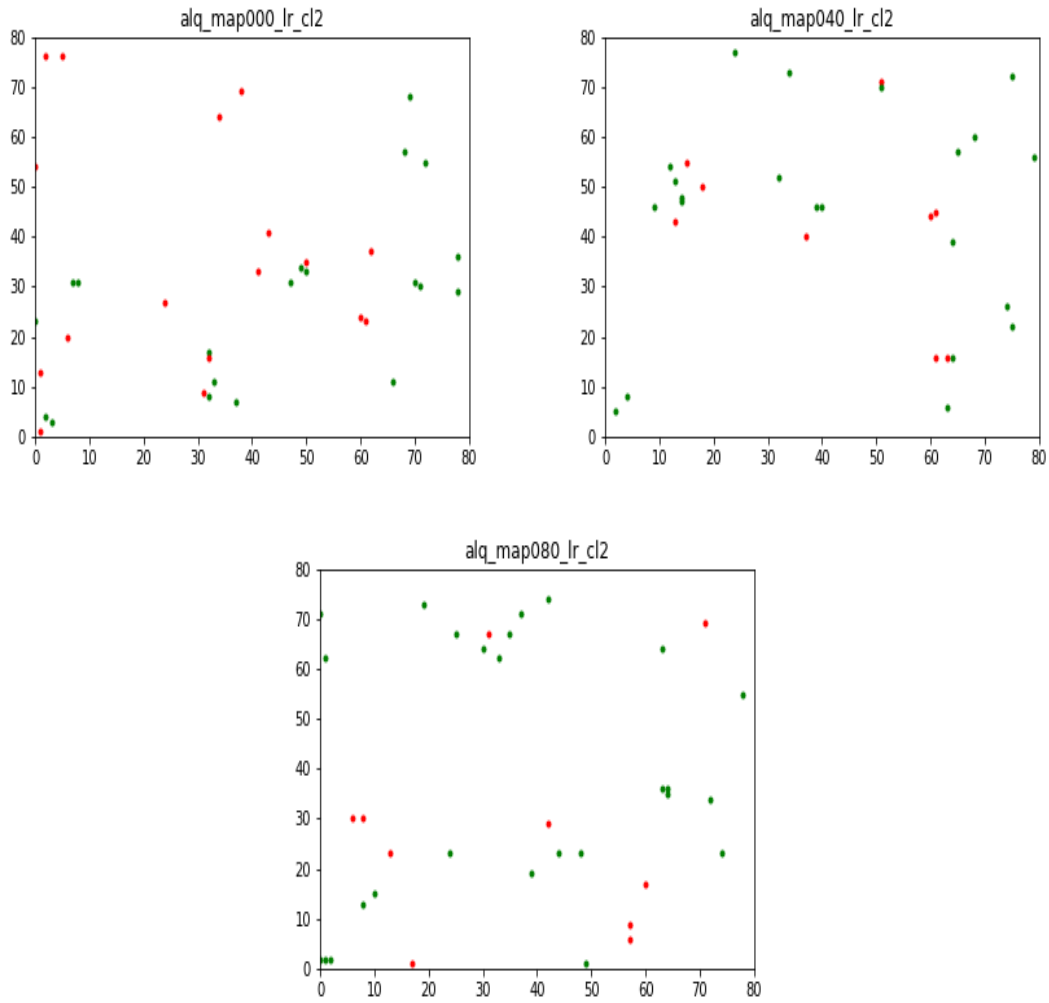


Figure 4.24: Examining applying the turning acceleration label over the maximal points set of driver 1
(a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 5 Altitude Acceleration: Labeling the models mentioned above based on when the driver accelerates over a higher altitude (positive values) (green) versus when the driver accelerates over a comparatively lower altitude (negative values) (red) as shown in figure 4.25. This is based on the acceleration along Z axis feature in the dataset.

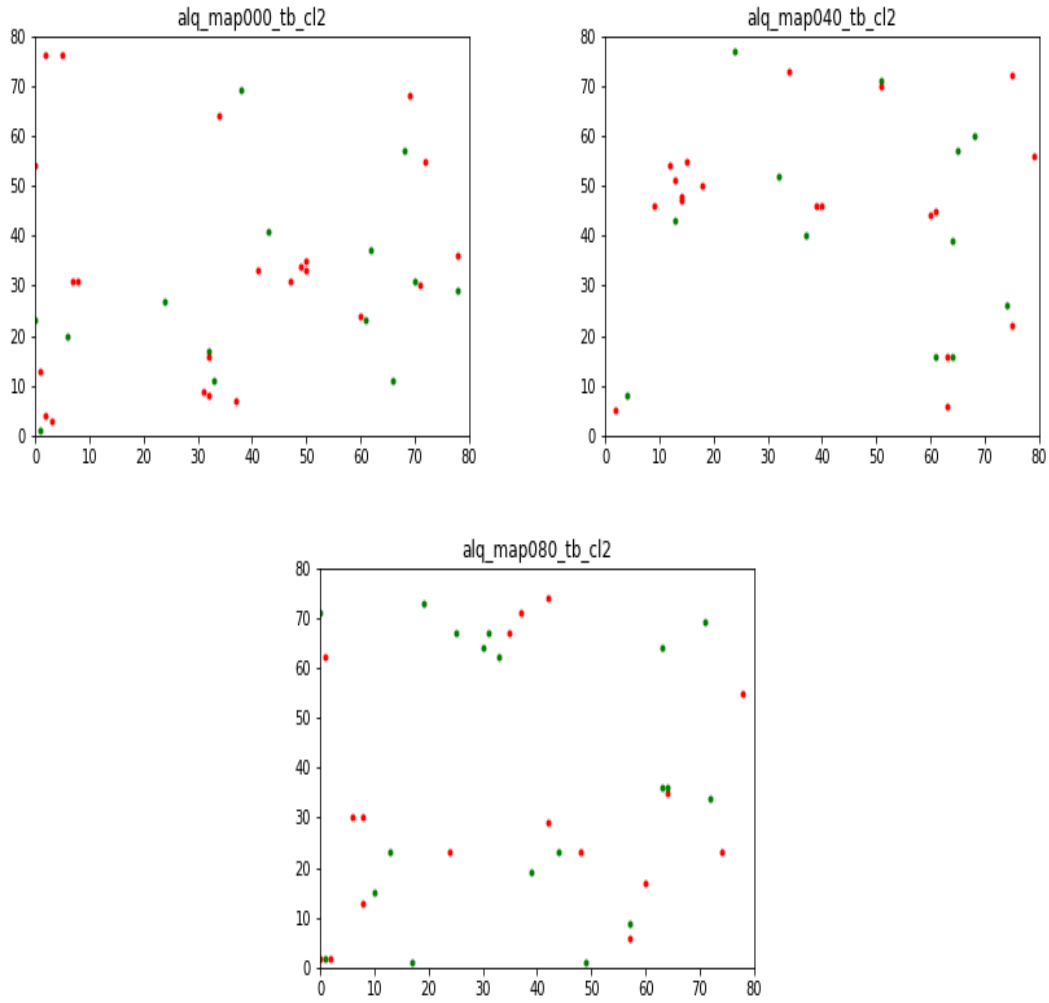


Figure 4.25: Examining applying the altitude acceleration label over the maximal points set of driver 1
 (a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 6 Gap between lanes: Labeling models mentioned above based on when the driver tends to drive to the right of the center of the lane (positive values) (green) versus when driving to the left side of the lane (negative values) (red) as shown in figure 4.26. This is based on the lane gap feature in the dataset.

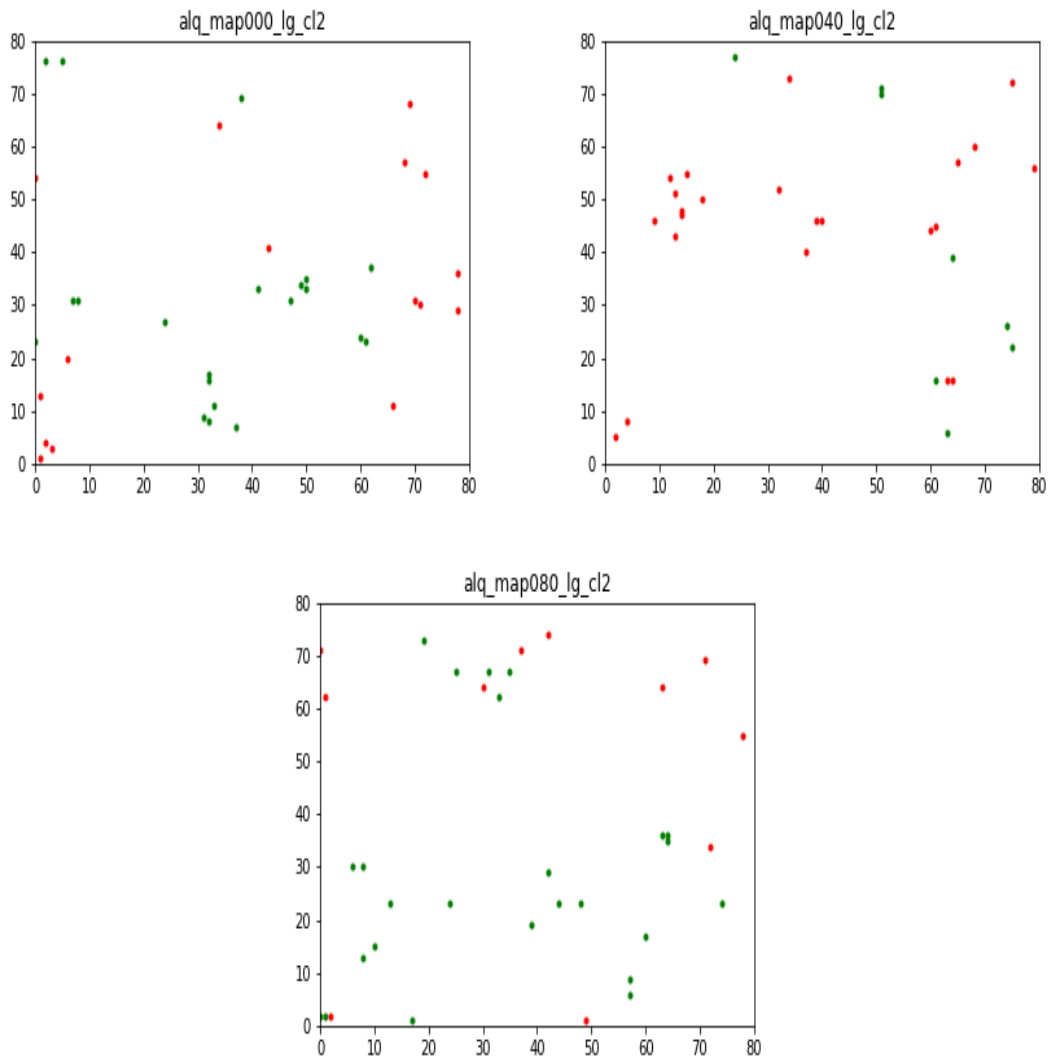


Figure 4.26: Examining applying the gap between lanes - label over the maximal points set of driver 1
 (a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 7 Split Sections: Labeling models mentioned above to understand the sequential flow of how each model adapts to the dataset. This is done by dividing the dataset into sections and labeling each sections on each model to understand when most actions are performed and how often the neurons are activated.

Label 7a: Dividing each dataset into two sections – the first section (red) and the second section (green) as seen in figure 4.27.

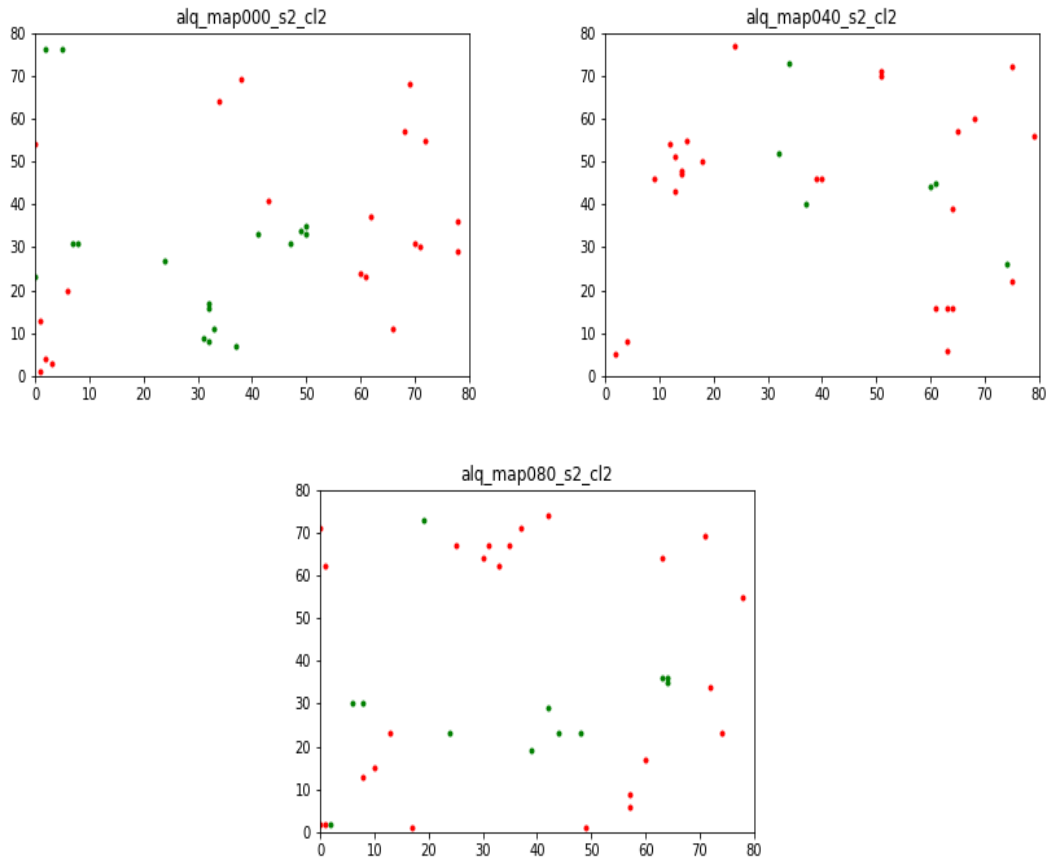


Figure 4.27: Examining applying the split into two sections - label over the maximal points set of driver 1

(a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 7b: Dividing the dataset into three sections – the first section (red), the second section (green) and the last section (blue) as shown in figure 4.28.

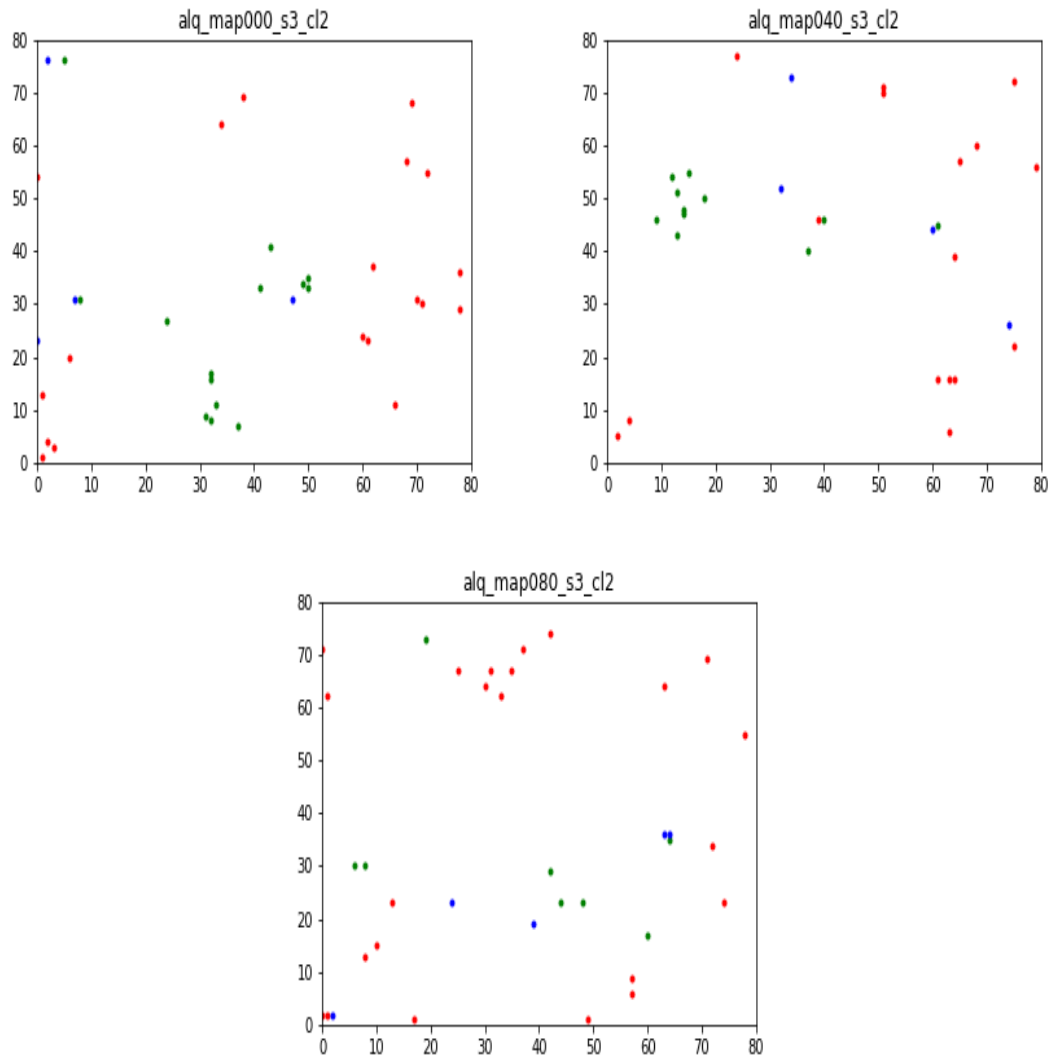


Figure 4.28: Examining applying the split into three segments - label over the maximal points set of driver 1
(a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 7c: Dividing the dataset into four sections – the first section (red), the second section (green), the third section (blue) and the last section (yellow) as seen in figure 4.29.

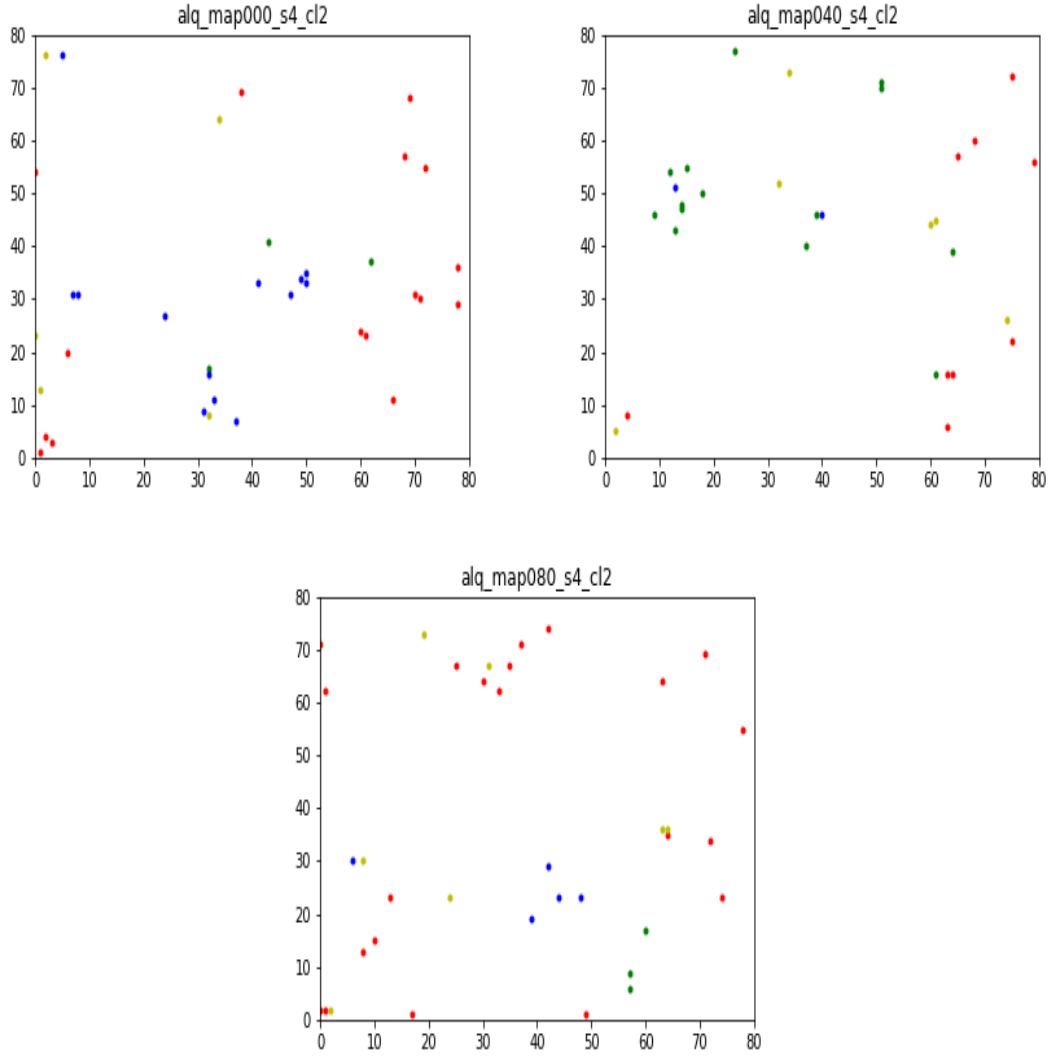


Figure 4.29: Examining applying the split into four segments - label over the maximal points set of driver 1
 (a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Label 7d: Dividing the dataset into five sections – the first section (red), the second section (green), the third section (blue), the fourth section (yellow) and the last section (black) as shown in figure 4.30.

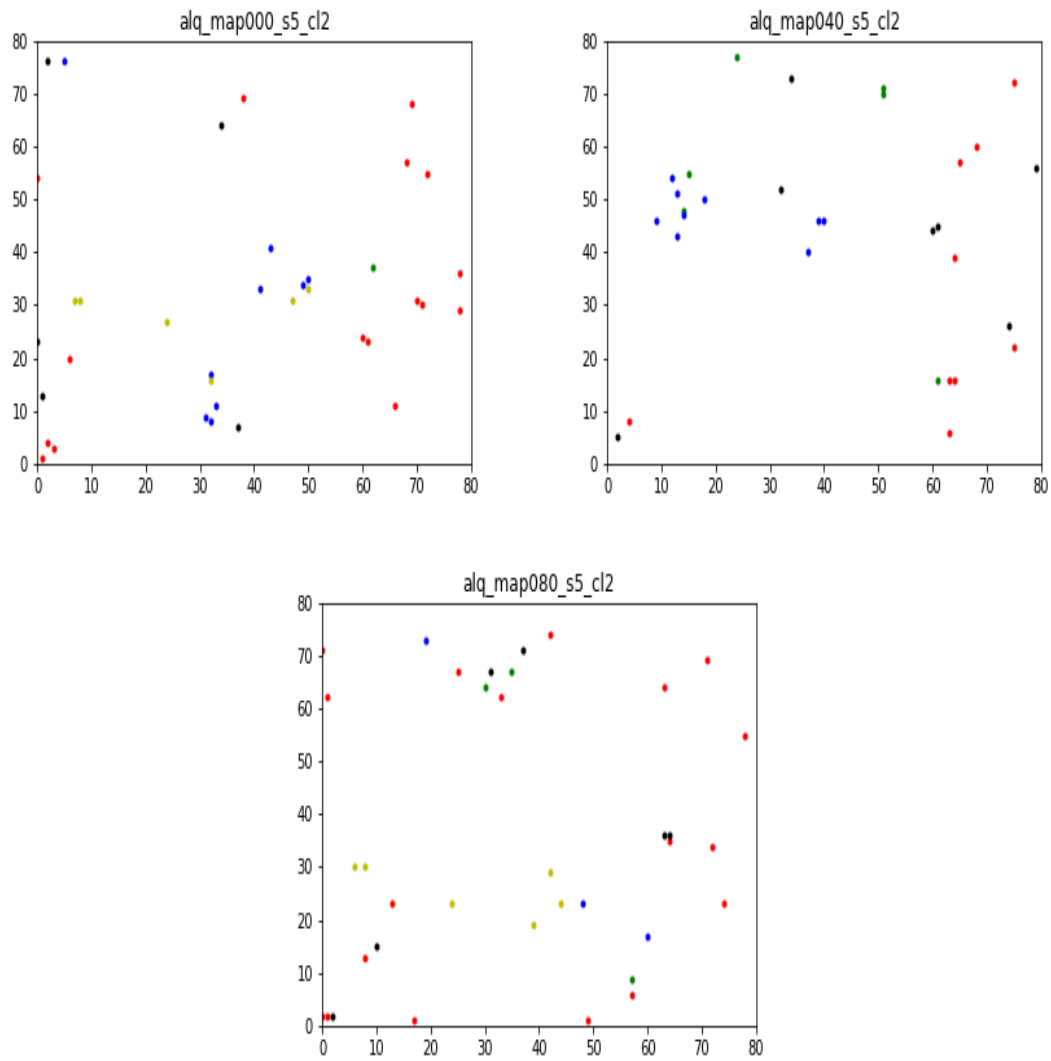


Figure 4.30: Examining applying the split into fives sections - label over the maximal points set of driver 1
(a) hands-free distraction (top-left), (b) music distraction (top-right) and (c) text distraction (bottom).

