

Configuration Manual

MSc Research Project
MSc. in Data Analytics

Richard Burke
Student ID: 15034097

School of Computing
National College of Ireland

Supervisor: Dr. Catherine Mulwa

National College of Ireland
MSc Project Submission Sheet



School of Computing

Student Name: Richard Burke
Student ID: 15034097
Programme: MSc. In Data Analytics **Year:** 2020
Module: Research Project
Lecturer: Catherine Mulwa
Submission Due Date: 17/08/2020
Project Title: The Significance of External Factors on an Individual's Health Determination
Word Count: 1027 **Page Count:** 11

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project.

ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature: Richard Burke
Date: 17/08/2020

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission, to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project, both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

Configuration Manual

Richard Burke
15034097

1 Introduction

The objective of this configuration manual is to offer the reader an understanding of the software and hardware requirements required for the delivery of the research project: 'Extraneous Factors in an Individual's Health Determination'. Furthermore, the reader will be able to implement independently post completion of this document.

2 System Configuration

2.1 Hardware

Model: Dell Latitude 7480, 2018, 256 GB

Processor: 2.8 GHz Intel Core i7-7600U

RAM: 16 GB 2904 MHz

Graphics: Intel HD Graphics 620 Dynamic ram allocation

Operating System: Microsoft Windows 10

2.2 Software

R Studio: Version 1.2.1335

R: R version 3.6.1 (2019-07-05)

3 Project Development

Following a Cross Industry Standard Process for Data Mining (CRISP- DM)

Methodology, the data collection, data transformation, data processing, and data visualisation are outlined below.

3.1 Datasets

Deprivation index data received from Jonathan Pratschke. Ireland shapefile and CSO data downloaded from the Irish Governmental Open Data site.¹

Data sets:

- Playgrounds (2011)
- Streetlights (2010)
- Parks (2011)
- Schools (2011)
- Health Centres (2011)
- Leisure Sites (2009)
- Trees (2011)
- Census 2016 Small Area Population Statistics
- Ireland Census 2016 Boundary File

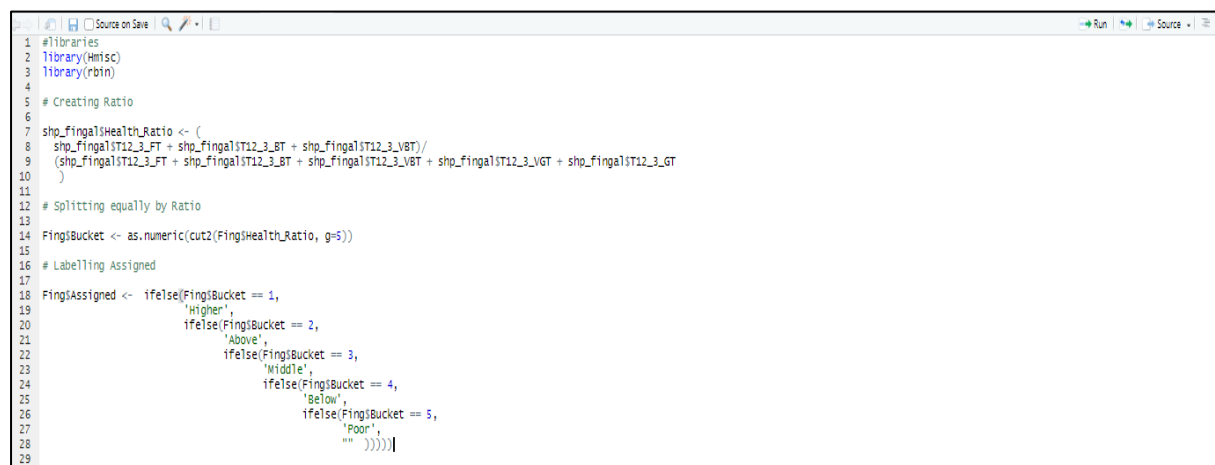
Jonathan Pratschke & Trutz Hasse (RIP):

- SAP Level Deprivation Index data – 2016²

3.2 Data Preparation

Data preparation is completed in R using R Studio. To generate a meaningful response variable, the self-assessed health determination from the census was ranked based on the proportion of total residents considering their health to be fair, bad, or very bad.