Calculating class percentages:

After all of the maximal set of points for all drivers and distractions are labeled, percentages are calculated as a measure of each class within a single label, and this is repeated for every label and every dataset. These percentages represent how each driver differs from their counterparts under that specific action (class within the label).

For example, consider a driver under music distraction, whose maximal points are computed and then labeled with the brake pressure label which results in having 85% towards when the brake is not applied and 15% towards when the brake is applied. This same driver when under the text distraction may tend to result in 70% towards when the brake is not applied and 30% towards when the brake is applied. From this the conclusion is made that the “action of applying the brake” is more frequent while the driver is under the distraction of texting as opposed to while he is under the distraction of listening to music.

Each percentage is calculated by using the basic percentage formula as seen in equation 4.1:

$$\text{Label A Percentage \%} = \frac{(\text{Number of points corresponding to label A})}{(\text{Number of points in the respective maximal set})} * 100 \quad (4.1)$$

All the results obtained by these experiments are shown below, separated by the labels applied. The results are represented in the form of tables whose respected bar plots are displayed only for the first label. For the remaining labels, only their respective tables of results are provided in this section and their respective bar plots can be found in Appendix B.

Label 1 Brake Analysis: The percentages for all the drivers under all the three

distractions for the label – brake analysis are compiled in table 4.1 given below where every pair of percentages are of the two classes (non-brake and brake) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	89.2% / 10.8%	87.1% / 12.9%	66.7% / 33.3%
Driver 2	83.3% / 16.7%	91.2% / 8.8%	74.0% / 26.0%
Driver 3	88.9% / 11.1%	86.8% / 13.2%	77.8% / 22.2%
Driver 4	85.0% / 15.0%	96.3% / 3.7%	88.4% / 11.6%
Driver 5	63.0% / 37.0%	87.8% / 12.2%	83.7% / 16.3%
Driver 6	81.2% / 18.8%	92.7% / 7.3%	69.0% / 31.0%
Driver 7	85.7% / 14.3%	88.5% / 11.5%	85.1% / 14.9%
Driver 8	67.7% / 32.3%	75.0% / 25.0%	67.4% / 32.6%
Driver 9	76.9% / 23.1%	87.5% / 12.5%	78.7% / 21.3%
Driver 10	85.4% / 14.6%	76.3% / 23.7%	82.4% / 17.6%
Driver 11	71.0% / 29.0%	100.0% / 0.0%	81.7% / 18.3%
Driver 12	89.7% / 10.3%	83.3% / 16.7%	80.0% / 20.0%
Driver 13	88.4% / 11.6%	92.3% / 7.7%	77.6% / 22.4%
Driver 14	86.2% / 13.8%	80.4% / 19.6%	67.4% / 32.6%
Driver 15	86.1% / 13.9%	77.8% / 22.2%	59.0% / 41.0%
Driver 16	80.0% / 20.0%	93.3% / 6.7%	91.7% / 8.3%
Driver 17	91.9% / 8.1%	84.2% / 15.8%	86.8% / 13.2%
Driver 18	65.4% / 34.6%	86.5% / 13.5%	80.9% / 19.1%
Driver 19	94.1% / 5.9%	75.0% / 25.0%	64.8% / 35.2%
Driver 20	70.6% / 29.4%	75.8% / 24.2%	78.1% / 21.9%
Driver 21	81.8% / 18.2%	77.4% / 22.6%	73.1% / 26.9%
Driver 22	70.7% / 29.3%	78.0% / 22.0%	66.7% / 33.3%
Driver 23	75.0% / 25.0%	75.0% / 25.0%	73.2% / 26.8%
Driver 24	75.0% / 25.0%	79.1% / 20.9%	85.7% / 14.3%
Driver 25	81.2% / 18.8%	66.0% / 34.0%	57.4% / 42.6%
Driver 26	79.1% / 20.9%	77.3% / 22.7%	69.5% / 30.5%
Driver 27	87.5% / 12.5%	78.7% / 21.3%	72.1% / 27.9%
Driver 28	69.4% / 30.6%	77.8% / 22.2%	70.9% / 29.1%
Driver 29	86.5% / 13.5%	66.7% / 33.3%	59.8% / 40.2%
Driver 30	81.4% / 18.6%	73.1% / 26.9%	84.2% / 15.8%
Driver 31	85.2% / 14.8%	78.3% / 21.7%	63.6% / 36.4%
Driver 32	70.7% / 29.3%	84.6% / 15.4%	82.3% / 17.7%
Driver 33	93.3% / 6.7%	88.7% / 11.3%	67.7% / 32.3%
Driver 34	81.0% / 19.0%	81.2% / 18.8%	68.4% / 31.6%
Driver 35	92.3% / 7.7%	91.4% / 8.6%	80.3% / 19.7%
Driver 36	89.1% / 10.9%	89.8% / 10.2%	80.5% / 19.5%
Driver 37	92.8% / 7.2%	91.7% / 8.3%	100.0% / 0.0%
Driver 38	86.1% / 13.9%	83.3% / 16.7%	59.1% / 40.9%
Driver 39	79.1% / 20.9%	84.1% / 15.9%	68.5% / 31.5%
Driver 40	71.7% / 28.3%	82.5% / 17.5%	81.3% / 18.7%

Table 4.1: Results from the Brake Analysis label

Label 2 Speed Limits: The percentages for all the drivers under all the three distractions for the label – speed limits are compiled in table 4.2 given below where every pair of percentages are of the two classes (below and above the average speed) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	54.1% / 45.9%	54.8% / 45.2%	80.6% / 19.4%
Driver 2	56.7% / 43.3%	44.1% / 55.9%	61.0% / 39.0%
Driver 3	48.1% / 51.9%	26.3% / 73.7%	47.2% / 52.8%
Driver 4	62.5% / 37.5%	66.7% / 33.3%	48.8% / 51.2%
Driver 5	70.4% / 29.6%	46.3% / 53.7%	51.0% / 49.0%
Driver 6	62.5% / 37.5%	29.3% / 70.7%	60.3% / 39.7%
Driver 7	50.0% / 50.0%	73.1% / 26.9%	58.1% / 41.9%
Driver 8	58.1% / 41.9%	54.2% / 45.8%	48.8% / 51.2%
Driver 9	56.4% / 43.6%	42.5% / 57.5%	49.3% / 50.7%
Driver 10	39.0% / 61.0%	44.7% / 55.3%	44.1% / 55.9%
Driver 11	58.1% / 41.9%	28.0% / 72.0%	45.1% / 54.9%
Driver 12	38.5% / 61.5%	55.0% / 45.0%	46.2% / 53.8%
Driver 13	41.9% / 58.1%	40.4% / 59.6%	70.7% / 29.3%
Driver 14	69.0% / 31.0%	43.1% / 56.9%	82.6% / 17.4%
Driver 15	38.9% / 61.1%	58.3% / 41.7%	65.6% / 34.4%
Driver 16	41.7% / 58.3%	28.0% / 72.0%	29.8% / 70.2%
Driver 17	56.8% / 43.2%	63.2% / 36.8%	50.0% / 50.0%
Driver 18	46.2% / 53.8%	37.8% / 62.2%	46.8% / 53.2%
Driver 19	31.4% / 68.6%	50.0% / 50.0%	58.0% / 42.0%
Driver 20	52.9% / 47.1%	57.6% / 42.4%	54.7% / 45.3%
Driver 21	60.6% / 39.4%	54.8% / 45.2%	46.2% / 53.8%
Driver 22	63.4% / 36.6%	48.8% / 51.2%	76.8% / 23.2%
Driver 23	81.2% / 18.8%	58.3% / 41.7%	67.6% / 32.4%
Driver 24	64.3% / 35.7%	67.4% / 32.6%	52.4% / 47.6%
Driver 25	53.1% / 46.9%	56.0% / 44.0%	67.6% / 32.4%
Driver 26	51.2% / 48.8%	65.9% / 34.1%	54.2% / 45.8%
Driver 27	40.6% / 59.4%	45.9% / 54.1%	65.6% / 34.4%
Driver 28	72.2% / 27.8%	52.8% / 47.2%	63.6% / 36.4%
Driver 29	44.2% / 55.8%	48.9% / 51.1%	59.8% / 40.2%
Driver 30	46.5% / 53.5%	71.2% / 28.8%	46.3% / 53.7%
Driver 31	33.3% / 66.7%	35.0% / 65.0%	55.8% / 44.2%
Driver 32	75.6% / 24.4%	40.4% / 59.6%	29.8% / 70.2%
Driver 33	25.0% / 75.0%	33.9% / 66.1%	67.7% / 32.3%
Driver 34	39.3% / 60.7%	34.1% / 65.9%	60.5% / 39.5%
Driver 35	36.5% / 63.5%	55.2% / 44.8%	63.4% / 36.6%
Driver 36	59.4% / 40.6%	42.4% / 57.6%	52.4% / 47.6%
Driver 37	21.7% / 78.3%	19.4% / 80.6%	39.3% / 60.7%
Driver 38	52.8% / 47.2%	81.0% / 19.0%	54.8% / 45.2%
Driver 39	46.5% / 53.5%	46.0% / 54.0%	66.7% / 33.3%
Driver 40	50.0% / 50.0%	52.5% / 47.5%	50.7% / 49.3%

Table 4.2: Results from Speed Limits label

Label 3 Linear Acceleration: The percentages for all the drivers under all the three distractions for the label – linear accerelation are compiled in table 4.3 given below where every pair of percentages are of the two classes (acceleration and deceleration) within the label.

Label 4 Turning Acceleration: The percentages for all the drivers under all the three distractions for the label – turning accerelation are compiled in table 4.4 given below where every pair of percentages are of the two classes (accelerating towards the left and towards the right) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	29.7% / 70.3%	35.5% / 64.5%	50.0% / 50.0%
Driver 2	50.0% / 50.0%	55.9% / 44.1%	58.4% / 41.6%
Driver 3	63.0% / 37.0%	57.9% / 42.1%	69.4% / 30.6%
Driver 4	45.0% / 55.0%	29.6% / 70.4%	39.5% / 60.5%
Driver 5	63.0% / 37.0%	41.5% / 58.5%	42.9% / 57.1%
Driver 6	50.0% / 50.0%	34.1% / 65.9%	39.7% / 60.3%
Driver 7	71.4% / 28.6%	38.5% / 61.5%	37.8% / 62.2%
Driver 8	71.0% / 29.0%	50.0% / 50.0%	41.9% / 58.1%
Driver 9	53.8% / 46.2%	67.5% / 32.5%	62.7% / 37.3%
Driver 10	43.9% / 56.1%	44.7% / 55.3%	50.0% / 50.0%
Driver 11	64.5% / 35.5%	44.0% / 56.0%	51.2% / 48.8%
Driver 12	33.3% / 66.7%	40.0% / 60.0%	40.0% / 60.0%
Driver 13	44.2% / 55.8%	42.3% / 57.7%	41.4% / 58.6%
Driver 14	41.4% / 58.6%	56.9% / 43.1%	41.3% / 58.7%
Driver 15	47.2% / 52.8%	52.8% / 47.2%	59.0% / 41.0%
Driver 16	45.0% / 55.0%	56.0% / 44.0%	45.2% / 54.8%
Driver 17	29.7% / 70.3%	31.6% / 68.4%	52.6% / 47.4%
Driver 18	76.9% / 23.1%	62.2% / 37.8%	42.6% / 57.4%
Driver 19	56.9% / 43.1%	43.2% / 56.8%	63.6% / 36.4%
Driver 20	47.1% / 52.9%	51.5% / 48.5%	40.6% / 59.4%
Driver 21	63.6% / 36.4%	56.5% / 43.5%	57.7% / 42.3%
Driver 22	48.8% / 51.2%	53.7% / 46.3%	52.2% / 47.8%
Driver 23	46.9% / 53.1%	41.7% / 58.3%	40.8% / 59.2%
Driver 24	50.0% / 50.0%	39.5% / 60.5%	47.6% / 52.4%
Driver 25	43.8% / 56.2%	60.0% / 40.0%	57.4% / 42.6%
Driver 26	53.5% / 46.5%	43.2% / 56.8%	59.3% / 40.7%
Driver 27	50.0% / 50.0%	42.6% / 57.4%	45.9% / 54.1%
Driver 28	44.4% / 55.6%	33.3% / 66.7%	45.5% / 54.5%
Driver 29	48.1% / 51.9%	55.6% / 44.4%	54.9% / 45.1%
Driver 30	48.8% / 51.2%	46.2% / 53.8%	42.1% / 57.9%
Driver 31	40.7% / 59.3%	43.3% / 56.7%	53.2% / 46.8%
Driver 32	63.4% / 36.6%	73.1% / 26.9%	53.2% / 46.8%
Driver 33	53.3% / 46.7%	48.4% / 51.6%	59.4% / 40.6%
Driver 34	44.0% / 56.0%	50.6% / 49.4%	48.7% / 51.3%
Driver 35	28.8% / 71.2%	39.7% / 60.3%	36.6% / 63.4%
Driver 36	39.1% / 60.9%	50.8% / 49.2%	36.6% / 63.4%
Driver 37	39.1% / 60.9%	36.1% / 63.9%	40.2% / 59.8%
Driver 38	44.4% / 55.6%	35.7% / 64.3%	57.0% / 43.0%
Driver 39	46.5% / 53.5%	52.4% / 47.6%	43.5% / 56.5%
Driver 40	52.2% / 47.8%	35.0% / 65.0%	54.7% / 45.3%

Table 4.3: Results from Linear Acceleration label

Driver ID	Hands-Free	Music	Text
Driver 1	45.9% / 54.1%	29.0% / 71.0%	27.8% / 72.2%
Driver 2	56.7% / 43.3%	64.7% / 35.3%	40.3% / 59.7%
Driver 3	51.9% / 48.1%	52.6% / 47.4%	63.9% / 36.1%
Driver 4	65.0% / 35.0%	40.7% / 59.3%	51.2% / 48.8%
Driver 5	44.4% / 55.6%	53.7% / 46.3%	49.0% / 51.0%
Driver 6	56.2% / 43.8%	43.9% / 56.1%	37.9% / 62.1%
Driver 7	42.9% / 57.1%	53.8% / 46.2%	39.2% / 60.8%
Driver 8	35.5% / 64.5%	37.5% / 62.5%	53.5% / 46.5%
Driver 9	53.8% / 46.2%	42.5% / 57.5%	53.3% / 46.7%
Driver 10	63.4% / 36.6%	42.1% / 57.9%	54.4% / 45.6%
Driver 11	54.8% / 45.2%	82.4% / 17.6%	62.2% / 37.8%
Driver 12	41.0% / 59.0%	55.0% / 45.0%	32.3% / 67.7%
Driver 13	41.9% / 58.1%	51.9% / 48.1%	51.7% / 48.3%
Driver 14	37.9% / 62.1%	62.7% / 37.3%	37.0% / 63.0%
Driver 15	58.3% / 41.7%	61.1% / 38.9%	52.5% / 47.5%
Driver 16	41.7% / 58.3%	48.0% / 52.0%	41.7% / 58.3%
Driver 17	48.6% / 51.4%	31.6% / 68.4%	42.1% / 57.9%
Driver 18	46.2% / 53.8%	51.4% / 48.6%	51.1% / 48.9%
Driver 19	49.0% / 51.0%	65.9% / 34.1%	48.9% / 51.1%
Driver 20	82.4% / 17.6%	66.7% / 33.3%	29.7% / 70.3%
Driver 21	33.3% / 66.7%	56.5% / 43.5%	38.5% / 61.5%
Driver 22	65.9% / 34.1%	58.5% / 41.5%	60.9% / 39.1%
Driver 23	71.9% / 28.1%	33.3% / 66.7%	32.4% / 67.6%
Driver 24	53.6% / 46.4%	51.2% / 48.8%	61.9% / 38.1%
Driver 25	50.0% / 50.0%	52.0% / 48.0%	45.6% / 54.4%
Driver 26	55.8% / 44.2%	34.1% / 65.9%	39.0% / 61.0%
Driver 27	56.2% / 43.8%	49.2% / 50.8%	45.9% / 54.1%
Driver 28	72.2% / 27.8%	69.4% / 30.6%	58.2% / 41.8%
Driver 29	44.2% / 55.8%	48.9% / 51.1%	47.6% / 52.4%
Driver 30	53.5% / 46.5%	57.7% / 42.3%	60.0% / 40.0%
Driver 31	29.6% / 70.4%	48.3% / 51.7%	42.9% / 57.1%
Driver 32	39.0% / 61.0%	42.3% / 57.7%	44.7% / 55.3%
Driver 33	46.7% / 53.3%	50.0% / 50.0%	41.7% / 58.3%
Driver 34	44.0% / 56.0%	48.2% / 51.8%	52.6% / 47.4%
Driver 35	57.7% / 42.3%	56.9% / 43.1%	36.6% / 63.4%
Driver 36	37.5% / 62.5%	47.5% / 52.5%	29.3% / 70.7%
Driver 37	42.0% / 58.0%	44.4% / 55.6%	76.8% / 23.2%
Driver 38	44.4% / 55.6%	47.6% / 52.4%	50.5% / 49.5%
Driver 39	51.2% / 48.8%	57.1% / 42.9%	42.6% / 57.4%
Driver 40	39.1% / 60.9%	50.0% / 50.0%	53.3% / 46.7%

Table 4.4: Results from Turning Acceleration label

Label 5 Altitude Acceleration: The percentages for all the drivers under all the three distractions for the label – altitude accerelation are compiled in table 4.5 given below where every pair of percentages are of the two classes (lower and higher altitudes) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	62.2% / 37.8%	61.3% / 38.7%	47.2% / 52.8%
Driver 2	43.3% / 56.7%	44.1% / 55.9%	39.0% / 61.0%
Driver 3	63.0% / 37.0%	50.0% / 50.0%	38.9% / 61.1%
Driver 4	50.0% / 50.0%	44.4% / 55.6%	65.1% / 34.9%
Driver 5	55.6% / 44.4%	34.1% / 65.9%	59.2% / 40.8%
Driver 6	65.6% / 34.4%	41.5% / 58.5%	44.8% / 55.2%
Driver 7	50.0% / 50.0%	46.2% / 53.8%	47.3% / 52.7%
Driver 8	58.1% / 41.9%	50.0% / 50.0%	55.8% / 44.2%
Driver 9	41.0% / 59.0%	37.5% / 62.5%	33.3% / 66.7%
Driver 10	43.9% / 56.1%	60.5% / 39.5%	39.7% / 60.3%
Driver 11	74.2% / 25.8%	99.2% / 0.8%	61.0% / 39.0%
Driver 12	59.0% / 41.0%	63.3% / 36.7%	67.7% / 32.3%
Driver 13	67.4% / 32.6%	38.5% / 61.5%	48.3% / 51.7%
Driver 14	55.2% / 44.8%	45.1% / 54.9%	52.2% / 47.8%
Driver 15	63.9% / 36.1%	61.1% / 38.9%	52.5% / 47.5%
Driver 16	53.3% / 46.7%	37.3% / 62.7%	39.3% / 60.7%
Driver 17	45.9% / 54.1%	36.8% / 63.2%	47.4% / 52.6%
Driver 18	69.2% / 30.8%	54.1% / 45.9%	46.8% / 53.2%
Driver 19	56.9% / 43.1%	63.6% / 36.4%	48.9% / 51.1%
Driver 20	52.9% / 47.1%	42.4% / 57.6%	53.1% / 46.9%
Driver 21	45.5% / 54.5%	30.6% / 69.4%	36.5% / 63.5%
Driver 22	58.5% / 41.5%	61.0% / 39.0%	50.7% / 49.3%
Driver 23	34.4% / 65.6%	37.5% / 62.5%	57.7% / 42.3%
Driver 24	75.0% / 25.0%	44.2% / 55.8%	50.0% / 50.0%
Driver 25	50.0% / 50.0%	58.0% / 42.0%	45.6% / 54.4%
Driver 26	46.5% / 53.5%	50.0% / 50.0%	52.5% / 47.5%
Driver 27	53.1% / 46.9%	52.5% / 47.5%	54.1% / 45.9%
Driver 28	55.6% / 44.4%	66.7% / 33.3%	61.8% / 38.2%
Driver 29	48.1% / 51.9%	51.1% / 48.9%	50.0% / 50.0%
Driver 30	53.5% / 46.5%	40.4% / 59.6%	43.2% / 56.8%
Driver 31	51.9% / 48.1%	60.0% / 40.0%	61.0% / 39.0%
Driver 32	39.0% / 61.0%	34.6% / 65.4%	39.7% / 60.3%
Driver 33	55.0% / 45.0%	67.7% / 32.3%	49.0% / 51.0%
Driver 34	45.2% / 54.8%	40.0% / 60.0%	32.9% / 67.1%
Driver 35	55.8% / 44.2%	50.0% / 50.0%	38.0% / 62.0%
Driver 36	57.8% / 42.2%	57.6% / 42.4%	42.7% / 57.3%
Driver 37	40.6% / 59.4%	33.3% / 66.7%	91.1% / 8.9%
Driver 38	52.8% / 47.2%	33.3% / 66.7%	39.8% / 60.2%
Driver 39	60.5% / 39.5%	63.5% / 36.5%	58.3% / 41.7%
Driver 40	58.7% / 41.3%	75.0% / 25.0%	60.0% / 40.0%

Table 4.5: Results from Altitude Acceleration label

Label 6 Gap between lanes: The percentages for all the drivers under all the three distractions for the label – gap between lanes are compiled in table 4.6 given below where every pair of percentages are of the two classes (left and right of the center of the lane) within the label.

Label 7 Split Sections: The percentages for all the drivers under all the three distractions for the label – split segments are described below depending on how many segments the dataset is split into.

Label 7a: The percentages for all the drivers under all the three distractions for the label – two segments are compiled in table 4.7 given below where every pair of percentages are of the two classes (the first and last segments) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	43.2% / 56.8%	74.2% / 25.8%	30.6% / 69.4%
Driver 2	46.7% / 53.3%	79.4% / 20.6%	26.0% / 74.0%
Driver 3	40.7% / 59.3%	39.5% / 60.5%	44.4% / 55.6%
Driver 4	77.5% / 22.5%	51.9% / 48.1%	53.5% / 46.5%
Driver 5	63.0% / 37.0%	34.1% / 65.9%	61.2% / 38.8%
Driver 6	46.9% / 53.1%	68.3% / 31.7%	17.2% / 82.8%
Driver 7	21.4% / 78.6%	80.8% / 19.2%	43.2% / 56.8%
Driver 8	67.7% / 32.3%	79.2% / 20.8%	69.8% / 30.2%
Driver 9	30.8% / 69.2%	60.0% / 40.0%	53.3% / 46.7%
Driver 10	51.2% / 48.8%	50.0% / 50.0%	54.4% / 45.6%
Driver 11	80.6% / 19.4%	32.0% / 68.0%	89.0% / 11.0%
Driver 12	23.1% / 76.9%	83.3% / 16.7%	33.8% / 66.2%
Driver 13	20.9% / 79.1%	73.1% / 26.9%	46.6% / 53.4%
Driver 14	37.9% / 62.1%	84.3% / 15.7%	73.9% / 26.1%
Driver 15	50.0% / 50.0%	69.4% / 30.6%	78.7% / 21.3%
Driver 16	20.0% / 80.0%	74.7% / 25.3%	26.2% / 73.8%
Driver 17	51.4% / 48.6%	73.7% / 26.3%	44.7% / 55.3%
Driver 18	61.5% / 38.5%	29.7% / 70.3%	36.2% / 63.8%
Driver 19	23.5% / 76.5%	86.4% / 13.6%	60.2% / 39.8%
Driver 20	58.8% / 41.2%	63.6% / 36.4%	18.8% / 81.2%
Driver 21	51.5% / 48.5%	67.7% / 32.3%	67.3% / 32.7%
Driver 22	80.5% / 19.5%	58.5% / 41.5%	59.4% / 40.6%
Driver 23	28.1% / 71.9%	54.2% / 45.8%	35.2% / 64.8%
Driver 24	42.9% / 57.1%	60.5% / 39.5%	64.3% / 35.7%
Driver 25	53.1% / 46.9%	56.0% / 44.0%	50.0% / 50.0%
Driver 26	46.5% / 53.5%	84.1% / 15.9%	39.0% / 61.0%
Driver 27	21.9% / 78.1%	78.7% / 21.3%	34.4% / 65.6%
Driver 28	55.6% / 44.4%	66.7% / 33.3%	58.2% / 41.8%
Driver 29	25.0% / 75.0%	62.2% / 37.8%	43.9% / 56.1%
Driver 30	39.5% / 60.5%	59.6% / 40.4%	47.4% / 52.6%
Driver 31	20.4% / 79.6%	90.0% / 10.0%	31.2% / 68.8%
Driver 32	48.8% / 51.2%	48.1% / 51.9%	43.3% / 56.7%
Driver 33	30.0% / 70.0%	95.2% / 4.8%	67.7% / 32.3%
Driver 34	46.4% / 53.6%	54.1% / 45.9%	55.3% / 44.7%
Driver 35	42.3% / 57.7%	67.2% / 32.8%	42.3% / 57.7%
Driver 36	39.1% / 60.9%	84.7% / 15.3%	25.6% / 74.4%
Driver 37	62.3% / 37.7%	61.1% / 38.9%	69.6% / 30.4%
Driver 38	50.0% / 50.0%	64.3% / 35.7%	47.3% / 52.7%
Driver 39	60.5% / 39.5%	71.4% / 28.6%	65.7% / 34.3%
Driver 40	23.9% / 76.1%	72.5% / 27.5%	44.0% / 56.0%

Table 4.6: Results from Gap between Lanes label

Driver ID	Hands-Free	Music	Text
Driver 1	54.1% / 45.9%	80.6% / 19.4%	66.7% / 33.3%
Driver 2	63.3% / 36.7%	76.5% / 23.5%	46.8% / 53.2%
Driver 3	51.9% / 48.1%	63.2% / 36.8%	52.8% / 47.2%
Driver 4	57.5% / 42.5%	59.3% / 40.7%	53.5% / 46.5%
Driver 5	81.5% / 18.5%	41.5% / 58.5%	55.1% / 44.9%
Driver 6	59.4% / 40.6%	51.2% / 48.8%	53.4% / 46.6%
Driver 7	50.0% / 50.0%	73.1% / 26.9%	60.8% / 39.2%
Driver 8	71.0% / 29.0%	66.7% / 33.3%	53.5% / 46.5%
Driver 9	53.8% / 46.2%	70.0% / 30.0%	49.3% / 50.7%
Driver 10	61.0% / 39.0%	65.8% / 34.2%	54.4% / 45.6%
Driver 11	71.0% / 29.0%	43.2% / 56.8%	52.4% / 47.6%
Driver 12	43.6% / 56.4%	68.3% / 31.7%	46.2% / 53.8%
Driver 13	41.9% / 58.1%	67.3% / 32.7%	65.5% / 34.5%
Driver 14	79.3% / 20.7%	58.8% / 41.2%	78.3% / 21.7%
Driver 15	50.0% / 50.0%	58.3% / 41.7%	70.5% / 29.5%
Driver 16	35.0% / 65.0%	50.7% / 49.3%	31.0% / 69.0%
Driver 17	64.9% / 35.1%	84.2% / 15.8%	47.4% / 52.6%
Driver 18	42.3% / 57.7%	43.2% / 56.8%	51.1% / 48.9%
Driver 19	35.3% / 64.7%	84.1% / 15.9%	63.6% / 36.4%
Driver 20	52.9% / 47.1%	60.6% / 39.4%	46.9% / 53.1%
Driver 21	66.7% / 33.3%	72.6% / 27.4%	51.9% / 48.1%
Driver 22	78.0% / 22.0%	51.2% / 48.8%	73.9% / 26.1%
Driver 23	71.9% / 28.1%	58.3% / 41.7%	66.2% / 33.8%
Driver 24	57.1% / 42.9%	55.8% / 44.2%	50.0% / 50.0%
Driver 25	65.6% / 34.4%	62.0% / 38.0%	63.2% / 36.8%
Driver 26	46.5% / 53.5%	88.6% / 11.4%	49.2% / 50.8%
Driver 27	46.9% / 53.1%	54.1% / 45.9%	57.4% / 42.6%
Driver 28	72.2% / 27.8%	44.4% / 55.6%	60.0% / 40.0%
Driver 29	46.2% / 53.8%	73.3% / 26.7%	57.3% / 42.7%
Driver 30	51.2% / 48.8%	71.2% / 28.8%	49.5% / 50.5%
Driver 31	27.8% / 72.2%	66.7% / 33.3%	46.8% / 53.2%
Driver 32	78.0% / 22.0%	46.2% / 53.8%	31.9% / 68.1%
Driver 33	31.7% / 68.3%	64.5% / 35.5%	63.5% / 36.5%
Driver 34	46.4% / 53.6%	54.1% / 45.9%	59.2% / 40.8%
Driver 35	40.4% / 59.6%	48.3% / 51.7%	63.4% / 36.6%
Driver 36	51.6% / 48.4%	83.1% / 16.9%	36.6% / 63.4%
Driver 37	30.4% / 69.6%	48.6% / 51.4%	45.5% / 54.5%
Driver 38	50.0% / 50.0%	69.0% / 31.0%	50.5% / 49.5%
Driver 39	58.1% / 41.9%	41.3% / 58.7%	73.1% / 26.9%
Driver 40	47.8% / 52.2%	50.0% / 50.0%	50.7% / 49.3%

Table 4.7: Results from Split two Segments label

Label 7b: The percentages for all the drivers under all the three distractions for the label – three segments are compiled in table 4.8 given below where every set of percentages are of the three classes (the first, middle and last segments) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	51.4% / 37.8% / 10.8%	51.6% / 35.5% / 12.9%	63.9% / 22.2% / 13.9%
Driver 2	56.7% / 33.3% / 10.0%	41.2% / 44.1% / 14.7%	33.8% / 31.2% / 35.1%
Driver 3	51.9% / 44.4% / 3.7%	39.5% / 55.3% / 5.3%	47.2% / 36.1% / 16.7%
Driver 4	47.5% / 40.0% / 12.5%	51.9% / 29.6% / 18.5%	37.2% / 51.2% / 11.6%
Driver 5	70.4% / 25.9% / 3.7%	31.7% / 43.9% / 24.4%	44.9% / 32.7% / 22.4%
Driver 6	50.0% / 34.4% / 15.6%	24.4% / 48.8% / 26.8%	41.4% / 19.0% / 39.7%
Driver 7	50.0% / 35.7% / 14.3%	53.8% / 30.8% / 15.4%	45.9% / 39.2% / 14.9%
Driver 8	58.1% / 32.3% / 9.7%	50.0% / 33.3% / 16.7%	37.2% / 39.5% / 23.3%
Driver 9	48.7% / 25.6% / 25.6%	45.0% / 35.0% / 20.0%	41.3% / 38.7% / 20.0%
Driver 10	39.0% / 56.1% / 4.9%	36.8% / 60.5% / 2.6%	42.6% / 38.2% / 19.1%
Driver 11	58.1% / 41.9% / 0.0%	13.6% / 72.0% / 14.4%	32.9% / 53.7% / 13.4%
Driver 12	35.9% / 41.0% / 23.1%	50.0% / 36.7% / 13.3%	41.5% / 43.1% / 15.4%
Driver 13	34.9% / 44.2% / 20.9%	38.5% / 59.6% / 1.9%	55.2% / 32.8% / 12.1%
Driver 14	72.4% / 24.1% / 3.4%	33.3% / 47.1% / 19.6%	73.9% / 8.7% / 17.4%
Driver 15	38.9% / 44.4% / 16.7%	41.7% / 44.4% / 13.9%	63.9% / 24.6% / 11.5%
Driver 16	35.0% / 45.0% / 20.0%	33.3% / 54.7% / 12.0%	25.0% / 41.7% / 33.3%
Driver 17	56.8% / 32.4% / 10.8%	52.6% / 42.1% / 5.3%	44.7% / 31.6% / 23.7%
Driver 18	30.8% / 50.0% / 19.2%	27.0% / 54.1% / 18.9%	44.7% / 40.4% / 14.9%
Driver 19	27.5% / 60.8% / 11.8%	52.3% / 40.9% / 6.8%	56.8% / 33.0% / 10.2%
Driver 20	41.2% / 35.3% / 23.5%	51.5% / 30.3% / 18.2%	40.6% / 21.9% / 37.5%
Driver 21	54.5% / 36.4% / 9.1%	53.2% / 32.3% / 14.5%	44.2% / 48.1% / 7.7%
Driver 22	61.0% / 39.0% / 0.0%	41.5% / 31.7% / 26.8%	62.3% / 21.7% / 15.9%
Driver 23	65.6% / 21.9% / 12.5%	33.3% / 41.7% / 25.0%	63.4% / 7.0% / 29.6%
Driver 24	50.0% / 25.0% / 25.0%	41.9% / 32.6% / 25.6%	42.9% / 42.9% / 14.3%
Driver 25	50.0% / 43.8% / 6.2%	40.0% / 32.0% / 28.0%	58.8% / 14.7% / 26.5%
Driver 26	41.9% / 30.2% / 27.9%	63.6% / 27.3% / 9.1%	45.8% / 10.2% / 44.1%
Driver 27	40.6% / 53.1% / 6.2%	31.1% / 54.1% / 14.8%	55.7% / 21.3% / 23.0%
Driver 28	63.9% / 25.0% / 11.1%	38.9% / 44.4% / 16.7%	56.4% / 21.8% / 21.8%
Driver 29	34.6% / 40.4% / 25.0%	51.1% / 31.1% / 17.8%	53.7% / 17.1% / 29.3%
Driver 30	46.5% / 37.2% / 16.3%	69.2% / 15.4% / 15.4%	43.2% / 31.6% / 25.3%
Driver 31	27.8% / 44.4% / 27.8%	35.0% / 51.7% / 13.3%	41.6% / 23.4% / 35.1%
Driver 32	75.6% / 9.8% / 14.6%	38.5% / 46.2% / 15.4%	26.2% / 34.8% / 39.0%
Driver 33	25.0% / 58.3% / 16.7%	30.6% / 64.5% / 4.8%	59.4% / 19.8% / 20.8%
Driver 34	25.0% / 56.0% / 19.0%	28.2% / 54.1% / 17.6%	53.9% / 21.1% / 25.0%
Driver 35	32.7% / 50.0% / 17.3%	44.8% / 31.0% / 24.1%	60.6% / 26.8% / 12.7%
Driver 36	45.3% / 32.8% / 21.9%	44.1% / 47.5% / 8.5%	31.7% / 14.6% / 53.7%
Driver 37	20.3% / 72.5% / 7.2%	20.8% / 73.6% / 5.6%	35.7% / 55.4% / 8.9%
Driver 38	38.9% / 55.6% / 5.6%	61.9% / 16.7% / 21.4%	39.8% / 26.9% / 33.3%
Driver 39	41.9% / 53.5% / 4.7%	30.2% / 41.3% / 28.6%	59.3% / 27.8% / 13.0%
Driver 40	43.5% / 32.6% / 23.9%	32.5% / 32.5% / 35.0%	41.3% / 12.0% / 46.7%

Table 4.8: Results from Split three Segments label

Label 7c: The percentages for all the drivers under all the three distractions for the label – four segments are compiled in table 4.9 given below where every set of percentages are of the four classes (the first, second, third and the last segments) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	43.2% / 8.1% / 35.1% / 13.5%	29.0% / 45.2% / 6.5% / 19.4%	58.3% / 8.3% / 13.9% / 19.4%
Driver 2	53.3% / 16.7% / 26.7% / 3.3%	35.3% / 38.2% / 14.7% / 11.8%	31.2% / 22.1% / 16.9% / 29.9%
Driver 3	44.4% / 14.8% / 33.3% / 7.4%	36.8% / 26.3% / 31.6% / 5.3%	41.7% / 13.9% / 30.6% / 13.9%
Driver 4	40.0% / 20.0% / 27.5% / 12.5%	51.9% / 7.4% / 25.9% / 14.8%	25.6% / 41.9% / 18.6% / 14.0%
Driver 5	59.3% / 25.9% / 7.4% / 7.4%	26.8% / 14.6% / 39.0% / 19.5%	44.9% / 14.3% / 24.5% / 16.3%
Driver 6	43.8% / 25.0% / 15.6% & 15.6%	24.4% / 19.5% / 39.0% / 17.1%	41.4% / 13.8% / 3.4% / 41.4%
Driver 7	50.0% / 28.6% / 21.4% / 0.0%	38.5% / 23.1% / 23.1% / 15.4%	35.1% / 32.4% / 10.8% / 21.6%
Driver 8	58.1% / 16.1% / 19.4% / 6.5%	45.8% / 20.8% / 29.2% / 4.2%	30.2% / 25.6% / 16.3% / 27.9%
Driver 9	38.5% / 7.7% / 23.1% / 30.8%	52.5% / 15.0% / 15.0% / 17.5%	34.7% / 14.7% / 33.3% / 17.3%
Driver 10	34.1% / 26.8% / 31.7% / 7.3%	28.9% / 31.6% / 18.4% / 21.1%	33.8% / 27.9% / 22.1% / 16.2%
Driver 11	51.6% / 22.6% / 22.6% / 3.2%	10.4% / 38.4% / 31.2% / 20.0%	29.3% / 30.5% / 24.4% / 15.9%
Driver 12	33.3% / 17.9% / 28.2% / 20.5%	50.0% / 20.0% / 16.7% / 13.3%	40.0% / 6.2% / 40.0% / 13.8%
Driver 13	32.6% / 9.3% / 44.2% / 14.0%	30.8% / 30.8% / 34.6% / 3.8%	44.8% / 20.7% / 17.2%17.2%
Driver 14	62.1% / 20.7% / 10.3% / 6.9%	29.4% / 21.6% / 33.3% / 15.7%	65.2% / 15.2% / 6.5% / 13.0%
Driver 15	27.8% / 27.8% / 33.3% / 11.1%	33.3% / 27.8% / 19.4% & 19.4%	59.0% / 11.5% / 11.5% / 18.0%
Driver 16	31.7% / 6.7% / 46.7% / 15.0%	26.7% / 17.3% / 48.0% / 8.0%	21.4% / 7.1% / 48.8% / 22.6%
Driver 17	45.9% / 32.4% / 13.5% / 8.1%	47.4% / 26.3% / 21.1% / 5.3%	42.1% / 7.9% / 23.7% / 26.3%
Driver 18	23.1% / 30.8% / 19.2% / 26.9%	27.0% / 13.5% / 43.2% / 16.2%	42.6% / 17.0% / 27.7% / 12.8%
Driver 19	29.4% / 23.5% / 35.3% / 11.8%	47.7% / 31.8% / 9.1% / 11.4%	53.4% / 15.9% / 22.7% / 8.0%
Driver 20	35.3% / 23.5% / 17.6% / 23.5%	42.4% / 18.2% / 21.2% / 18.2%	35.9% / 14.1% / 20.3% / 29.7%

Driver 21	51.5% / 18.2% / 27.3% / 3.0%	56.5% / 16.1% / 14.5% / 12.9%	44.2% / 23.1% / 25.0% / 7.7%
Driver 22	56.1% / 22.0% / 22.0% / 0.0%	29.3% / 19.5% / 17.1% / 34.1%	53.6% / 18.8% / 10.1% / 17.4%
Driver 23	59.4% / 9.4% / 15.6% & 15.6%	33.3% / 16.7% / 25.0% & 25.0%	52.1% / 12.7% / 8.5% / 26.8%
Driver 24	46.4% / 17.9% / 14.3% / 21.4%	25.6% / 23.3% / 16.3% / 34.9%	42.9% / 11.9% / 31.0% / 14.3%
Driver 25	40.6% / 34.4% / 12.5% & 12.5%	34.0% / 36.0% / 8.0% / 22.0%	55.9% / 13.2% / 2.9% / 27.9%
Driver 26	37.2% / 14.0% / 32.6% / 16.3%	54.5% / 27.3% / 4.5% / 13.6%	40.7% / 6.8% / 23.7% / 28.8%
Driver 27	37.5% / 12.5% / 46.9% / 3.1%	31.1% / 23.0% / 34.4% / 11.5%	52.5% / 4.9% / 19.7% / 23.0%
Driver 28	52.8% / 19.4% / 8.3% / 19.4%	27.8% / 22.2% / 22.2% / 27.8%	49.1% / 9.1% / 14.5% / 27.3%
Driver 29	28.8% / 19.2% / 36.5% / 15.4%	48.9% / 20.0% / 20.0% / 11.1%	47.6% / 11.0% / 11.0% / 30.5%
Driver 30	32.6% / 30.2% / 25.6% / 11.6%	57.7% / 11.5% / 15.4% / 15.4%	30.5% / 17.9% / 22.1% / 29.5%
Driver 31	24.1% / 7.4% / 48.1% / 20.4%	21.7% / 40.0% / 28.3% / 10.0%	40.3% / 13.0% / 14.3% / 32.5%
Driver 32	65.9% / 7.3% / 7.3% / 19.5%	32.7% / 7.7% / 42.3% / 17.3%	20.6% / 9.9% / 45.4% / 24.1%
Driver 33	23.3% / 11.7% / 56.7% / 8.3%	21.0% / 37.1% / 30.6% / 11.3%	54.2% / 10.4% / 9.4% / 26.0%
Driver 34	14.3% / 40.5% / 21.4% / 23.8%	25.9% / 31.8% / 21.2% / 21.2%	46.1% / 11.8% / 15.8% / 26.3%
Driver 35	32.7% / 15.4% / 46.2% / 5.8%	39.7% / 8.6% / 34.5% / 17.2%	59.2% / 8.5% / 18.3% / 14.1%
Driver 36	45.3% / 9.4% / 28.1% / 17.2%	37.3% / 35.6% / 18.6% / 8.5%	24.4% / 12.2% / 13.4% / 50.0%
Driver 37	18.8% / 30.4% / 44.9% / 5.8%	20.8% / 22.2% / 51.4% / 5.6%	31.2% / 25.0% / 41.1% / 2.7%
Driver 38	36.1% / 19.4% / 27.8% / 16.7%	45.2% / 7.1% / 11.9% / 35.7%	35.5% / 16.1% / 14.0% / 34.4%
Driver 39	37.2% / 25.6% / 27.9% / 9.3%	23.8% / 22.2% / 28.6% / 25.4%	48.1% / 25.0% / 14.8% / 12.0%
Driver 40	37.0% / 17.4% / 21.7% / 23.9%	25.0% / 22.5% / 22.5% / 30.0%	32.0% / 14.7% / 5.3% / 48.0%

Table 4.9: Results from Split four Segments label

Label 7d: The percentages for all the drivers under all the three distractions for the label – five segments are compiled in table 4.10 given below where every set of percentages are of the five classes (the first, second, third, fourth and the last segments) within the label.

Driver ID	Hands-Free	Music	Text
Driver 1	43.2% / 2.7% / 24.3% / 16.2% / 13.5%	29.0% / 19.4% / 29.0% / 0.0% / 22.6%	50.0% / 8.3% / 8.3% / 16.7%16.7%
Driver 2	36.7% / 23.3% / 30.0% / 10.0% / 0.0%	35.3% / 32.4% / 14.7% / 0.0% / 17.6%	20.8% / 22.1% / 15.6% / 10.4% / 31.2%
Driver 3	40.7% / 7.4% / 40.7% / 7.4% / 3.7%	31.6% / 2.6% / 44.7% / 15.8% / 5.3%	38.9% / 11.1% / 19.4% / 16.7% / 13.9%
Driver 4	32.5% / 20.0% / 25.0% / 10.0% / 12.5%	40.7% / 0.0% / 22.2% / 11.1% / 25.9%	23.3% / 25.6% / 14.0% / 18.6%18.6%
Driver 5	48.1% / 22.2% / 18.5% / 0.0% / 11.1%	9.8% / 22.0% / 31.7% / 17.1% / 19.5%	38.8% / 12.2% / 16.3% / 10.2% / 22.4%
Driver 6	37.5% / 28.1% / 12.5% / 6.2% / 15.6%	19.5% / 14.6% / 26.8% / 22.0% / 17.1%	36.2% / 10.3% / 10.3% / 5.2% / 37.9%
Driver 7	28.6% / 28.6% / 28.6% / 7.1% & 7.1%	15.4% / 38.5% / 19.2% / 11.5% / 15.4%	31.1% / 24.3% / 14.9% / 4.1% / 25.7%
Driver 8	58.1% / 6.5% / 16.1% / 16.1% / 3.2%	41.7% / 12.5% / 20.8% / 20.8% / 4.2%	27.9% / 20.9% / 11.6% / 9.3% / 30.2%
Driver 9	20.5% / 20.5% / 7.7% / 20.5% / 30.8%	52.5% / 5.0% / 27.5% / 7.5% / 7.5%	28.0% / 14.7% / 29.3% / 10.7% / 17.3%
Driver 10	34.1% / 29.3% / 17.1% / 17.1% / 2.4%	34.2% / 23.7% / 21.1% / 2.6% / 18.4%	25.0% / 20.6% / 20.6% / 14.7% / 19.1%
Driver 11	45.2% / 6.5% / 32.3% / 6.5% / 9.7%	8.0% / 9.6% / 59.2% / 11.2% / 12.0%	20.7% / 20.7% / 34.1% / 7.3% / 17.1%
Driver 12	28.2% / 5.1% / 33.3% / 12.8% / 20.5%	43.3% / 13.3% / 18.3% / 10.0% / 15.0%	33.8% / 7.7% / 24.6% / 20.0% / 13.8%
Driver 13	25.6% / 9.3% / 39.5% / 7.0% / 18.6%	28.8% / 28.8% / 23.1% / 19.2% / 0.0%	39.7% / 27.6% / 12.1% / 3.4% / 17.2%
Driver 14	41.4% / 37.9% / 13.8% / 3.4% & 3.4%	27.5% / 11.8% / 23.5% / 17.6% / 19.6%	47.8% / 26.1% / 0.0% / 6.5% / 19.6%
Driver 15	22.2% / 19.4% / 38.9% / 13.9% / 5.6%	22.2% / 19.4% / 33.3% / 8.3% / 16.7%	42.6% / 18.0% / 11.5% / 6.6% / 21.3%
Driver 16	26.7% / 8.3% / 23.3% / 23.3% / 18.3%	28.0% / 14.7% / 41.3% / 9.3% / 6.7%	16.7% / 8.3% / 39.3% / 13.1% / 22.6%
Driver 17	37.8% / 18.9% / 21.6% / 13.5% / 8.1%	42.1% / 21.1% / 26.3% / 5.3% / 5.3%	31.6% / 18.4% / 15.8% / 10.5% / 23.7%
Driver 18	23.1% / 15.4% / 30.8% / 7.7% / 23.1%	21.6% / 5.4% / 51.4% / 10.8% / 10.8%	36.2% / 12.8% / 19.1% / 14.9% / 17.0%
Driver 19	21.6% / 7.8% / 56.9% / 9.8% / 3.9%	54.5% / 9.1% / 27.3% / 0.0% / 9.1%	51.1% / 9.1% / 22.7% / 11.4% / 5.7%

Driver 20	17.6% / 35.3% / 17.6% / 5.9% / 23.5%	39.4% / 18.2% / 12.1% / 15.2% / 15.2%	28.1% / 14.1% / 10.9% / 15.6% / 31.2%
Driver 21	36.4% / 21.2% / 33.3% / 6.1% / 3.0%	38.7% / 17.7% / 21.0% / 11.3% / 11.3%	36.5% / 11.5% / 25.0% / 13.5% / 13.5%
Driver 22	46.3% / 22.0% / 22.0% / 7.3% / 2.4%	22.0% / 12.2% / 17.1% / 12.2% / 36.6%	39.1% / 26.1% / 11.6% / 4.3% / 18.8%
Driver 23	43.8% / 21.9% / 15.6% / 6.2% / 12.5%	20.8% / 20.8% / 20.8% / 4.2% / 33.3%	40.8% / 21.1% / 4.2% / 7.0% / 26.8%
Driver 24	39.3% / 14.3% / 21.4% / 14.3% / 10.7%	27.9% / 25.6% / 16.3% / 9.3% / 20.9%	38.1% / 7.1% / 11.9% / 23.8% / 19.0%
Driver 25	34.4% / 25.0% / 25.0% / 3.1% / 12.5%	32.0% / 10.0% / 28.0% / 8.0% / 22.0%	44.1% / 13.2% / 8.8% / 1.5% / 32.4%
Driver 26	32.6% / 11.6% / 18.6% / 18.6% & 18.6%	45.5% / 15.9% / 25.0% / 0.0% / 13.6%	37.3% / 5.1% / 11.9% / 23.7% / 22.0%
Driver 27	31.2% / 6.2% / 50.0% / 9.4% / 3.1%	26.2% / 13.1% / 26.2% / 19.7% / 14.8%	42.6% / 6.6% / 13.1% / 11.5% / 26.2%
Driver 29	23.1% / 15.4% / 25.0% / 23.1% / 13.5%	37.8% / 15.6% / 15.6% / 11.1% / 20.0%	39.0% / 15.9% / 9.8% / 8.5% / 26.8%
Driver 30	34.9% / 23.3% / 30.2% / 7.0% / 4.7%	48.1% / 15.4% / 7.7% / 9.6% / 19.2%	21.1% / 23.2% / 20.0% / 14.7% / 21.1%
Driver 31	22.2% / 7.4% / 18.5% / 38.9% / 13.0%	20.0% / 26.7% / 26.7% / 13.3% / 13.3%	32.5% / 15.6% / 7.8% / 11.7% / 32.5%
Driver 32	43.9% / 26.8% / 2.4% / 7.3% / 19.5%	25.0% / 17.3% / 21.2% / 15.4% / 21.2%	17.7% / 13.5% / 24.1% / 24.8% / 19.9%
Driver 33	18.3% / 6.7% / 31.7% / 33.3% / 10.0%	22.6% / 22.6% / 41.9% / 3.2% / 9.7%	46.9% / 14.6% / 8.3% / 5.2% / 25.0%
Driver 34	10.7% / 33.3% / 26.2% / 13.1% / 16.7%	21.2% / 20.0% / 31.8% / 7.1% / 20.0%	40.8% / 18.4% / 9.2% / 5.3% / 26.3%
Driver 35	25.0% / 9.6% / 26.9% / 28.8% / 9.6%	27.6% / 15.5% / 24.1% / 15.5% / 17.2%	52.1% / 9.9% / 14.1% / 14.1% / 9.9%
Driver 36	40.6% / 9.4% / 18.8% / 15.6% / 15.6%	32.2% / 16.9% / 37.3% / 5.1% / 8.5%	19.5% / 12.2% / 4.9% / 14.6% / 48.8%
Driver 37	15.9% / 13.0% / 63.8% / 4.3% / 2.9%	11.1% / 8.3% / 63.9% / 6.9% / 9.7%	23.2% / 15.2% / 42.0% / 17.0% / 2.7%
Driver 38	30.6% / 13.9% / 27.8% / 8.3% / 19.4%	26.2% / 23.8% / 11.9% / 7.1% / 31.0%	34.4% / 12.9% / 12.9% / 3.2% / 36.6%
Driver 39	30.2% / 14.0% / 30.2% / 18.6% / 7.0%	20.6% / 9.5% / 22.2% / 28.6% / 19.0%	39.8% / 27.8% / 13.0% / 7.4% / 12.0%
Driver 40	23.9% / 17.4% / 17.4% / 13.0% / 28.3%	22.5% / 20.0% / 7.5% / 22.5% / 27.5%	25.3% / 18.7% / 0.0% / 5.3% / 50.7%

Table 4.10: Results from Split five Segments label

Results Summary:

The results obtained in this section are summarised as follows:

- Applying the threshold over the set of neurons activated by a single dataset results in only the maximally activated points that particular dataset which represent the distinctive patterns for an individual dataset. This means that only the neurons that are matched to the most vectors of the particular dataset will be activated while the rest are deactivated.
- Applying the threshold identifies the nodes which are most frequently activated by each feature that is labeled in the data. This represents a distinctive pattern for an individual feature.
- This set of points obtained after thresholding are then labeled to show which feature they represent.
- After applying the labels, each feature within a label are treated as separate classes and a calculation is done to denote what percentage the set of points relates to the features within each label.
- The above steps are repeated for all the three distractions and their results are shown separately as seen above.