These were then split into five equal-sized groups based on their ratio and assigned categories.



```
1 #libraries
2 library(Hmisc)
3 library(rbin)
4
5 # Creating Ratio
6
7 shp_fingalHealthRatio <- (
8   shp_fingalST12_3_FT + shp_fingalST12_3_BT + shp_fingalST12_3_VBT)/
9   (shp_fingalST12_3_FT + shp_fingalST12_3_BT + shp_fingalST12_3_VBT + shp_fingalST12_3_VGT + shp_fingalST12_3_GT
10 )
11
12 # Splitting equally by Ratio
13
14 FingalBucket <- as.numeric(cut2(FingalHealthRatio, g=5))
15
16 # Labelling Assigned
17
18 FingalAssigned <- ifelse(FingalBucket == 1,
19   'Higher',
20   ifelse(FingalBucket == 2,
21     'Above',
22     ifelse(FingalBucket == 3,
23       'Middle',
24       ifelse(FingalBucket == 4,
25         'Below',
26         ifelse(FingalBucket == 5,
27           'Poor',
28           "" )))))
29 ]
```

Figure 1: Assigned Categories in R Studio

¹ <https://data.gov.ie/organization/fingal-county-council>

² <http://trutzhase.eu/deprivation-index/the-2016-pobal-hp-deprivation-index-for-small-areas/>

Every facility or amenity needed to be mapped to a small area to generate aggregate statistics. These are completed in R Studio using the spatial packages (sp and rgdal). This was completed for each amenity separately before being amalgamated. Example of transformation for play areas included in Figure 2.

```

1 #libraries
2 library(sp)
3 library(rgdal)
4
5 # Play Area
6 play <- read.csv("C:/Users/burkeri/Desktop/Project_2/fingal_play.csv")
7
8 # Rename to lat/long
9 colnames(play)[23] <- 'Latitude'
10 colnames(play)[24] <- 'Longitude'
11
12 # Remove non-geocoded
13 play <- subset(play, !is.na(play$Latitude))
14
15 # add coordinates to file creating spatial points data frame
16 coordinates(play) <- c("Longitude", "Latitude")
17
18 #read in ireland shapefile
19 ir1_shp <- readOGR(".", "Small_Areas_Generalised_20m_OSI_National_Boundaries")
20 ir1_shp_wgs84 <- spTransform(ir1_shp, CRS("+proj=longlat +datum=WGS84"))
21
22 #apply same projection
23 proj4string(play) <- proj4string(ir1_shp_wgs84)
24
25 # overlay play areas with small areas and merge
26 play_geo <- over(play, ir1_shp_wgs84)
27

```

Figure 2: Overlay play areas and small areas in R Studio

The streetlight data set was mapped using easting and northing rather the longitude and latitude. This required a spatial transformation to generate their respective longitude and latitude references.

```

1 #Read in streetlight files
2 streetlights <- read.csv("C:/Users/burkeri/Desktop/Project_2/fingal_streetlights.csv")
3 streetlights <- streetlights[complete.cases(streetlights[, 7:8]),]
4
5 # Extract easting and northing
6 holes <- subset(streetlights, select = c("Easting", "Northing"))
7 holes$holes_ID <- 1:nrow(holes)
8 coords <- cbind(EASTING = round(as.numeric(as.character(holes$Easting)),0),NORTHING = round(as.numeric(as.character(holes$Northing)),0))
9
10 #Project to standard epsg
11 holes_SP <- SpatialPointsDataFrame(coords, data = data.frame(holes$holes_ID), proj4string = CRS("+init=epsg:29902"))
12 holes_SP_LL <- spTransform(holes_SP, CRS("+proj=longlat"))
13 points(holes_SP_LL, col=2, pch=18)
14 round(as.numeric(holes$Northing),0)
15
16 # Create points, rename and merge back to streetlight data
17 streetlights_points <- cbind(holes_SP_LL@coords, holes_SP_LL@data)
18 colnames(streetlights)[1] <- 'id'
19 colnames(streetlights_points)[3] <- 'id'
20 streetlights_new <- merge(streetlights_points, streetlights, by = 'id', all.x = TRUE)
21 streetlight <- subset(streetlights_new, !is.na(streetlights_points$NORTHING))
22 coordinates(streetlight) <- c("EASTING", "NORTHING")
23
24 # Match projection to Ireland shapefile
25 ir1_shp_wgs84 <- spTransform(ir1_shp, CRS("+proj=longlat +datum=WGS84"))
26 proj4string(streetlight) <- proj4string(ir1_shp_wgs84)
27
28 #Layer over small areas
29 streetlight_geo <- over(streetlight, ir1_shp_wgs84)