4.5 Driver and Distraction results comparison

In this section each label is broken down into individual classes and comparisons are made between drivers and distractions. Each class (within each label) for individual drivers will be compared with all the distractions at the same time. A measure (the percentage described in the previous section) is required to compare drivers in their

respective maximal sets. As described in section 4.4, each driver has a percentage denoting their characteristics with respect to each class and each distraction.

The main goal of this thesis is to distinguish between drivers and how similar they are relative to the three distractions. To do this, another measure is required that combines drivers from the three distractions (under a single class) to give conclusive results on which distraction is most disruptive or tends to require the most attention. The measure to do this is the percentage of how many drivers are most active for their respective class (within each label). This is given below in equation 4.2:

$$\text{Distraction Percentage \%} = \frac{(\text{Number of drivers most activated})}{(\text{Total number of drivers})} * 100 \quad (4.2)$$

Where:

Number of drivers most activated: The count of drivers who are most activated for a particular distraction under an individual class.

Total number of drivers: The count of drivers under each individual distraction on which the model is trained and this is always 40.

The results begin by first calculating the percentages of each driver under each distraction (from section 4.4) and plotting them on a bar chart based on a single class. Then the number of drivers that are maximally activated for a single class under each distraction are counted. Then each of their percentage is calculated using the formula above in equation 4.2. This percentage is a score of how much drivers under different distractions tend to be activated when analysing each individual class, one at a time. The algorithm for this is described below:

DistractionTally(DriverSet, LabelSet, DriverPercentages):

// **DriverSet:** Set of drivers in the combined dataset
// **LabelSet:** Set of all the labels that each driver is tested on.
// **DriverPercentages:** List containing all the percentages for each class for each driver under all of the three distractions.

```
final_results ← [ ]
for Label in LabelSet:
  for Class in Label:
    hands - free_tally ← 0
    music_tally ← 0
    text_tally ← 0
    class_results ← [ ]
    for Driver in DriverSet:
      percentages ← DriverPercentages[Class][Driver]
      find the maximum percent in percentages
      update the distraction tally having the maximum percent by 1
    end for
    class_percents ← percentages of the distraction counters (equation 4.2)
    for result in class_percents
      add result to class_results list
    end for
  end for
  add class_results to the final_results
end for
return final_results
```

Some important points to note before analysing the results are as follows:

- There are also cases where more than one distraction is maximally activated for a particular class label.
- Some distractions show the same number of drivers being activated. This does not mean that both distractions activate the same set of drivers although there can be overlap between sets of drivers activating different distractions.

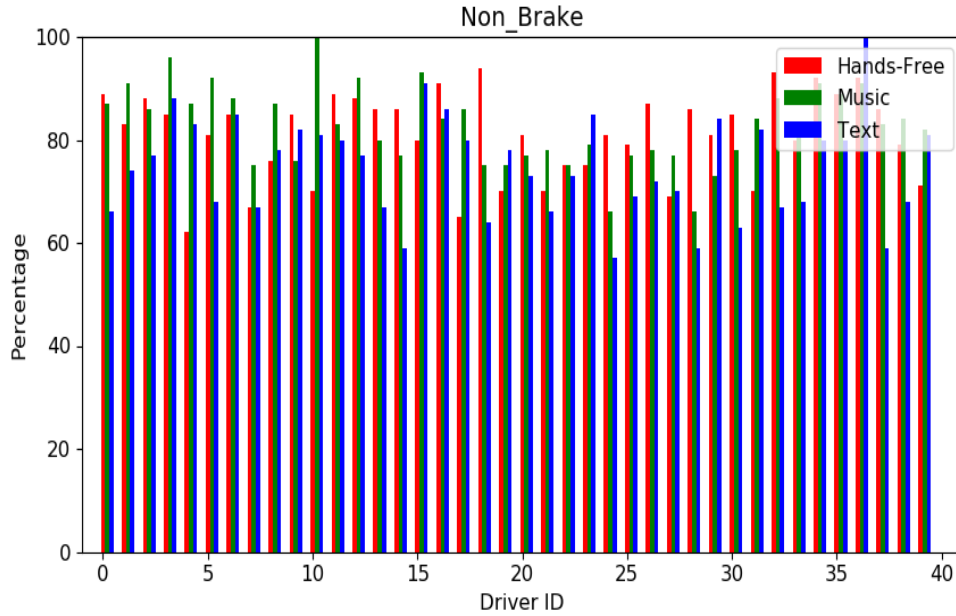


Figure 4.31: Examining the activity of drivers under all the distractions under the non-brake class label.

4.5.1 Individual class analysis

After applying all the labels and calculating their respective percentages a comparison is done between each of the drivers for all three distractions where one class is examined at a time. This is also a way to test specific behaviour of drivers under the three distractions. Figure 4.31 represents the comparison of all the drivers under the three distractions while testing the action of not applying the brake denoted by the non-brake class under the brake pressure label. The music distraction shows that the highest proportion of drivers tend not to apply the brakes very much while the text distraction shows that the lowest proportion of drivers tend not to apply the brakes very much. This supports the claim that texting while driving is the most disruptive distraction as a very small proportion of drivers tend not to apply the brakes very often. Figure 4.32 represents the comparison of all the drivers under

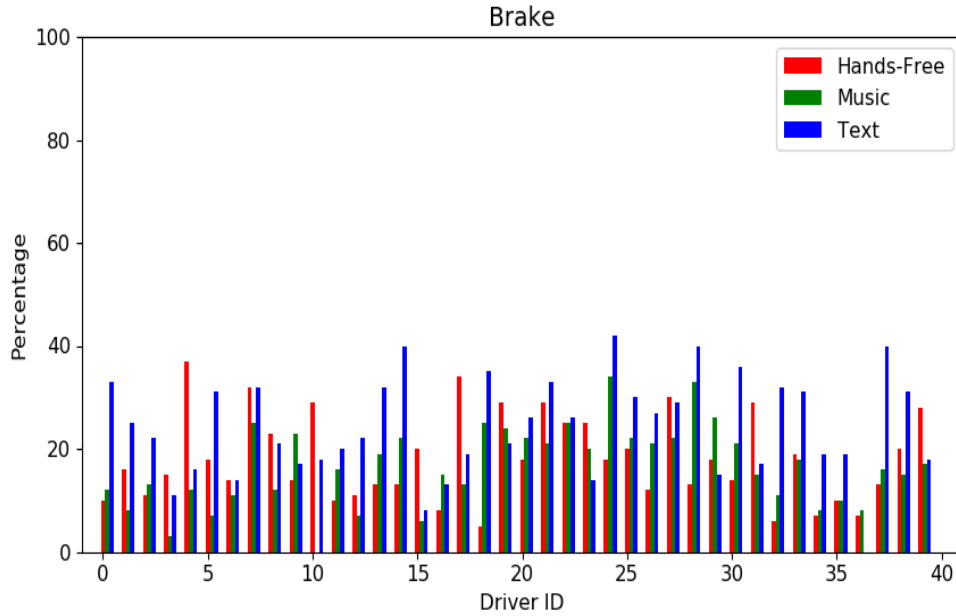


Figure 4.32: Examining the activity of drivers under all the distractions under the brake class label.

the three distractions while testing the action of applying the brake denoted by the brake class under the brake pressure label. The music distraction shows that the lowest proportion of drivers tend to apply the brakes while the text distraction shows that the highest proportion of drivers tend to apply the brakes. This supports the claim that texting while driving is the most disruptive distraction as a very large proportion of drivers tend to frequently apply the brakes.

The other classes are examined in the same way as figures 4.31 and 4.32, however it is found that visualizing them in the form of bar plots or tables does not make it easy to make conclusions among all the drivers for each distraction under each class. This is because not all the bar plots obtained are as clearly distinguishable as figures 4.44 and 4.45. Hence the number of drivers that are most activated under all three distractions are counted based on the distraction separation algorithm specified above

and their respective percentages are recorded below in table 4.11. For reference, the bar plots of all the remaining classes are documented in Appendix C.

No.	Class Name	Hands-Free	Music	Text
1.	Non-Brake	47.5% (19/40)	47.5% (19/40)	10.0% (4/40)
2.	Brake	32.5% (13/40)	10.0% (4/40)	62.5% (25/40)
3.	Below Average Speed	32.5% (13/40)	27.5% (11/40)	45.0% (18/40)
4.	Above Average Speed	45.0% (18/40)	45.0% (18/40)	17.5% (7/40)
5.	Deceleration	32.5% (13/40)	32.5% (13/40)	37.5% (15/40)
6.	Acceleration	37.5% (15/40)	37.5% (15/40)	27.5% (11/40)
7.	Left Acceleration	32.5% (13/40)	45.0% (18/40)	30.0% (12/40)
8.	Right Acceleration	40.0% (16/40)	17.5% (7/40)	45.0% (18/40)
9.	Lower Altitude Acceleration	47.5% (19/40)	30.0% (12/40)	27.5% (11/40)
10.	Higher Altitude Acceleration	17.5% (7/40)	42.5% (17/40)	40.0% (16/40)
11.	Left Center Lane	12.5% (5/40)	72.5% (29/40)	20.0% (8/40)
12.	Right Center Lane	55.0% (22/40)	20.0% (8/40)	27.5% (11/40)
13.	Splits Two 1st	27.5% (11/40)	57.5% (23/40)	17.5% (7/40)
14.	Splits Two 2nd	45.0% (18/40)	22.5% (9/40)	32.5% (13/40)
15.	Splits Three 1st	40.0% (16/40)	17.5% (7/40)	42.5% (17/40)
16.	Splits Three 2nd	35.0% (14/40)	52.5% (21/40)	17.5% (7/40)
17.	Splits Three 3rd	17.5% (7/40)	27.5% (11/40)	57.5% (23/40)

18.	Splits Four 1st	40.0% (16/40)	25.0% (10/40)	35.0% (14/40)
19.	Splits Four 2nd	32.5% (13/40)	57.5% (23/40)	15.0% (6/40)
20.	Splits Four 3rd	52.5% (21/40)	37.5% (15/40)	15.0% (6/40)
21.	Splits Four 4th	12.5% (5/40)	40.0% (16/40)	57.5% (23/40)
22.	Splits Five 1st	30.0% (12/40)	27.5% (11/40)	45.0% (18/40)
23.	Splits Five 2nd	40.0% (16/40)	47.5% (19/40)	27.5% (11/40)
24.	Splits Five 3rd	47.5% (19/40)	50.0% (20/40)	5.0% (2/40)
25.	Splits Five 4th	35.0% (14/40)	35.0% (14/40)	37.5% (15/40)
26.	Splits Five 5th	10.0% (4/40)	27.5% (11/40)	65.0% (26/40)

Table 4.11: Results containing percentages and number of drivers distinctive under every class under each distraction

4.5.2 Results Analysis and Findings

The results in table 4.11 are studied and conclusions are derived on how disruptive each distraction would be depending on each class of behaviour which represent the actions performed by the drivers.

1. Brake Analysis:

Non-Brake class: 47.5% of drivers (19 out of 40) are most distinctive while being distracted by the “hands-free” distraction and the same result is seen while being distracted by the “music” distraction, while drivers under the “text” distraction were the least distinctive where only 10% of the driver patterns (4 out of 40) were maximally activated.

Brake class: 62.5% of drivers (25 out of 40) are most distinctive while under the “text” distraction, 32.5% of drivers (13 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “music” distraction were the least distinctive where only 10% of the driver patterns (4 out of 40) were maximally activated.

Looking into the braking patterns of drivers, it is found that 47.5% of drivers under the distractions of hands-free and music tend not to apply the brake very much and comparatively more than drivers under the text distraction which is about 10% of the drivers. This implies that drivers under text distraction tend to apply the brakes more in comparison to drivers under hands free and music distractions. This in turn indicates that texting is more disruptive than the other distractions. This is further strengthened in the next result where 62.5% of drivers tend to exhibit more frequent braking while under the text distraction. This is much lesser under the hands free distraction with 32.5% and the music distraction with just 10%. This illustrates that listening to music while driving is the least distracting compared with driving while talking on a hands-free device while texting while driving requires the most brake presses and hence it is the most distracting while analysing the braking behaviour.

2. Speed Limits:

Below Average Speed class: 45% of drivers (18 out of 40) are most distinctive while under the “text” distraction, 32.5% of drivers (13 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “music” distraction were the least distinctive where only 27.5% of the driver patterns (13 out of 40) were maximally activated.

Above Average Speed class: 45% of drivers (18 out of 40) are most distinctive

while under the “hands-free” distraction and the same results are seen for the “music” distraction, while drivers under the “text” distraction were the least distinctive where only 17.5% of the driver patterns (7 out of 40) were maximally activated.

When analysing the changing speeds of the drivers, it is found that 32.5% of drivers under the hands-free distraction and 27.5% of drivers under the music distraction exhibit distinctive patterns where drivers drive at speeds below the average speed limit according to their recorded data, while the most distinctive patterns are observed under text distraction with 47.5% of the drivers driving below the average speed.

However when comparing this to patterns observed when driving above the average speed limit, the most distinct patterns found were under the hands-free and music distraction with both being at 45%. The least distinctive patterns over the speed limit were found under the text distraction with 17.5%. Once again the results suggest that texting while driving appears to be most distracting where most drivers are unable to go over the speed limit.

3. Linear Acceleration:

Deceleration class: 37.5% of drivers (15 out of 40) are most distinctive while under the “text” distraction, while 32.5% of drivers (13 out of 40) are distinctive while under the “hands-free” distraction were maximally activated and these are the same results found under the “music” distraction.

Acceleration class: 37.5% of drivers (15 out of 40) are most distinctive while under the “hands-free” distraction and the same results are seen for the “music” distraction, while drivers under the “text” distraction were the least distinctive where only 27.5% of the driver patterns (11 out of 40) were maximally activated.

32.5% of drivers under either the hands-free or the music distraction show patterns where they to decelerate or slow down during the course while 37.5% of the drivers under the text distraction showed patterns containing the action of deceleration. This shows that drivers under text distraction tend to slow down more than while they are under the hands-free and music distractions.

Contradictory to the above results, 37.5% of drivers' patterns under either the distraction of hands-free or music were found to display patterns of acceleration while only 27.5% of drivers displayed similar behaviour under the text distraction, thus showing that drivers listening to music or using a hands-free device recorded patterns containing lesser deceleration but more acceleration behaviour while the inverse is found for drivers who text while driving.

4. Turning Acceleration:

Left Acceleration class: 45% of drivers (18 out of 40) are most distinctive while under the “music” distraction, 32.5% of drivers (13 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “text” distraction were the least distinctive where only 30% of the driver patterns (12 out of 40) were maximally activated.

Right Acceleration class: 45% of drivers (18 out of 40) are most distinctive while under the “text” distraction, 40% of drivers (16 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “music” distraction were the least distinctive where only 17.5% of the driver patterns (7 out of 40) were maximally activated.

Looking into the acceleration behaviour of drivers while turning, distinctive patterns signifying acceleration towards the left are found with 32.5% of drivers being under

the hands-free distraction, 45% of drivers under the music distraction but only 30% of drivers under text distraction. This shows that there was less activity while accelerating to the left (including situation where a left turn is made) under text distraction and most activity under the music distraction.

Looking at accelerating to the right, 40% of drivers under the hands-free distraction, 17.5% drivers under the music distraction and 45% of drivers under the text distraction revealed most distinctive patterns. So most drivers who were distracted by texting showed the most acceleration towards the right while the least distinctive patterns were found by the music distraction. Therefore observing the activity of drivers under music distraction shows that they tend to display patterns of turning towards the left while those under the text distraction show more activity while turning towards the right.

5. Altitude Acceleration:

Lower Altitude Acceleration class: 47.5% of drivers (19 out of 40) are most distinctive while under the “hands-free” distraction, 30% of drivers (12 out of 40) are distinctive while under the “music” distraction, while drivers under the “text” distraction were the least distinctive where only 27.5% of the driver patterns (11 out of 40) were maximally activated.

Higher Altitude Acceleration class: 42.5% of drivers (17 out of 40) are most distinctive while under the “music” distraction, 40% of drivers (16 out of 40) are distinctive while under the “text” distraction, while drivers under the “hands-free” distraction were the least distinctive where only 17.5% of the driver patterns (7 out of 40) were maximally activated.

While analysing the acceleration with respect to the altitude, comparisons are

made as to when the driver goes between lower to higher altitudes and vice versa, most activity is found with drivers under the hands-free distraction at 47.5% drivers, under the music distraction at 30% drivers and the least is found under the text distraction at only 27.5%.

Similarly under the activity of acceleration while moving towards higher altitudes, just 17.5% of drivers under the hands-free distraction were distinguishable but 42.5% of drivers under music distraction and 40% under text distraction were recorded. This is one of the first results to show a clear divide between behaviour between drivers under hands-free and music distractions where drivers were more distracted by using a hands-free device versus listening to music while testing behaviour found while speeding up over higher altitudes.

6. Gap between Lanes:

Left Center Lane class: 72.5% of drivers (29 out of 40) are most distinctive while under the “music” distraction, 20% of drivers (8 out of 40) are distinctive while under the “text” distraction, while drivers under the “hands-free” distraction were the least distinctive where only 12.5% of the driver patterns (5 out of 40) were maximally activated.

Right Center Lane class: 55% of drivers (22 out of 40) are most distinctive while under the “hands-free” distraction, 27% of drivers (11 out of 40) are distinctive while under the “text” distraction, while drivers under the “music” distraction were the least distinctive where only 20% of the driver patterns (8 out of 40) were maximally activated.

Looking into how drivers drove within their respective lane, it is found that most drivers up to 72.5% tended to drive closer to the left side of the lane while listening

to music while results under the hands-free and music distractions revealed far lesser scores of 12.5% and 20% respectively. This shows that while listening to music, most drivers drive closer the left of the lane.

However while under hands-free distraction, 55% of drivers tended to drive closer to the right of the lane, while only 20% and 27.5% tended to do the same under the music and text distraction. Once again a big difference is seen between the drivers under the music and hands-free distractions. The results show that driver distracted by music drive more to the left of the lane while those distracted by hands-free drive more to the right of the lane.

7. Split into two segments:

Two split first section class: 57.5% of drivers (23 out of 40) are most distinctive while under the “music” distraction, 27.5% of drivers (11 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “text” distraction were the least distinctive where only 17.5% of the driver patterns (7 out of 40) were maximally activated.

Two split second section class: 45% of drivers (18 out of 40) are most distinctive while under the “hands-free” distraction, 32.5% of drivers (13 out of 40) are distinctive while under the “text” distraction, while drivers under the “music” distraction were the least distinctive where only 22.5% of the driver patterns (9 out of 40) were maximally activated.

When dividing the dataset into two parts, the starting half and the ending half, it is found that while under the hands-free distraction, the drivers’ patterns are less distinct in the first half with only 27.5% activity and are more distinct in the ending half with 45% activity. For the music distraction, in the starting half, drivers exhibit

most activity with 57.5% and least activity in the ending half with 22.5%. Finally for the text distraction, 17.5% driver patterns were distinct in the starting half while 32.5% were more distinct in the ending half. Most distinct patterns were located at the ending, starting and ending halves for the hands-free, music and text distractions respectively.

7. Split into three segments:

Three split first section class: 42.5% of drivers (17 out of 40) are most distinctive while under the “text” distraction, 40% of drivers (16 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “music” distraction were the least distinctive where only 17.5% of the driver patterns (7 out of 40) were maximally activated.

Three split second section class: 52.5% of drivers (21 out of 40) are most distinctive while under the “music” distraction, 35% of drivers (14 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “text” distraction were the least distinctive where only 17.5% of the driver patterns (7 out of 40) were maximally activated.

Three split third section class: 57.5% of drivers (23 out of 40) are most distinctive while under the “text” distraction, 27.5% of drivers (11 out of 40) are distinctive while under the “music” distraction, while drivers under the “hands-free” distraction were the least distinctive where only 17.5% of the driver patterns (7 out of 40) were maximally activated.

Similar to the previous split, this time the dataset is divided into three parts being the start, middle and end parts. It is found that under the hands-free distraction, 40% of the driver patterns are active in the first part, 35% of the patterns

are active in the middle part and 17.5% are active at the end. Under the music distraction, 17.5% of drivers show patterns active in the first section, 52.5% in the middle section and 27.5% show active patterns in the last section. Under the text distraction, 42.5% of drivers show active patterns in the first section, 17.5% in the middle section and 52.5% in the ending section. This shows that drivers using hands-free devices tend to be more active towards the starting and middle of their drive, while listening to music were more active at the middle of their drive while those texting were more active at the first and last part of their respective drives.

8. Split into four segments:

Four split first section class: 40% of drivers (16 out of 40) are most distinctive while under the “hands-free” distraction, 35% of drivers (14 out of 40) are distinctive while under the “text” distraction, while drivers under the “music” distraction were the least distinctive where only 25% of the driver patterns (10 out of 40) were maximally activated.

Four split second section class: 57.5% of drivers (23 out of 40) are most distinctive while under the “music” distraction, 32.5% of drivers (13 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “text” distraction were the least distinctive where only 15% of the driver patterns (6 out of 40) were maximally activated.

Four split third section class: 52.5% of drivers (21 out of 40) are most distinctive while under the “hands-free” distraction, 37.5% of drivers (15 out of 40) are distinctive while under the “music” distraction, while drivers under the “text” distraction were the least distinctive where only 15% of the driver patterns (6 out of 40) were maximally activated.

Four split fourth section class: 57.5% of drivers (23 out of 40) are most distinc-

tive while under the “text” distraction, 40% of drivers (16 out of 40) are distinctive while under the “music” distraction, while drivers under the “hands-free” distraction were the least distinctive where only 12.5% of the driver patterns (5 out of 40) were maximally activated.

In this split, each dataset is divided and labeled by four equal sections, the first part, the second part, the third part and the ending part. Under the hands-free distraction, driver patterns in the first part were 40% active, in the second part were 32.5% active, in the third part were 52.5% and finally in the last part were 12.5% active, thus showing that under the hands-free distraction, drivers are most active at the third part. Under the music distraction, active driver patterns were found by 25% drivers in the first part, by 57.5% drivers in the second part, by 37.5% of drivers in the third part and by 40% of drivers in the last part. Under the text distraction, active patterns in the first part were found by 35%, 15% in both the second and third parts and 57.5% in the last part. This shows that drivers under hands-free distraction are most active or distinct at the third part, under music distraction were most distinct in the second part and under the text distraction were most active in the last part.

9. Split into five segments:

Five split first section class: 45% of drivers (18 out of 40) are most distinctive while under the “text” distraction, 30% of drivers (12 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “music” distraction were the least distinctive where only 27.5% of the driver patterns (11 out of 40) were maximally activated.

Five split second section class: 47.5% of drivers (19 out of 40) are most distinct-

tive while under the “music” distraction, 40% of drivers (16 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “text” distraction were the least distinctive where only 27.5% of the driver patterns (11 out of 40) were maximally activated.

Five split third section class: 50% of drivers (20 out of 40) are most distinctive while under the “music” distraction, 47.5% of drivers (19 out of 40) are distinctive while under the “hands-free” distraction, while drivers under the “text” distraction were the least distinctive where only 5% of the driver patterns (2 out of 40) were maximally activated.

Five split fourth section class: 37.5% of drivers (15 out of 40) are most distinctive while under the “text” distraction, while drivers under the “hands-free” distraction were the least distinctive where 35% of the driver patterns (14 out of 40) were maximally activated and these results are the same as those found while under the music distraction.

Five split fifth section class: 65% of drivers (26 out of 40) are most distinctive while under the “text” distraction, 27.5% of drivers (11 out of 40) are distinctive while under the “music” distraction, while drivers under the “hands-free” distraction were the least distinctive where only 10% of the driver patterns (4 out of 40) were maximally activated.

In this final split, each dataset is divided and labeled by five equal sections, the first part, second, third fourth and the fifth which is the last part. Under the hands-free distraction, driver patterns in the first part were 30% active, in the second part were 40% active, in the third part were 47%, in the fourth part were 35% and finally in the last part were 10% active, thus showing that under the hands-free distraction, drivers are most active at the third part. Under the music distraction, active driver

patterns were found by 27.5% drivers in the first part, by 47.5% drivers in the second part, by 50% of drivers in the third part, by 35% in the fourth part and by 27.5% of drivers in the last part. Under the text distraction, active patters in the first part were found by 45%, 27.5% in the second part, 5% in the third part, 37.5% in the fourth part and 65% in the last part. This shows that drivers under hands-free distraction are most active or distinct at the third part, under music distraction were once again more active in the third part and under the text distraction were most active in the last part.

Chapter 5

Conclusions & Future Work

The ultimate chapter briefly encompasses all the results of this research described in the previous chapter and what was learned from it including brief descriptions on the SOM architecture and the way the model is trained. After summarising the results, the different ideas to further build on the existing work are proposed while briefly describing the potential for future growth.

5.1 Conclusions

Training the SOM over time series datasets and analysing the features one at a time shows that the SOM is able to understand the structure of a time series dataset. By applying a window of size five over the dataset and then sliding the window one step at a time over the dataset in order to provide context to each individual event (vector in the dataset), the SOM also learns contextual information, once again by the numerous feature analysis shown in section 4.1 of the fourth chapter.

In order to study multiple drivers' behaviour in regards to understanding both

the differences between individual drivers and the same driver under different distractions, the SOM was trained on the first four driver datasets which also included the same drivers under different distractions as seen in section 4.2. The results once again showed how the SOM adapted to the datasets with similar feature analysis. It was also found in feature analysis that some features were stronger than others in that they activated the SOM more than other features. For individual drivers, the SOM showed many smaller clusters. The smaller clusters were usually not near each other which in turn shows how unique each driver is in regards to the action they perform. This shows that a SOM trained with the first four drivers represents the specific attributes of each dataset and this is why the SOM is used to pick out distinct driver patterns. Tests done on analysing the patterns of distractions on the map showed that the “music” distraction was the most visible because of having the most clusters closest to each other, followed by the “text” distraction and lastly the “hands-free” distraction.

Finally the SOM is trained on all 40 with each of the three distractions. Once again similar tests were done through feature analysis where the features that most activated the SOM were further strengthened while the effect of other features is further weakened. This highlighted the gap between stronger and weaker features which can be seen in section 4.3 of the fourth chapter.

The most significant result was found while identifying each driver in the map using the driver ID to label the clusters in the map. The focus is on differentiating between different drivers was difficult due to the overlap between used ID labels. To reveal only the distinct patterns of each driver and each dataset, a threshold was applied on the map such that each node only has one label which is the driver ID

that most often activates that particular neuron in the map. This process is then applied in such a way that each of the three datasets is used to activate to the map followed by thresholding it to obtain a set of points that activate and most represent each particular dataset. This process is repeated for all the 120 datasets (40 drivers under three distractions) such that each dataset now has their own set of maximally activated points which is seen in at the beginning of section 4.4.

The threshold results in each driver having their own set of points (that indicate distinct patterns) but they were so sparse and varied that almost none of the drivers had their own large clusters forming out of individual smaller clusters. The datasets of the same drivers but under different distractions did not form larger clusters nor did their patterns overlap in most cases, all the results of which are documented in section 4.4. Numerous attempts were made to cause larger clusters of drivers to form by augmenting the dataset and training a new model, however they affected the natural structure created by the SOM by adding a layer of complexity and by adversely affecting the strength of features already present.

Applying the threshold ensures that distinct patterns of clusters for each driver under different distractions are obtained. It is on these distinct patterns that further analysis can occur since they are representative of each dataset, which can be seen at the beginning of section 4.4. Previous results have shown that the SOM adapts differently to all the features it is trained on. To understand the significance of each set of patterns, multiple labels are applied, that features of the drivers and labels that represent the activity of the drivers in different sections of the course.

When the labels are applied to the patterns, percentages are calculated based on how

much a pattern relates with each class within a label. These percentages represent a score of how much a particular dataset relates to each class which represents of specific actions which indicate a drivers' behaviour at any given time. This process is then repeated for all the labels which are seen in section 4.4 and then the percentages of how much each pattern belongs to a class, are calculated. This process is repeated for all the driver patterns that are obtained after applying the threshold. All the percentages of drivers based on the action performed (each class) thus obtained now denote the relationship between each class and how much a particular driver's pattern tends to relate to it. Each driver has a percentage for each class which can now be compared with other drivers and classes to identify how the drivers relate the most or the least to an individual class. All of the results are seen in section 4.4.

The percentages obtained serve as measures to compare drivers under different distractions. This is seen in section 4.5 where the percentages for all the drivers a comparison is made of all drivers under each of the three distractions. This comparison is represented as a percentage for each distraction that contains the number of drivers that are most active (having the highest percentage) for each distraction under a specific class. Once these are computed (as seen in table 4.1) the following conclusions are made for the distinct driver patterns under each distraction:

- Drivers distracted by texting tend to use the brake more often and the least non-braking activity. Under the music distraction, the fewest number of drivers apply the brakes and also show the most non-braking behaviour which is similar to the hands-free distraction.
- Drivers under the text distraction had most of their recorded speeds below the average and the least above the average. Similarly under the music distraction

drivers recorded higher speeds, mostly above the average and the least number of drivers displayed speeds below the average.

- Drivers under the text distraction showed most deceleration behaviour and then least acceleration behaviour. While under hands-free and music distractions drivers showed least deceleration behaviour and the most acceleration behaviour.
- Drivers under the music distraction showed the most acceleration towards the left while those under the hands-free and text distraction displayed more behaviour in accelerating towards the right.
- Drivers under the hands-free distraction had the most acceleration from lower to higher altitudes while those under the music and text distraction showed the most behaviour in accelerating from a higher to lower altitudes.
- Drivers under the music distraction showed that most of them tended to drive to the left of the center of the lane while most drivers under the hands-free distraction tend to driver closer to the right of the center of the lane.
- When examining driver's data while dividing them into two parts – a starting and an ending half. Drivers under the music distraction had the most activity at the starting half of the drive while under the hands-free distraction, the most activity is recorded in the ending half.
- When examining driver's data while dividing them into three parts – a starting part, a middle portion and an ending portion. Drivers under the hands-free and text distractions showed the most activating in the starting portion. Drivers under the music distraction showed the most activity in the middle portion and drivers under the text distraction showed the most activity in the ending portion.

- When examining driver's data while dividing them into four parts – a starting portion, a second portion, a third portion and finally the ending portion. Drivers under the hands-free distraction showed the most activity in the starting portion. Drivers under the music distraction showed the most activity in the second portion and drivers under the hands-free distraction showed the most activity in the third portion. Finally drivers under the text distraction showed the most activity in the last portion.
- When examining driver's data while dividing them into five parts – a starting portion, a second portion, a third portion, a forth portion and finally the ending portion. Drivers under the text distraction showed the most activity in the starting portion. Drivers under the music distraction showed the most activity in both the second and third portions. Drivers under the all the distractions showed almost equal activity in the forth portion. Finally drivers under the text distraction showed the most activity in the last portion.

5.2 Future Work

Numerous other interesting ideas that build upon the existing work of finding distinctive patterns using SOMs came up during the course of this research. This section briefly discusses those ideas and other possible routes that the research could have taken in terms of extending the applications of the SOM beyond its usage as a tool of visualization and dimensionality reduction.

- One of the main contributions is the set of maximal points that indicate the distinct patterns of each driver under different distractions. While different percentages were used to understand the value of these distinct points, other metrics or analysis could be used that would further bring out their value.

- In addition to the existing three distraction datasets for each individual driver, if there are experiments resulting in datasets of the same drivers not under any distractions, comparisons could be made as to when the drivers actually become distracted during the course of the experiment.
- When labeling these distinct patterns, the labels used could be further modified as required in the following ways:
 - Analyse the data within the each specific distinct pattern to understand more about driver behaviour and generate more specialized labels to test them on the map.
 - Experiment with labels of more than two classes would increase the possibility of detecting both general and specific patterns within the driver patterns.
 - Expanding on the previous point to build a multi class classifier and test the map on drivers that have not yet been tested. This can further be expanded by either comparing or even augmenting the SOM with any classification algorithms.
- Using the existing labels (or even labels customized as mentioned in the previous point) to create a vector for each driver and dataset that would serve as a unique identifier. Provided the labels applied are well defined, this unique identifier can be used to provide a description of a particular driver under any specific distraction.
- To get a more generalized understanding of driver traits and characteristics, it would have been valuable to have a larger number of drivers in the overall dataset. This would have also enabled possible cases of using training and

testing sets for making predictions (explained further in the next two points) and in making a comparatively more accurate consensus.

- Given enough datasets, it would be interesting create two SOM models, one trained on all the datasets and the other on most of the datasets. Then both the models are tested on the datasets that the latter model is not trained on. The results of the predictions are compared to examine if the results in the latter SOM is capable of generalizing and can predict driver patterns on data that it is not trained on.
- Once a SOM is sufficiently trained on numerous drivers, techniques such as Linear Vector Quantization could be applied to generate vectors closest to the ones used to train the model and these vectors could be tested to understand if they replicate driver behaviour at a very small scale since the result are limited by the six features that the model is trained on.

References

1. Chang, Kuang-Hua; “e-Design”, Part II – Product Performance Evaluation, Chapter 8 – Motion Analysis, pp. 391 - 395, 2015, ISBN: 978-0-12-382038-9.
2. Václav, Linkov; Petr, Zámečník; Aleš, Zaoral; “The use of Driving Simulators in Psychological Research”, Siberian Psychological Journal (SPJ), No. 64, pp. 65 – 75, 2017, DOI: 10.17223/17267080/64/4.
3. Nowosielski, Robert J.; Trick, Lana M.; Toxopeus, Ryan; “Good Distractions: Testing the effects of listening to an Audiobook on Driving Performance in Simple and Complex Road Environments”. Accident Analysis and Prevention 111, pp. 202 – 209, 2018.
4. Lochner, Martin J.; Trick, Lana M.; “Multiple-Object Tracking while Driving: The Multiple-Vehicle Tracking Task”, Attention, Perception & Psychophysics - 76, pp. 2326 – 2345, 2014, DOI: 10.3758/s13414-014-0694-3.
5. Rodd, Heather E.K.; Trick, Lana M.; “Predictors of Mind-Wandering While Driving”, 9th International Driving Symposium on Human Factors in Driver Assessment, Training and Vehicle Design, pp. 326 – 332, 2017, DOI: 10.17077/drivingassessment.1654.
6. Ramkhalawansingh, Robert C.; Trick, Lana M.; “Too Close for Comfort:

- Evaluating a Reward-Based Approach to Increase Drivers' Headway", 8th International Driving Symposium on Human Factors in Driver Assessment, Training and Vehicle Design, pp. 226–232, 2015, DOI: 10.17077/drivingassessment.1576.
7. Mueller, Alexander S.; Trick, Lana M.; "Effects of driving experience on change detection based on target relevance and size", PROCEEDINGS of the Seventh International Driving Symposium on Human Factors in Driver Assessment, Training, and Vehicle Design, pp. 341-347, 2013.
 8. Haykin, Simon; "Neural Networks and Learning Machines, 3rd edition", Pearson Education Inc., 2009.
 9. Hebb, Donald O.; "The Organization of Behavior – A Neuropsychological Theory", New York: Wiley & Sons, 1949.
 10. Hinton, Geoffrey E.; Sejnowski, Terrence J.; "Unsupervised Learning: Foundations of Neural Computation", The MIT Press, Cambridge, Massachusetts, pp. 398 – 408, 1999, ISBN 0-262-58168-X.
 11. Neilson, Michael; "Neural Networks and Deep Learning", Chapter-4, <http://neuralnetworksanddeeplearning.com>, 2019.
 12. Simpson, Patrick, K.; "Artificial Neural Systems: Foundations, Paradigms, Applications and Implementations" Pergamon Pr, 1st edition, 1990, ISBN-13: 978-0080378947.
 13. Sejnowski, Terrence J., Roseberg, Charles R.; "NETtalk: A Parallel Network that Learns to Read Aloud", Neuro Computing: Foundations of Research, pp. 661-672, 1988, ISBN:0-262-01097-6.

14. Rumelhart David E.; Hinton, Geoffrey E.; Williams, Ronald J.; "Learning Representations by Back-Propagating Errors", *Nature* Vol. 323, pp. 533-536, 1986.
15. LeCun, Y., Boser, B., Denker, J. S., Henderson, D., Howard, R. E., Hubbard, W., Jackel, L. D.; "Backpropagation Applied to Handwritten Zip Code Recognition", *Neural Computation*, Vol. 1, Issue 4, pp. 541-551, 1989.
16. Grossberg, S.; "Adaptive Pattern Classification and Universal Recoding, II: Feedback, Expectation, Olfaction, and Illusions", *Biological Cybernetics*, 23, 187-202, 1976b.
17. Kohonen, Teuvo; "Self-Organized formation of Topologically Correct Maps", *Biological Cybernetics*, Volume 43, pp 59-69, 1982.
18. Kohonen, Teuvo; "The self-organized map", *Proceedings of the IEEE*, Vol. 78, pp. 1464-1480, 1990.
19. Hopfield, J. J., Tank, D. W.; "Neural Computation of Decisions in Optimization Problems", *Biological Cybernetics*, Vol. 52, pp. 141-152, 1985, DOI: 10.1007/BF00339943.
20. Barreto, Guilherme A.; "Time Series Prediction with the Self-Organizing Map: A Review", *Studies in Computational Intelligence*, pp. 135-158, 2007.
21. Lang, Kevin J., Waibel, Alex H., Hinton, Geoffrey E.; "A Time-Delay Neural Network Architecture for Isolated Word Recognition", *Journal - Neural Networks*, Vol. 3 Issue 1, pp. 23-43, 1990.
22. Henderson, Paul; "Sammon Mapping", 2010.

23. McCulloch, Warren S., Pitts, Walter; "A Logical Calculus of the Ideas Immanent in Nervous Activity", Bulletin of Mathematical Biophysics, Vol. 5, pp. 115-133, 1943.
24. le Cun, Yann; "A Theoretical Framework for Back-Propagation". Proceedings of the 1988 Connectionist Models Summer School, pp 21-28, 1988.
25. Rosenblatt, Franklin; "The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain", Psychological Review, Vol. 65, No. 6, 1958.

Appendices

A Splitting driver 1's dataset into segments

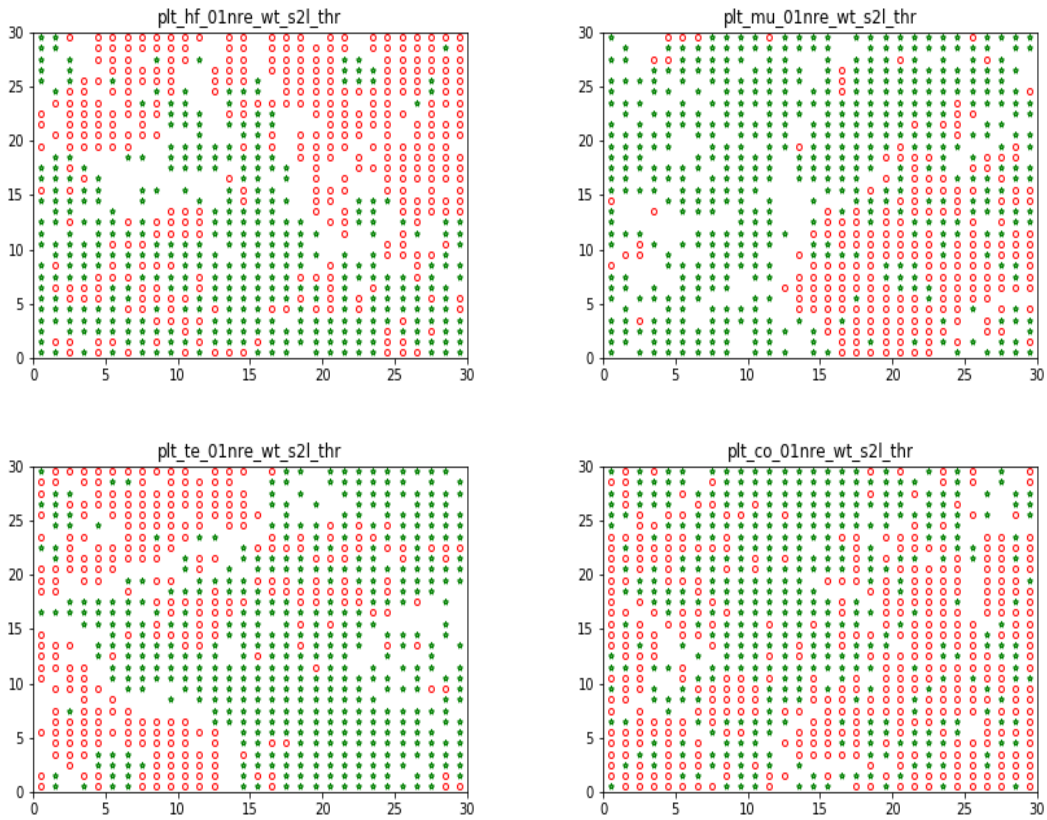


Figure A.1: Examining splitting the dataset of driver 1 into two segments

- (a) Hands-Free distraction (top-left), (b) Music distraction (top-right), (c) Text distraction (bottom-left) and (d) combination of datasets (bottom-right).

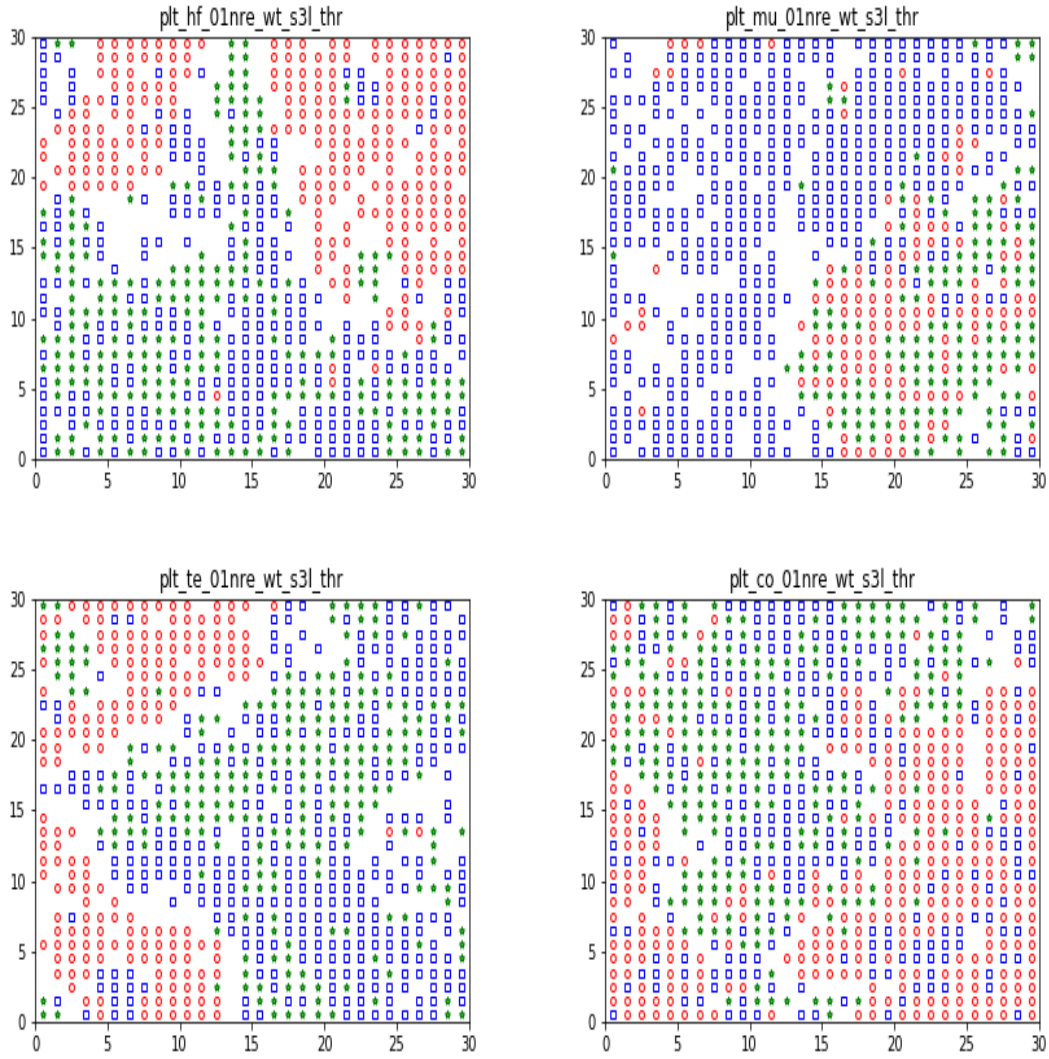


Figure A.2: Examining splitting the dataset of driver 1 into three segments
 (a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
 (c) Text distraction (bottom-left) and (d) combination of datasets (bottom-right).

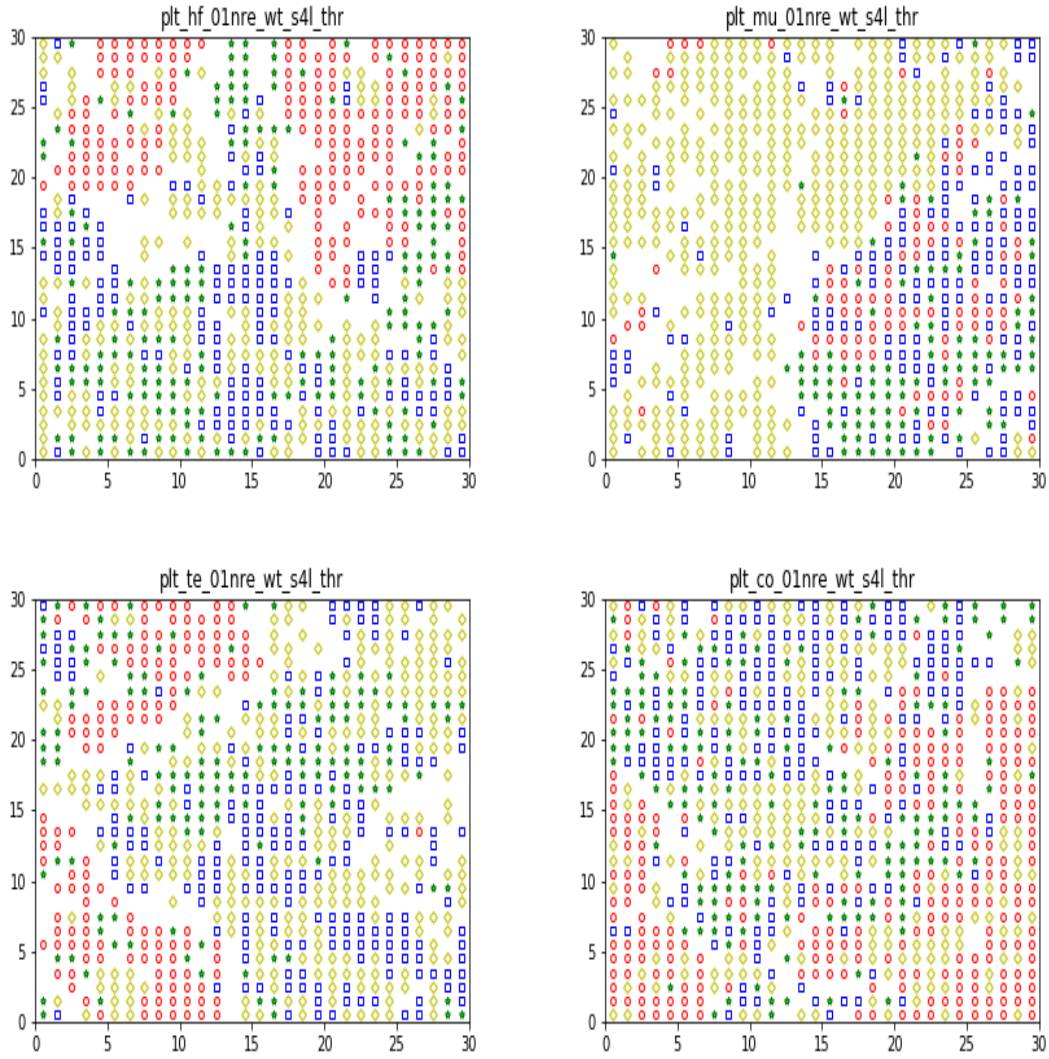


Figure A.3: Examining splitting the dataset of driver 1 into four segments
 (a) Hands-Free distraction (top-left), (b) Music distraction (top-right),
 (c) Text distraction (bottom-left) and (d) combination of datasets (bottom-right).

B Driver and distraction analysis by labeling



Figure B.1: Examining the braking behaviour by separating the classes within the brake pressure label based on percentages
(a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

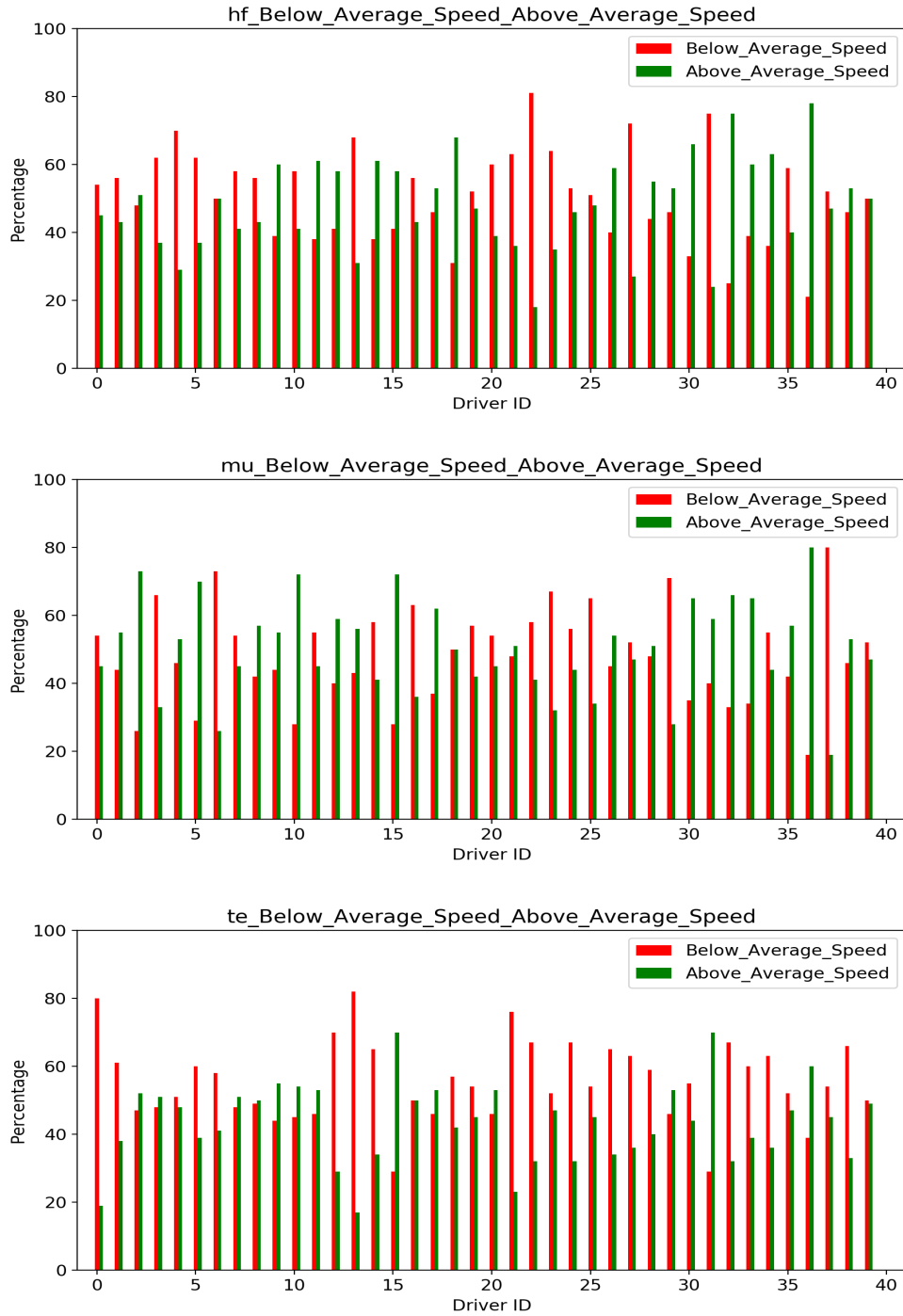


Figure B.2: Examining the speeding behaviour by separating the classes within the speed limits label based on percentages
 (a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

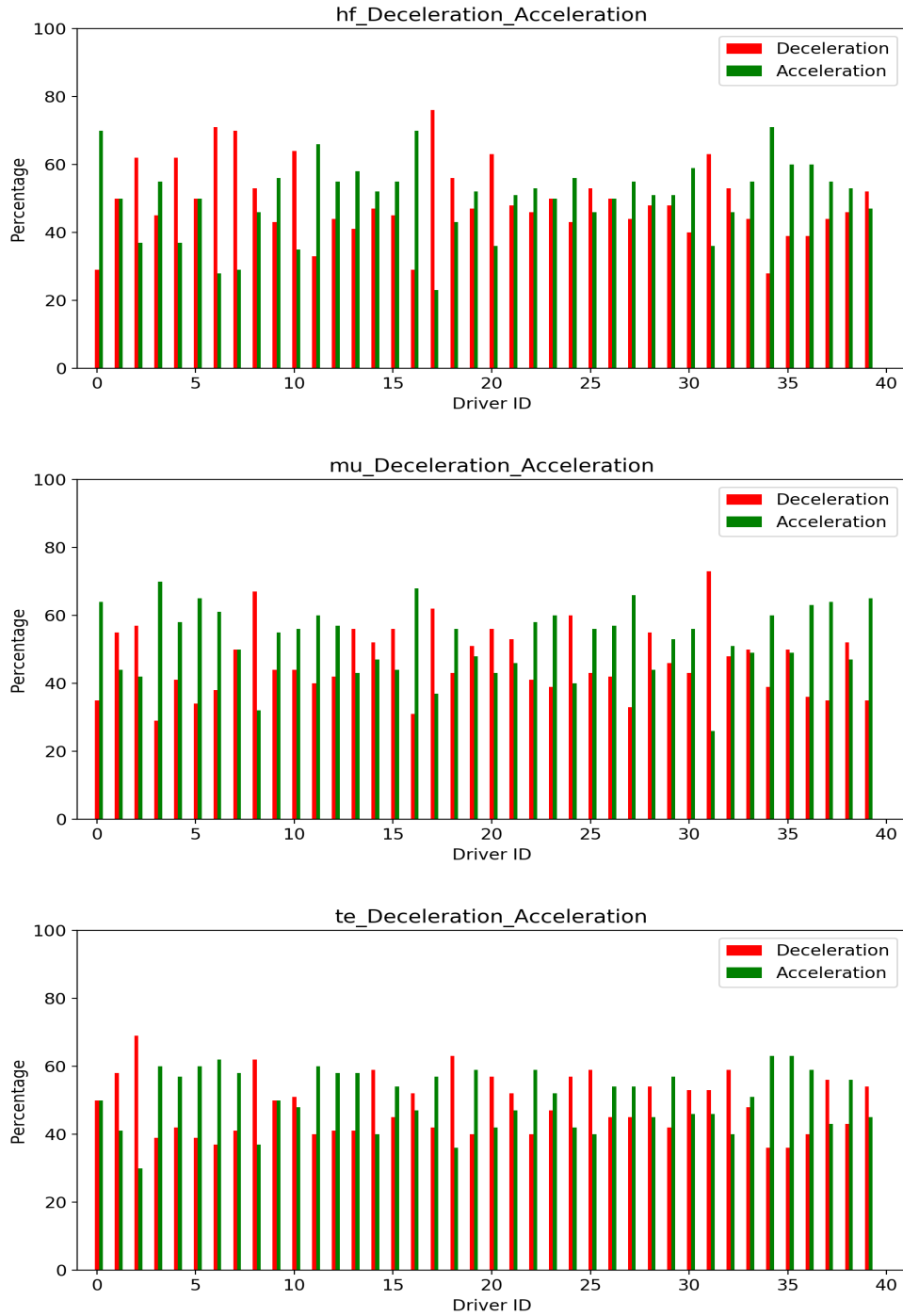


Figure B.3: Examining the acceleration along X axis behaviour by separating the classes within the linear acceleration label based on percentages
(a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

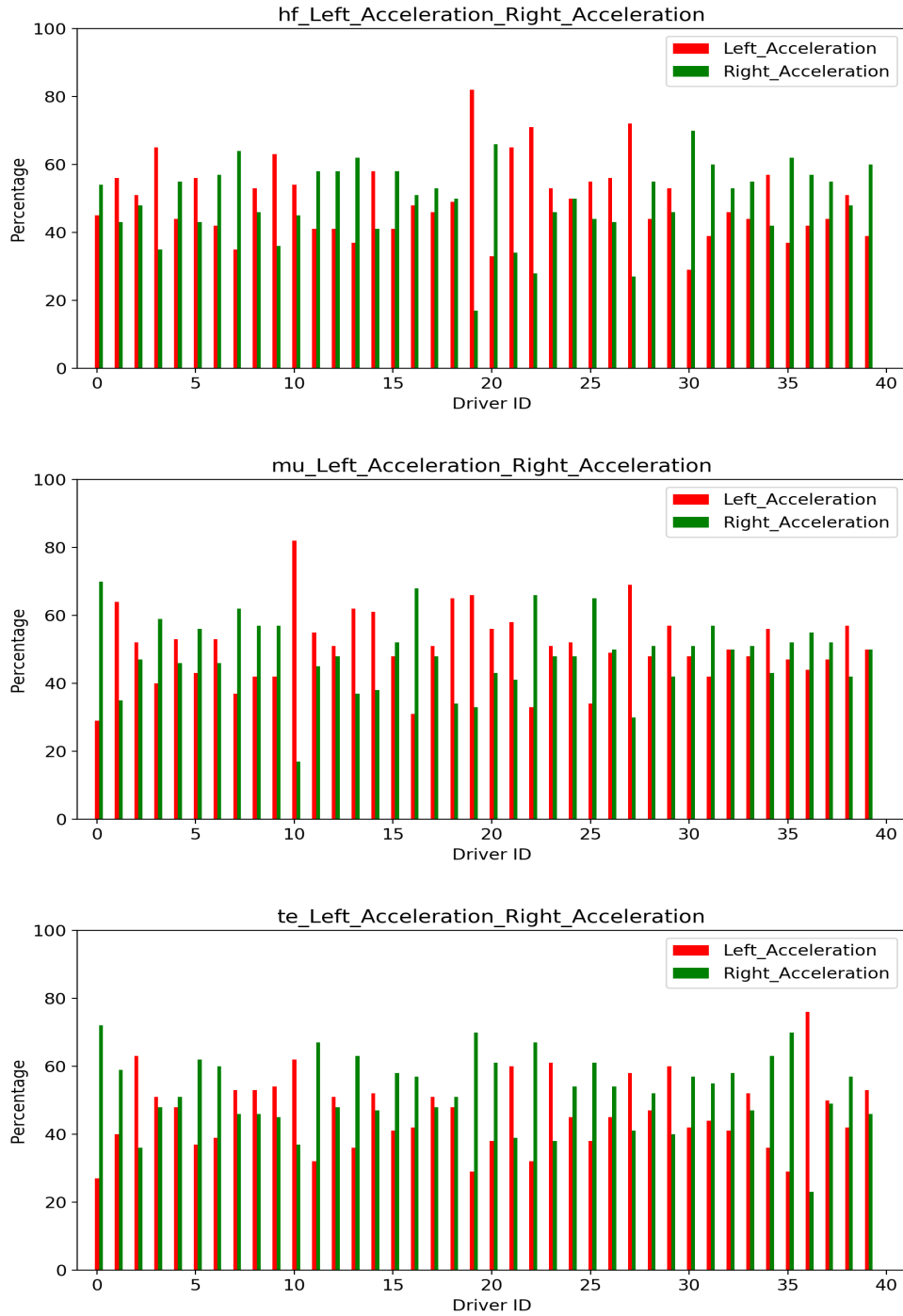


Figure B.4: Examining the acceleration along Y axis behaviour by separating the classes within the turning acceleration label based on percentages
 (a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

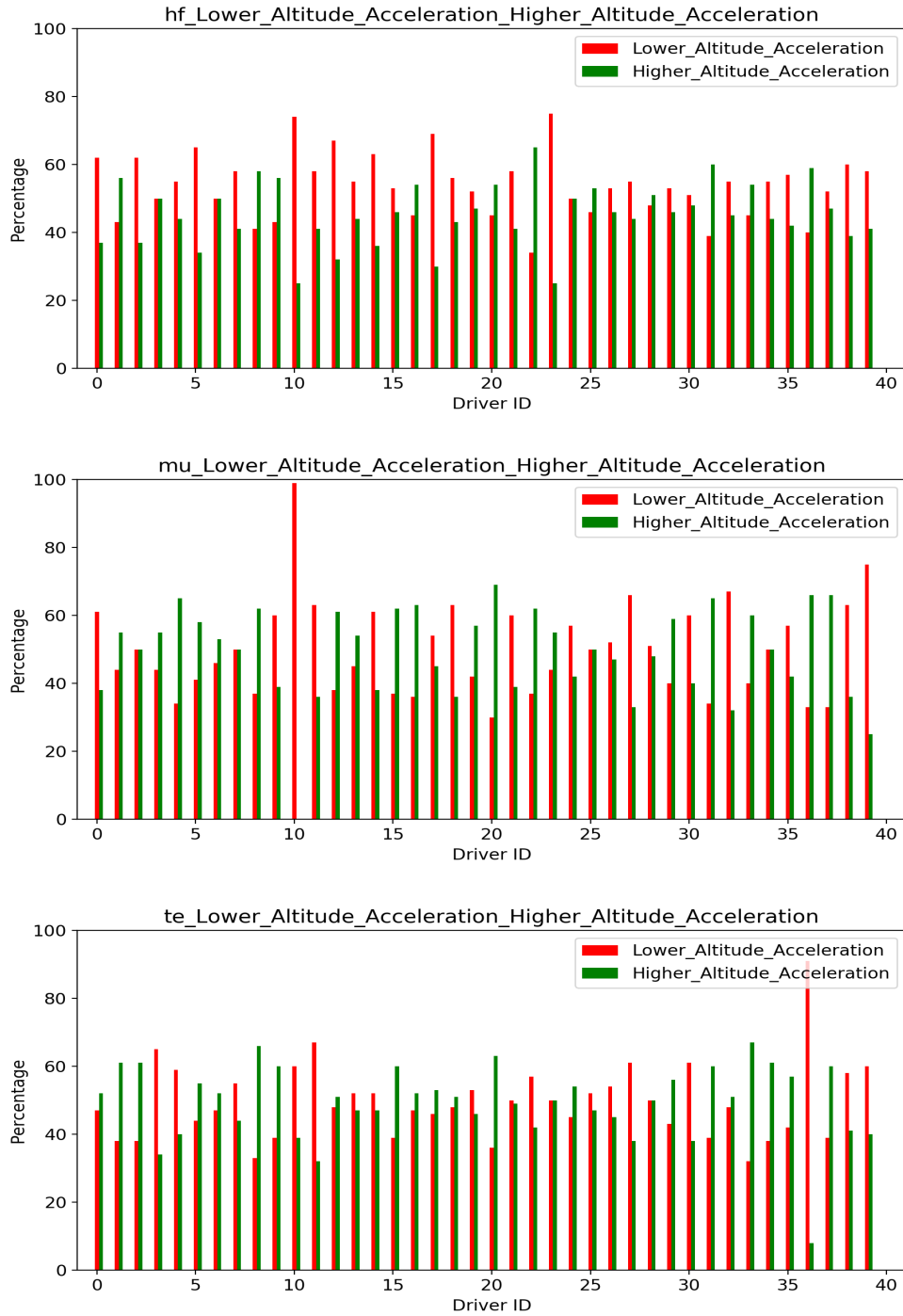


Figure B.5: Examining the acceleration along Z axis behaviour by separating the classes within the altitude acceleration label based on percentages
(a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

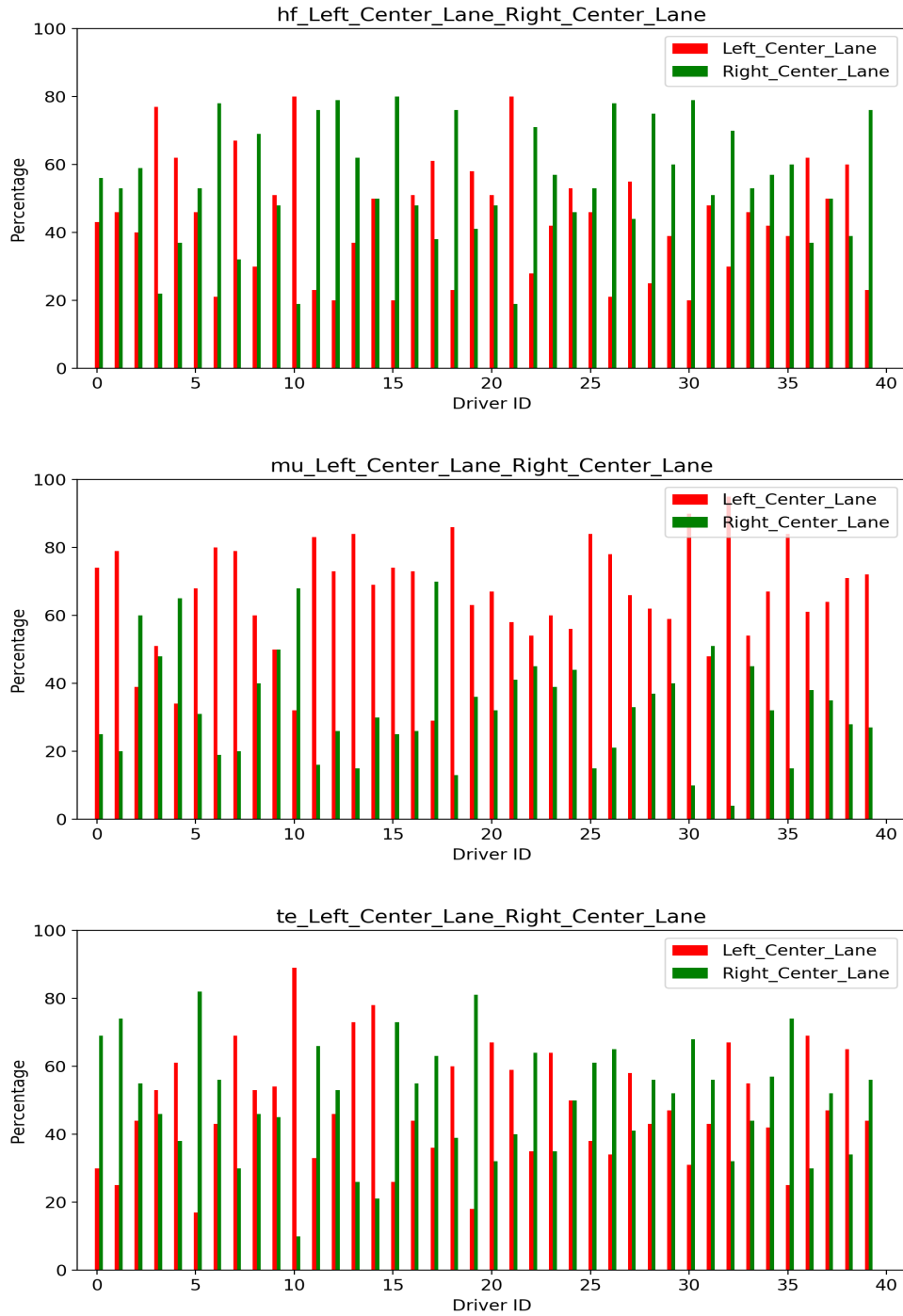


Figure B.6: Examining the lane gap behaviour by separating the classes within the gap between lanes label based on percentages
 (a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

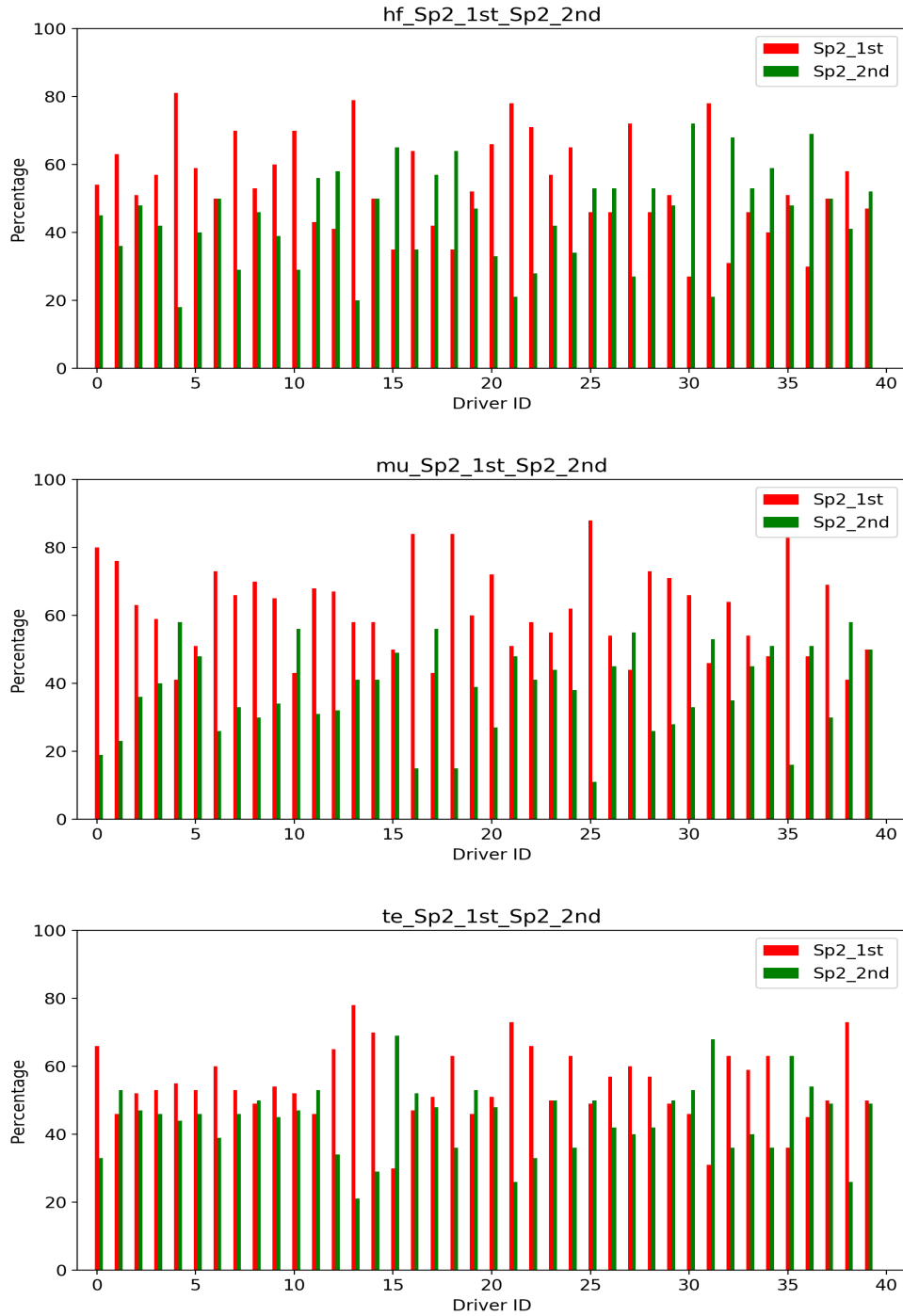


Figure B.7: Examining the activity of drivers under different section of their dataset by separating the classes within the split into two sections label based on percentages
(a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

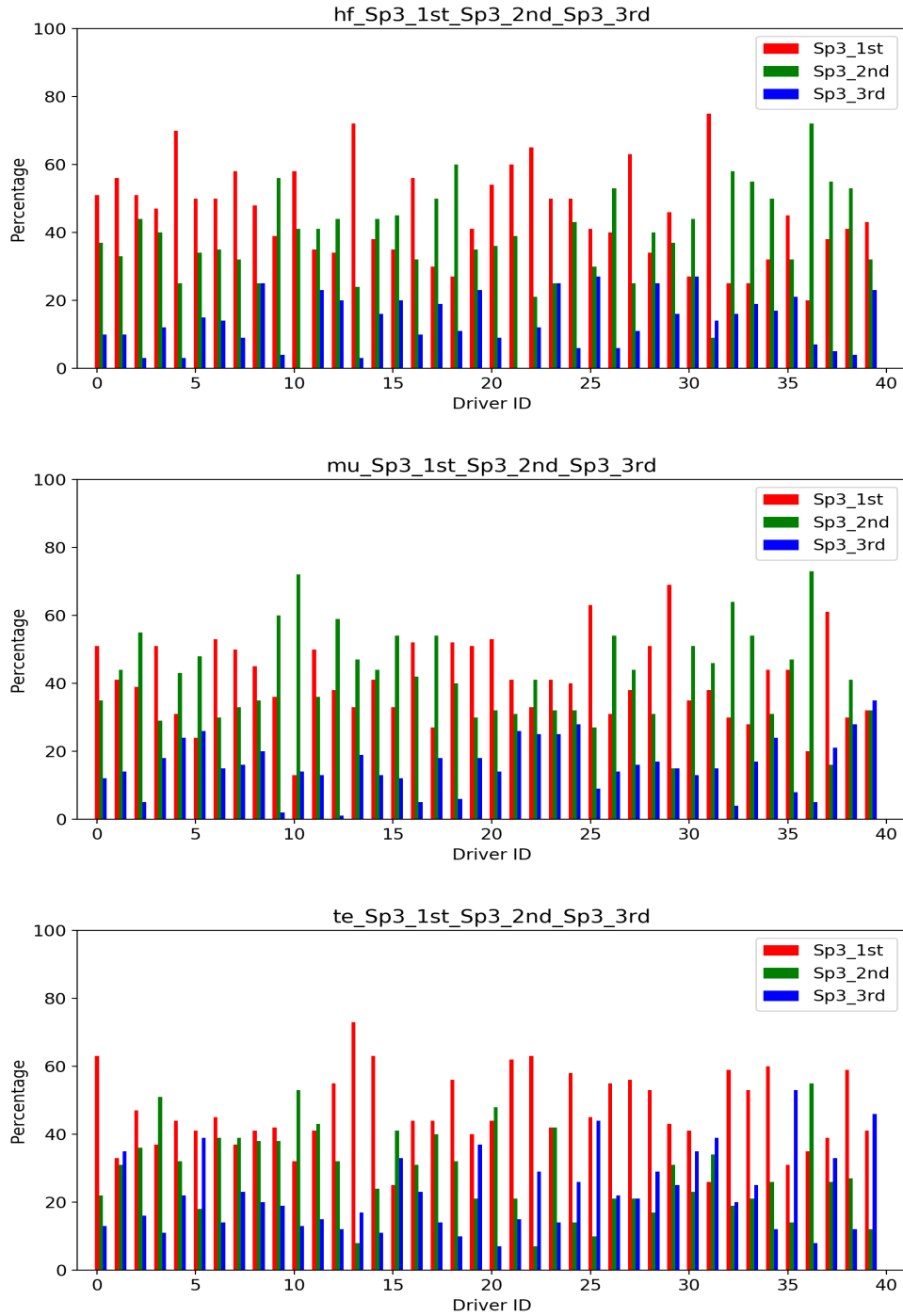


Figure B.8: Examining the activity of drivers under different section of their dataset by separating the classes within the split into three sections label based on percentages
 (a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

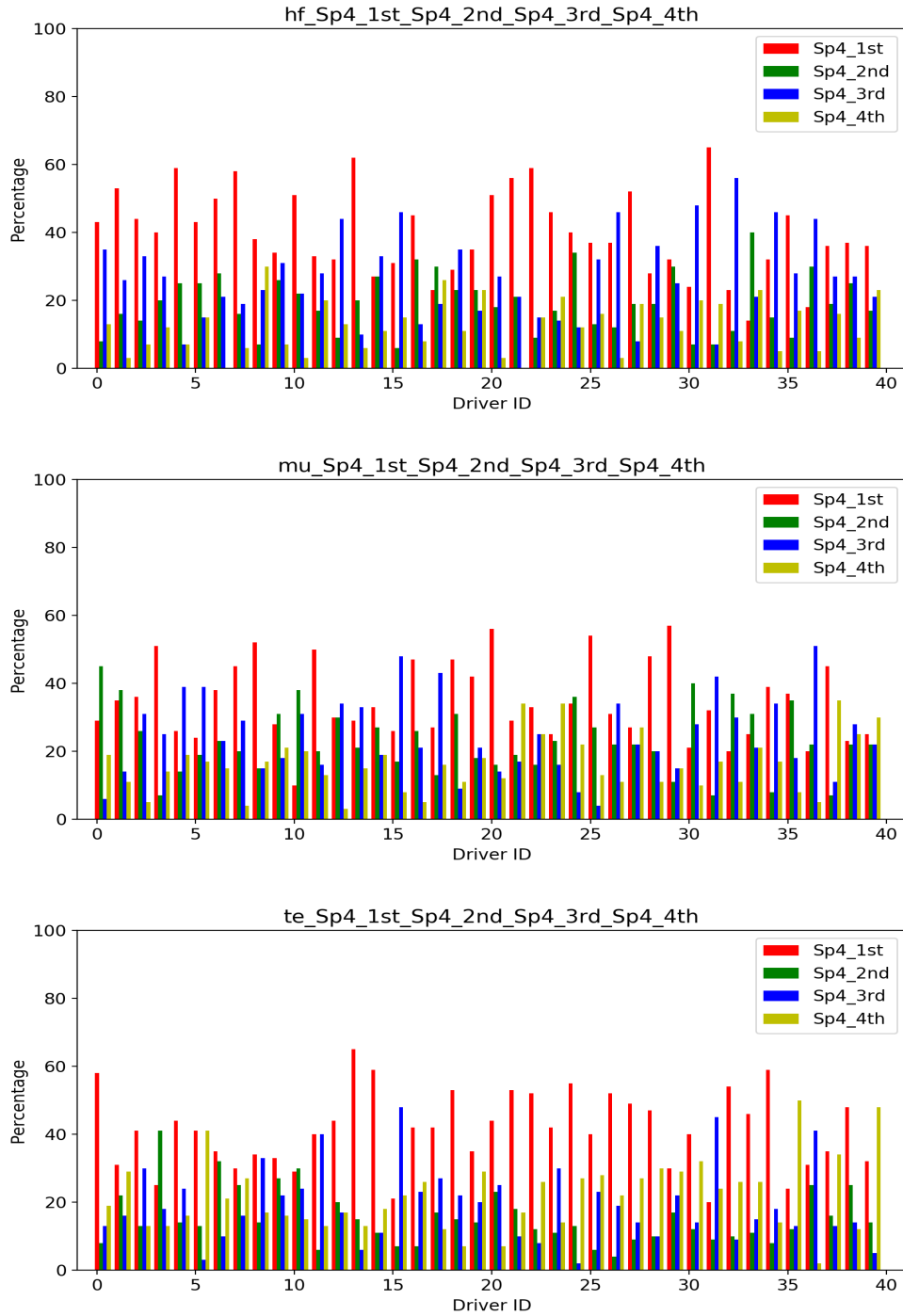


Figure B.9: Examining the activity of drivers under different section of their dataset by separating the classes within the split into three sections label based on percentages
 (a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

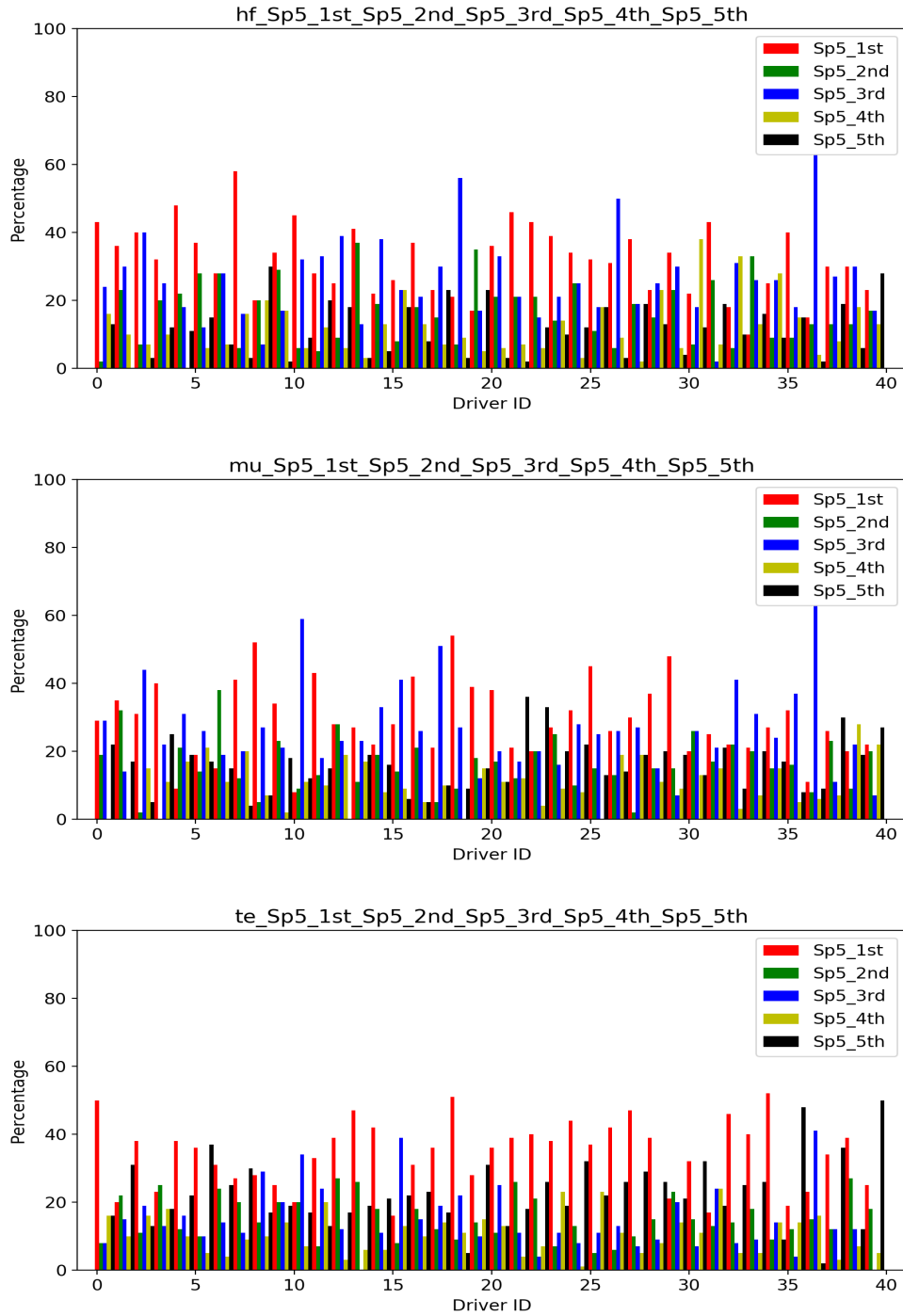


Figure B.10: Examining the activity of drivers under different section of their dataset by separating the classes within the split into three sections label based on percentages
(a) hands-free distraction (top), (b) music distraction (middle) and (c) text distraction (bottom).

C Combined analysis under each class

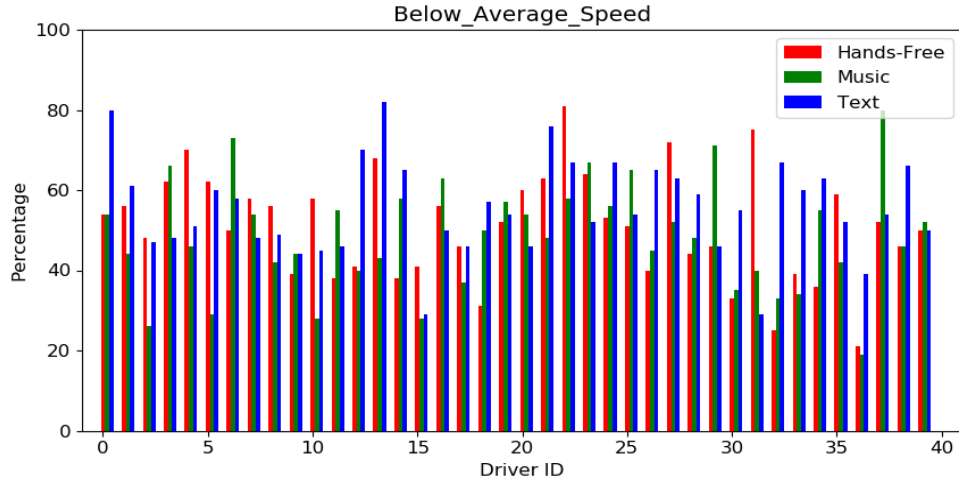


Figure C.1: Examining the activity of drivers under all the distractions under the below average speed class label.

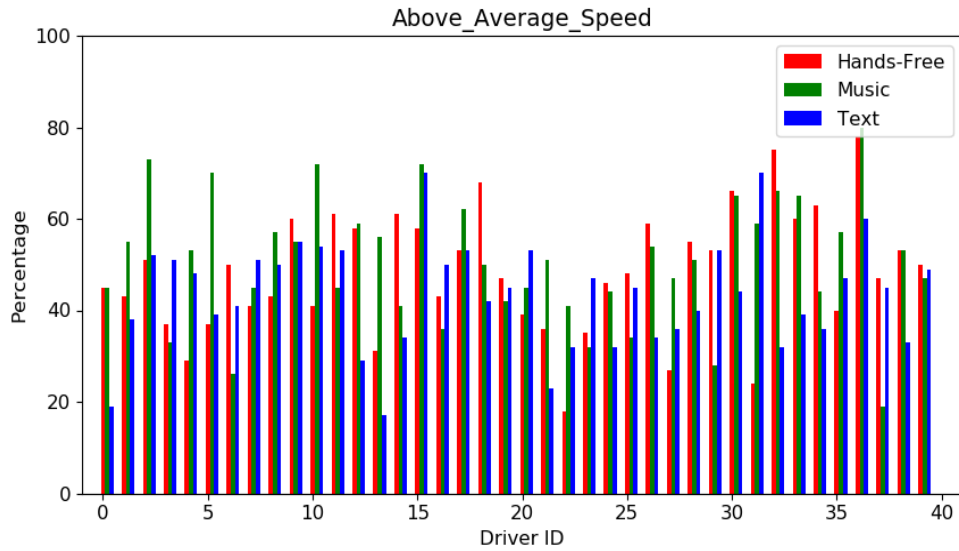


Figure C.2: Examining the activity of drivers under all the distractions under the above average speed class label.

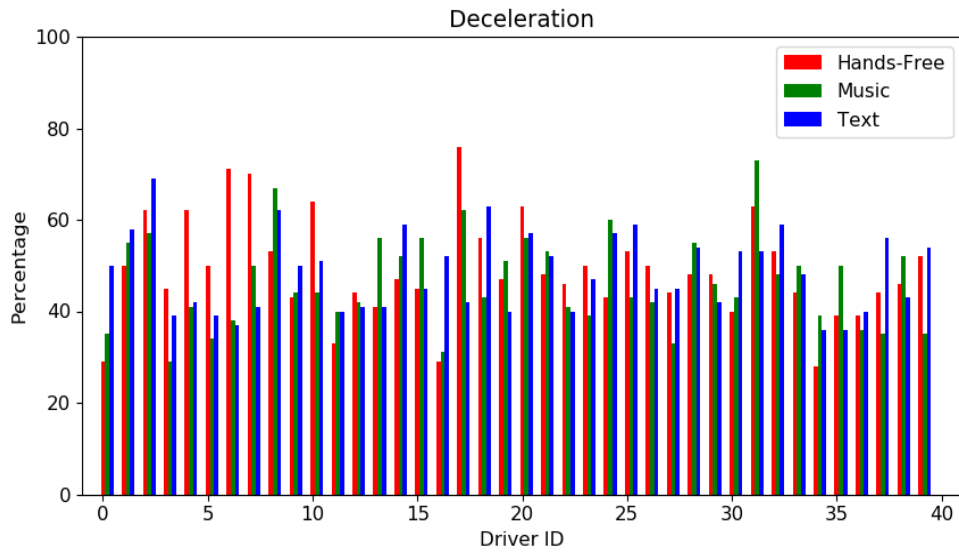


Figure C.3: Examining the activity of drivers under all the distractions under the backward deceleration class label.

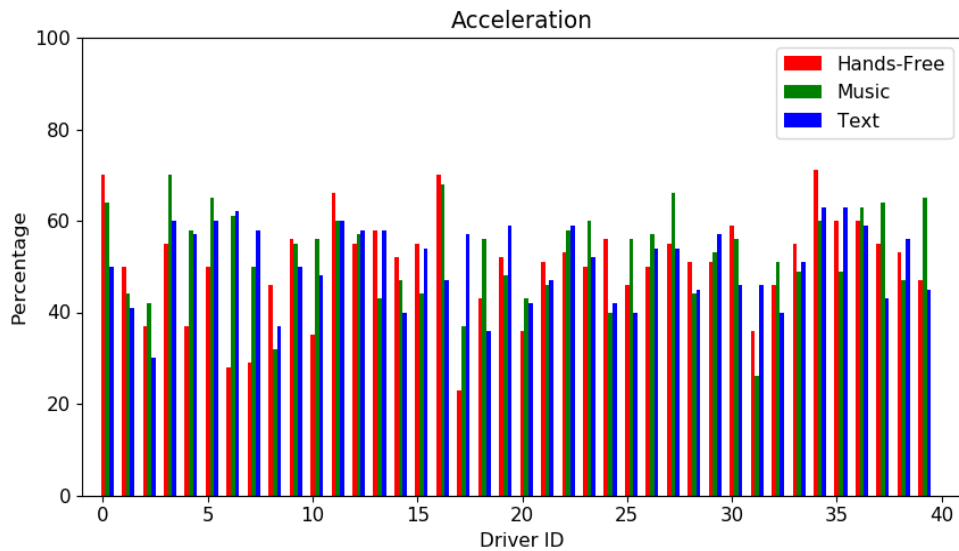


Figure C.4: Examining the activity of drivers under all the distractions under the forward acceleration class label.

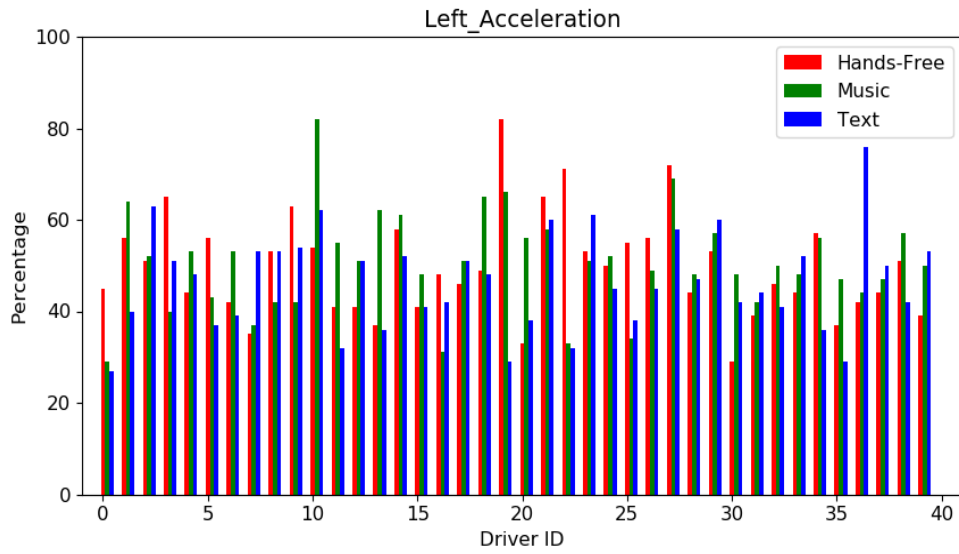


Figure C.5: Examining the activity of drivers under all the distractions under the left acceleration class label.

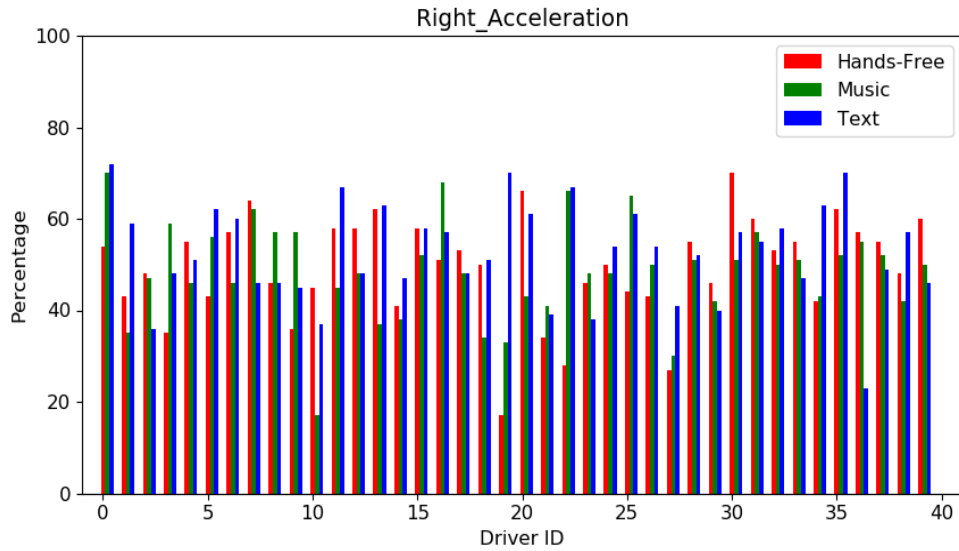


Figure C.6: Examining the activity of drivers under all the distractions under the right acceleration class label.

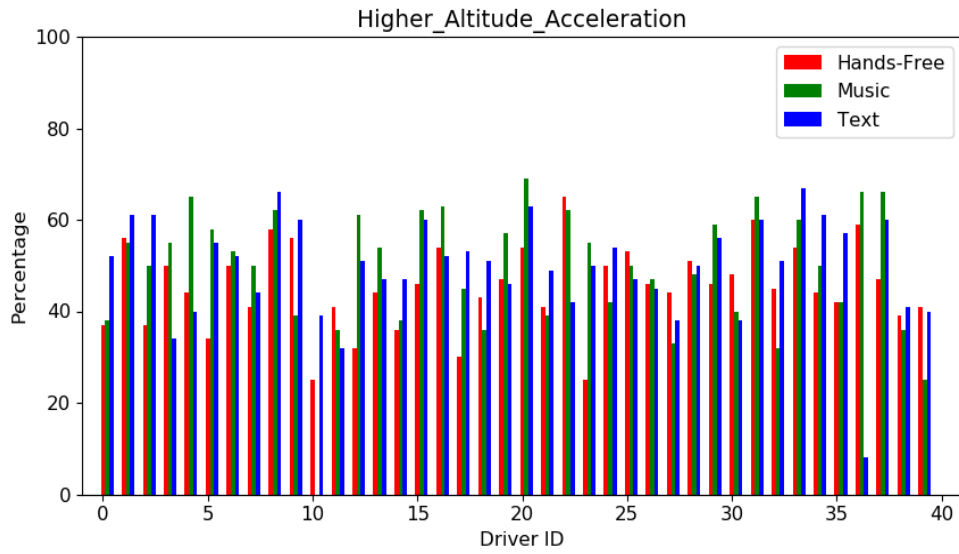


Figure C.8: Examining the activity of drivers under all the distractions under the higher altitude acceleration class label.

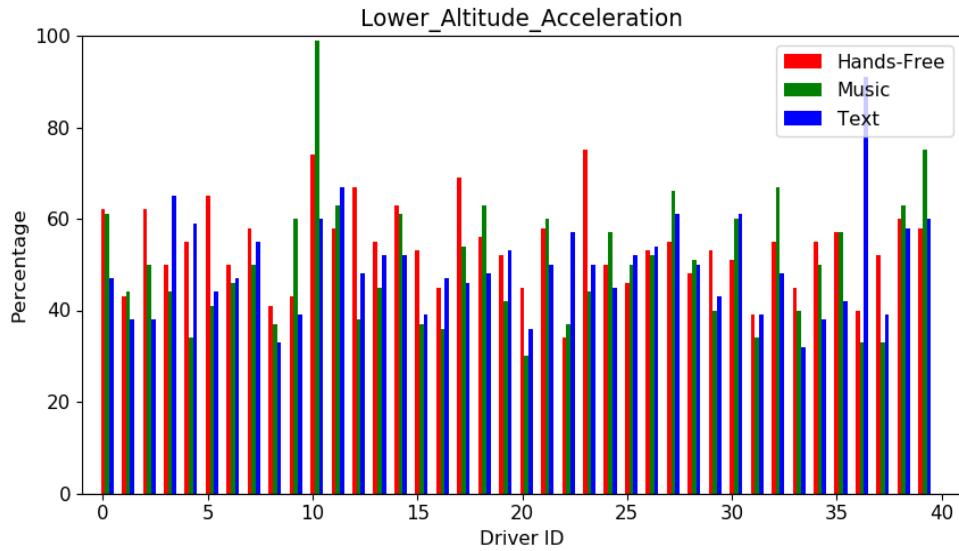


Figure C.7: Examining the activity of drivers under all the distractions under the lower altitude acceleration class label.

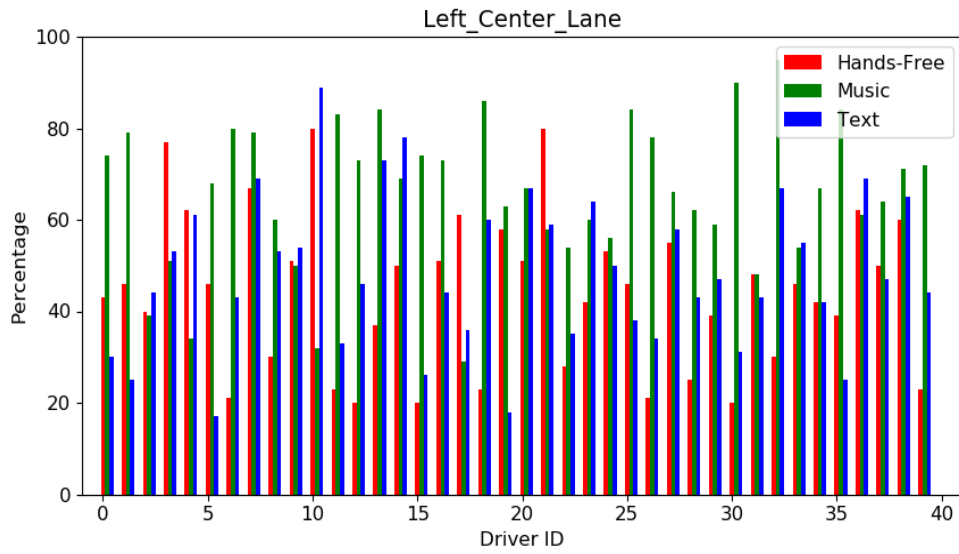


Figure C.9: Examining the activity of drivers under all the distractions under the left center lane class label.

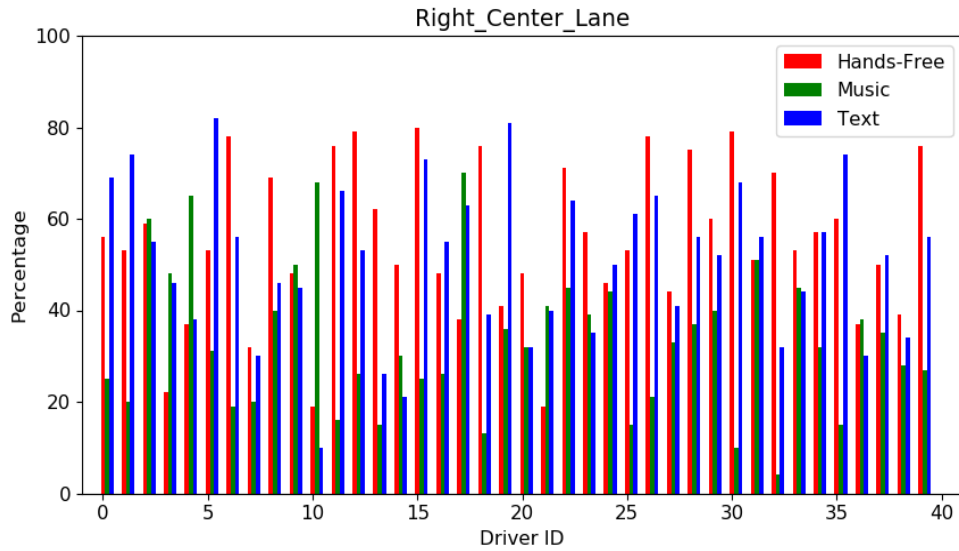


Figure C.10: Examining the activity of drivers under all the distractions under the right center lane class label.

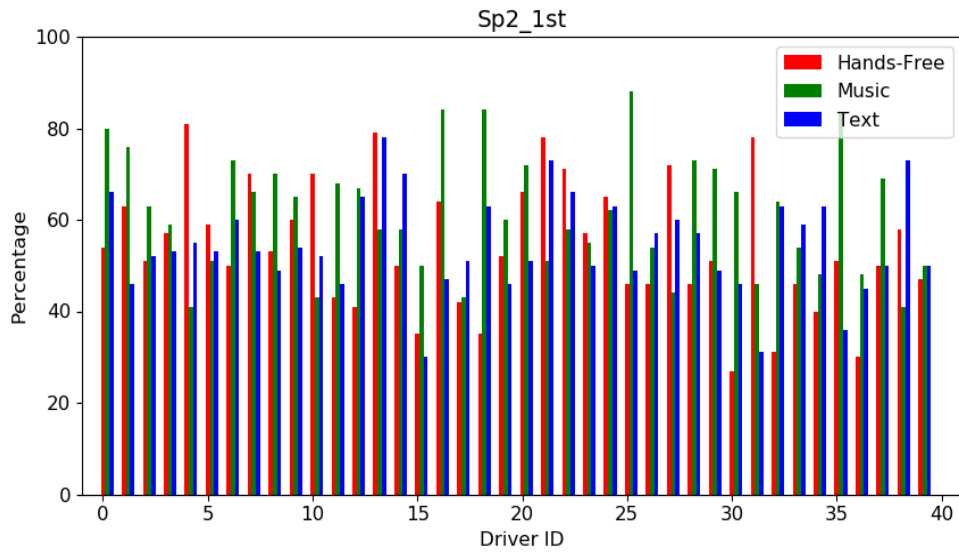


Figure C.11: Examining the activity of drivers under all the distractions under the split two 1st section class label.

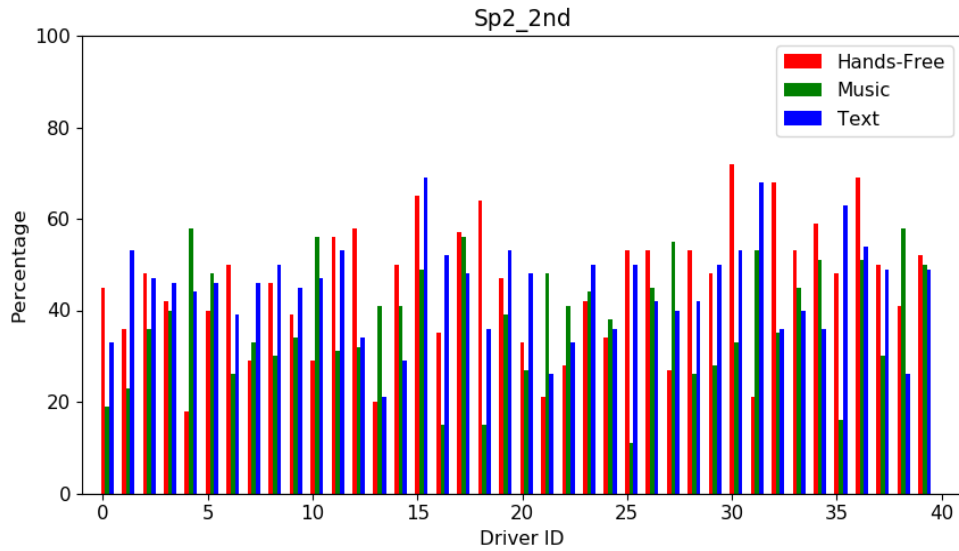


Figure C.12: Examining the activity of drivers under all the distractions under the split two 2nd section class label.

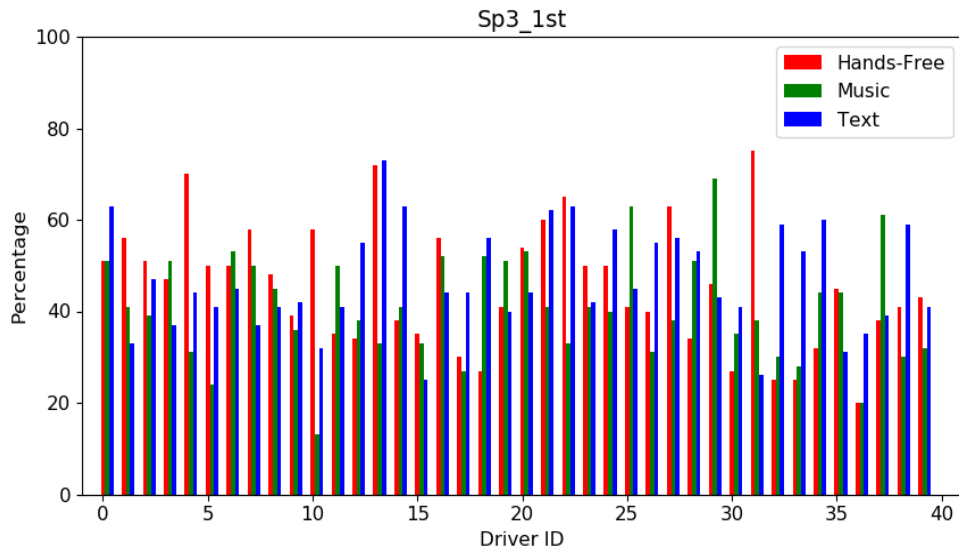


Figure C.13: Examining the activity of drivers under all the distractions under the split three 1st section class label.

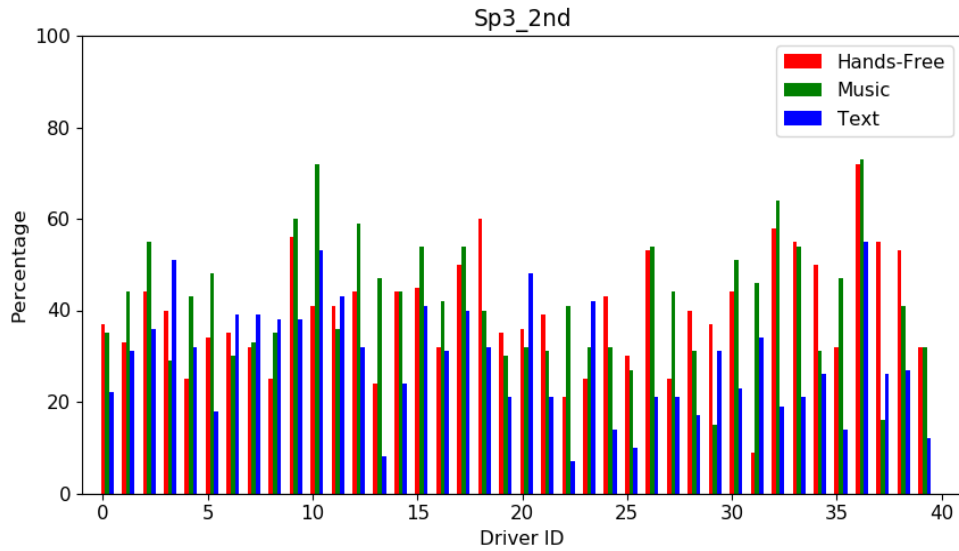


Figure C.14: Examining the activity of drivers under all the distractions under the split three 2nd section class label.

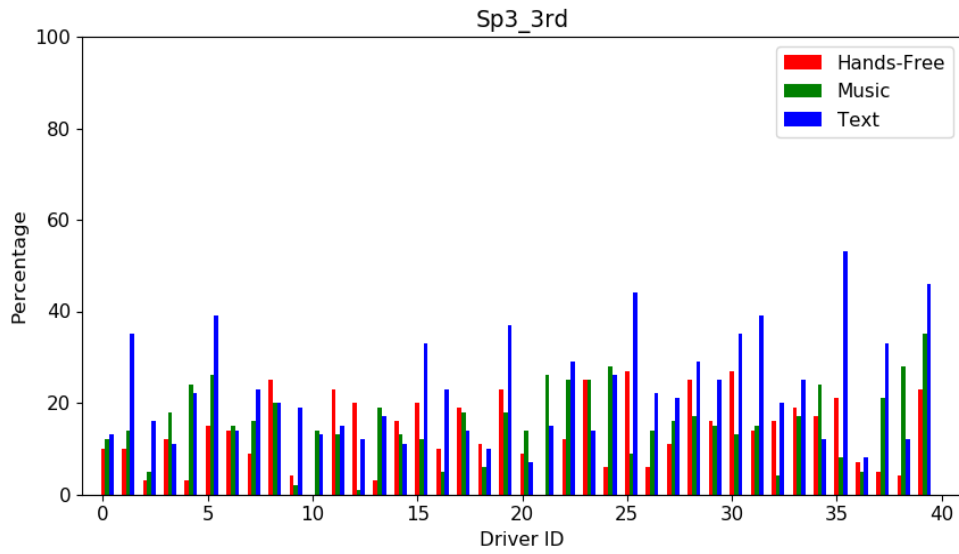


Figure C.15: Examining the activity of drivers under all the distractions under the split three 3rd section class label.

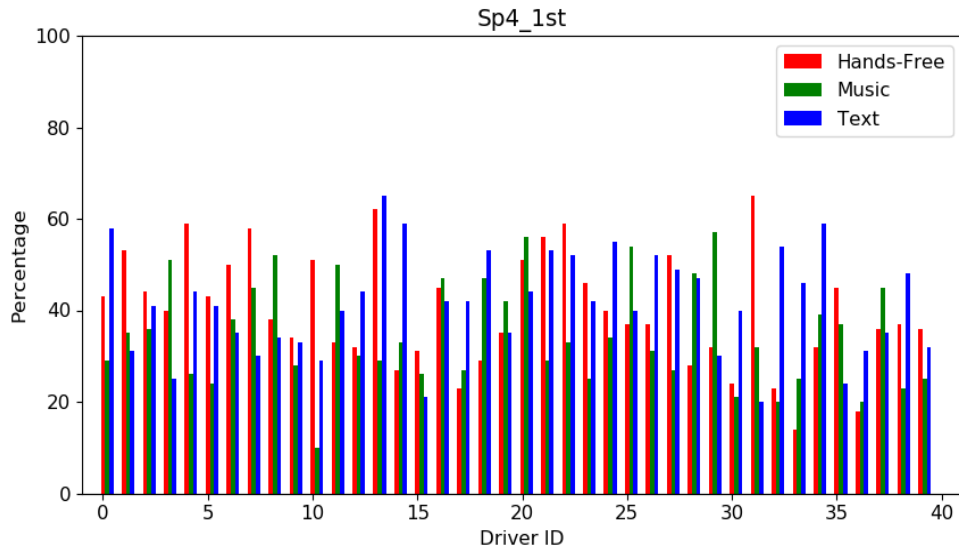


Figure C.16: Examining the activity of drivers under all the distractions under the split four 1st section class label.

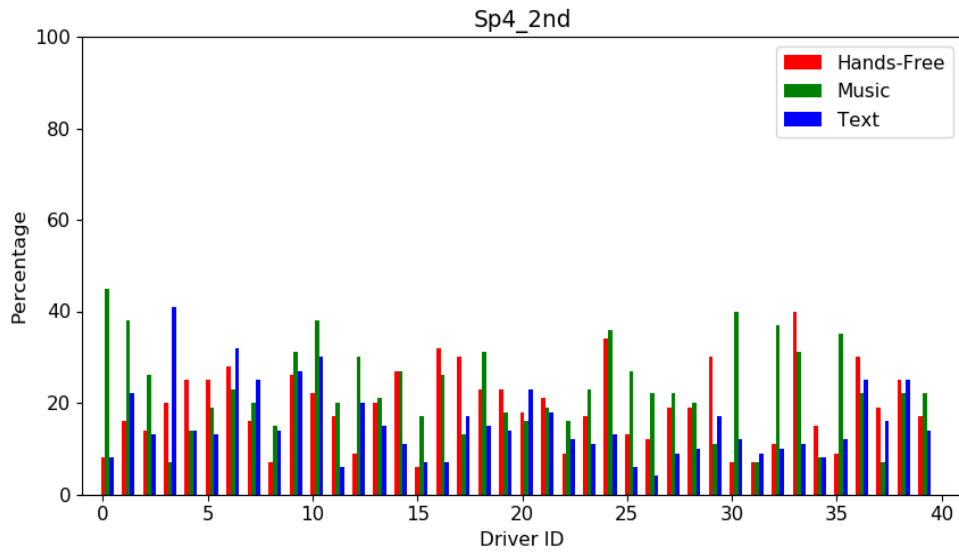


Figure C.17: Examining the activity of drivers under all the distractions under the split four 2nd section class label.

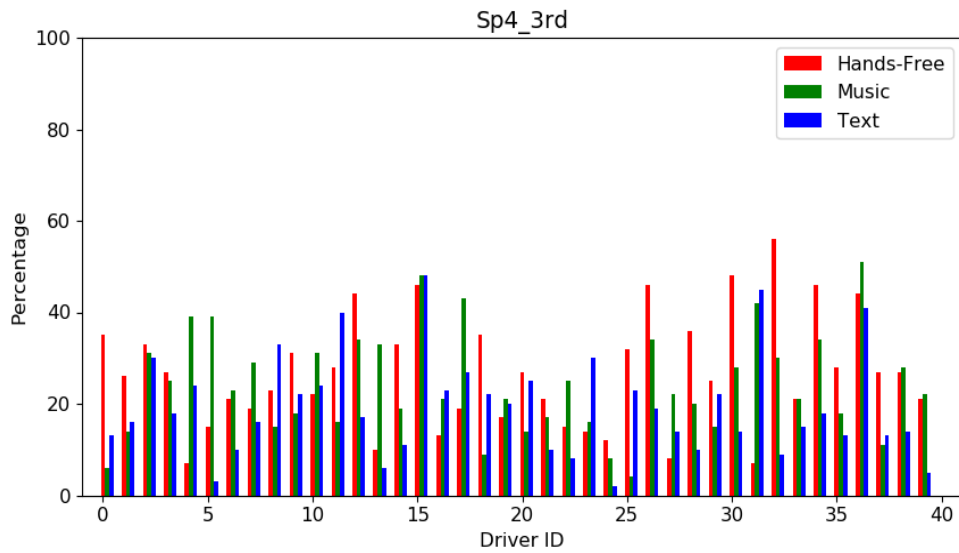


Figure C.18: Examining the activity of drivers under all the distractions under the split four 3rd section class label.

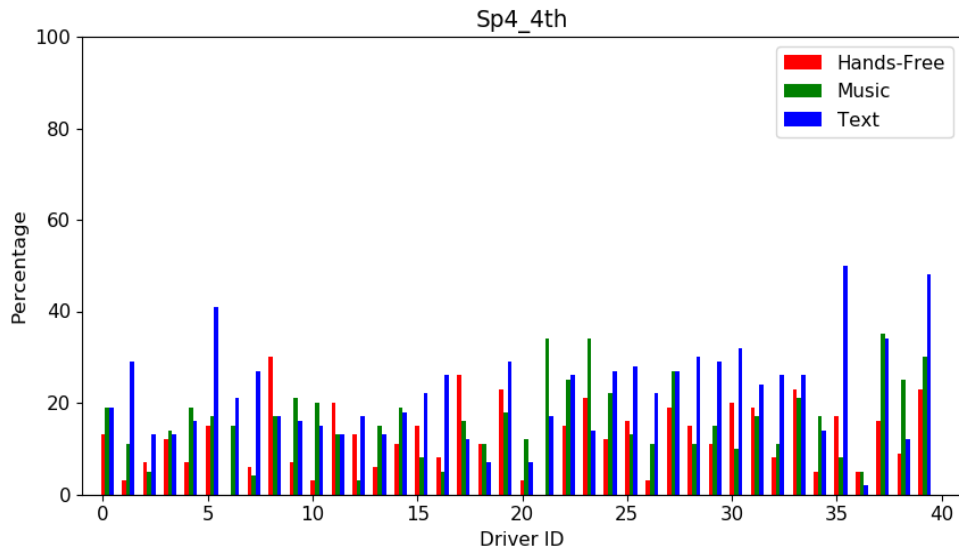


Figure C.19: Examining the activity of drivers under all the distractions under the split four 4th section class label.

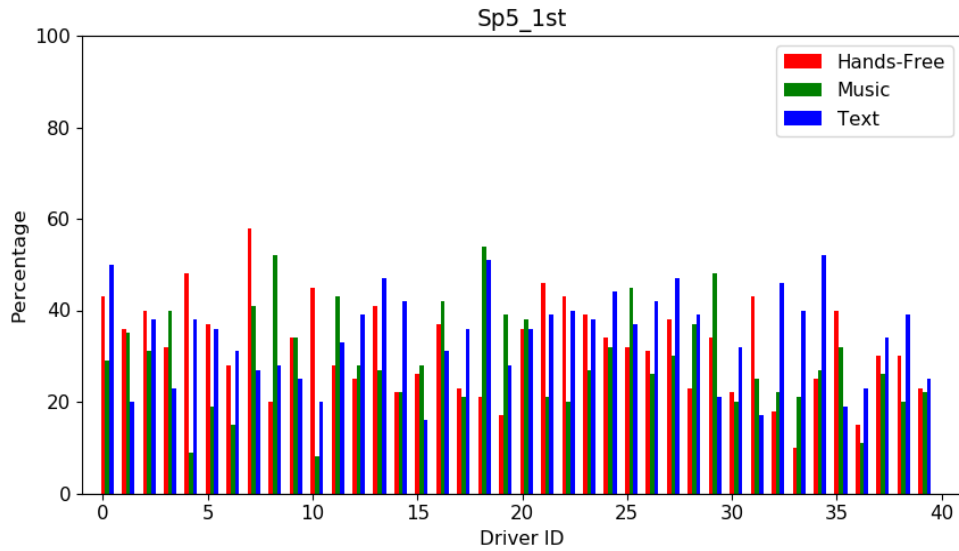


Figure C.20: Examining the activity of drivers under all the distractions under the split five 1st section class label.

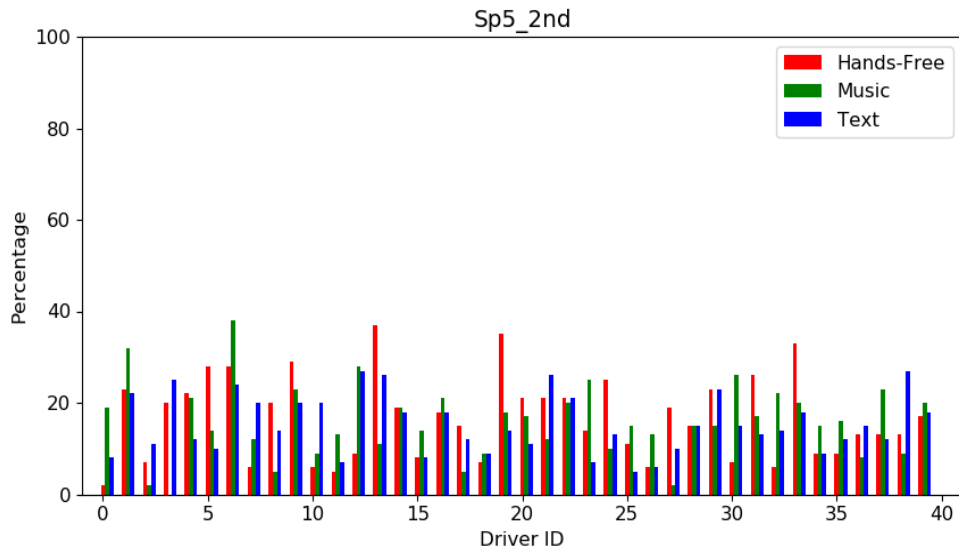


Figure C.21: Examining the activity of drivers under all the distractions under the split five 2nd section class label.

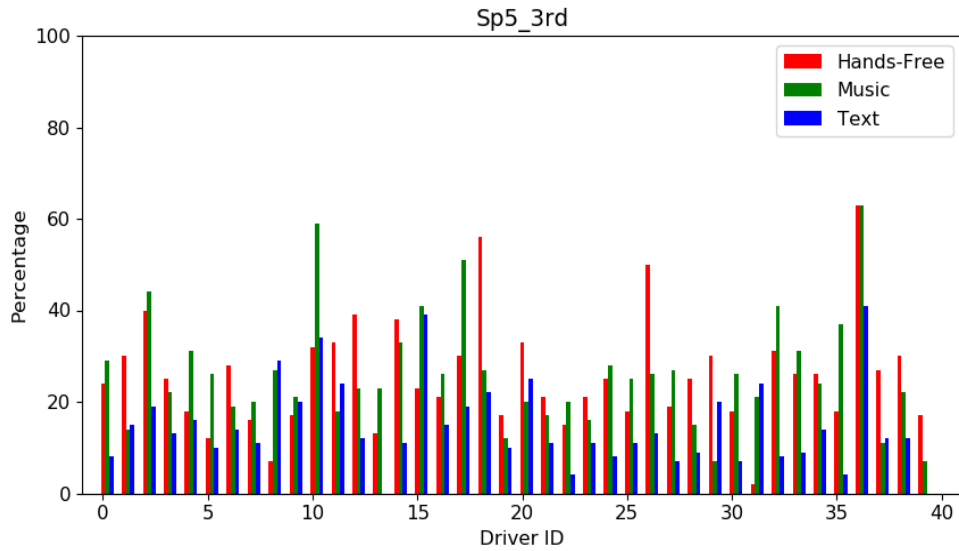


Figure C.22: Examining the activity of drivers under all the distractions under the split five 3rd section class label.

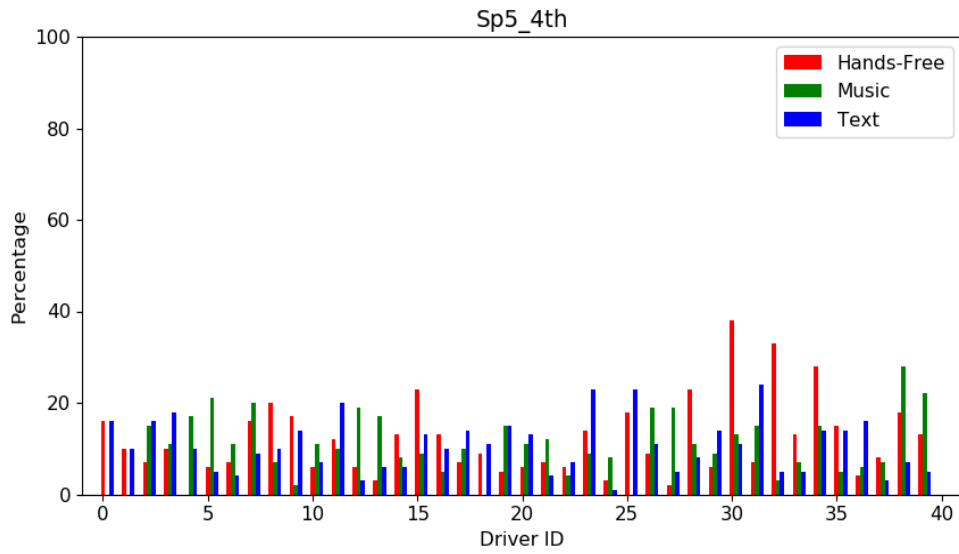


Figure C.23: Examining the activity of drivers under all the distractions under the split five 4th section class label.

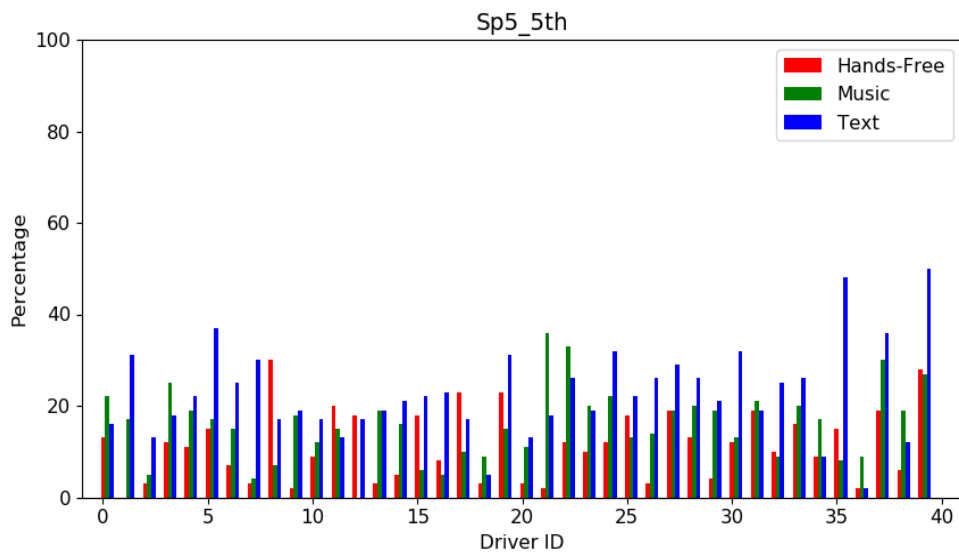


Figure C.24: Examining the activity of drivers under all the distractions under the split five 5th section class label.