```

Figure 3: Reprojection of streetlights to latitude longitude

The points for all amenities were aggregated by each small area in Fingal.

```

1 library(dplyr)
2
3 play_count <- play_geo %>% count(SMALL_AREA)
4 colnames(play_count)[2] <- 'Play_Area'
5
6 streetlight_count <- streetlight_geo %>% count(SMALL_AREA)
7 colnames(streetlight_count)[2] <- 'Streetlight_Area'
8
9 parks_count <- parks_geo %>% count(SMALL_AREA)
10 colnames(parks_count)[2] <- 'Parks_Area'
11
12 school_count <- schools_geo %>% count(SMALL_AREA)
13 colnames(school_count)[2] <- 'School_Area'
14
15 health_count <- health_geo %>% count(SMALL_AREA)
16 colnames(health_count)[2] <- 'Health_Area'
17
18 leisure_count <- leisure_geo %>% count(SMALL_AREA)
19 colnames(leisure_count)[2] <- 'Leisure_Area'
20
21 tree_count <- tree_geo %>% count(SMALL_AREA)
22 colnames(tree_count)[2] <- 'Tree_Area'
23
24
25 Fingal_Areas <- as.data.frame(Fingal$SMALL_AREA)
26 colnames(Fingal_Areas)[1] <- 'SMALL_AREA'
27

```

Figure 4: Aggregation to small area

Due to the sparse nature of the data sets (apart from trees and streetlights), secondary factors were created to represent the nearest distance from each polygon (small area) to the amenity.

```

1 ##### Park_Distance #####
2
3 # Define global reference ellipsoid and a gravity field model
4 utmStr <- "+proj=utm +zone=%d +datum=NAD83 +units=m +no_defs +ellps=GRS80"
5 crs <- CRS(sprintf(utmStr, 29))
6
7 # Apply CRS
8 proj4string(parks) <- CRS(proj4string(Fingal_Shp_File))
9 parks <- spTransform(parks, crs)
10 Fingal_Shp_File <- spTransform(Fingal_Shp_File, crs)
11
12 # Determine distance
13 a <- gDistance(Fingal_Shp_File, parks, byid=TRUE)
14 park_name <- as.data.frame(row.names(a))
15 polygon_name <- as.data.frame(colnames(a))
16 # Vecotr of column names and min distance
17 vec = colnames(a)
18 min_dist_park <- as.data.frame(apply(a[,vec], 2, min))
19 colnames(min_dist_park)[1] <- 'min_dist_park'
20 min_dist_park_num <- as.data.frame(apply(a[,vec], 2, which.min))
21 min_dist_park$row_id <- row.names(min_dist_park)
22 Fingal_Shp_Frm <- Fingal_Shp_File@data
23 Fingal_Shp_Frm$row_id <- row.names(Fingal_Shp_Frm)
24
25 #merge back to main shapefile
26 dist_to_park <- merge(min_dist_park, Fingal_Shp_Frm[, c('SMALL_AREA', 'row_id')], by = 'row_id')
27

```

Figure 5: Derive minimum distance between points and polygons

Finally, combine all the data derived data sets ready for processing and visualisation.

```

1 #Merge data sets
2 HP_Index <- read.csv("c:/Users/burkeri/Desktop/Project_2/Copy of HP Index 2006-2016 HP Index Scores
3
4 HP_Index$Small_Area_New <- ifelse(HP_Index$SmallArea2017 == '',
5                                   as.character(HP_Index$SmallArea2011),
6                                   as.character(HP_Index$SmallArea2017))
7
8 Fingal_Areas <- merge(Fingal_Areas, play_count, by = 'SMALL_AREA', all.x = TRUE)
9 Fingal_Areas <- merge(Fingal_Areas, streetlight_count, by = 'SMALL_AREA', all.x = TRUE)
10 Fingal_Areas <- merge(Fingal_Areas, parks_count, by = 'SMALL_AREA', all.x = TRUE)
11 Fingal_Areas <- merge(Fingal_Areas, school_count, by = 'SMALL_AREA', all.x = TRUE)
12 Fingal_Areas <- merge(Fingal_Areas, health_count, by = 'SMALL_AREA', all.x = TRUE)
13 Fingal_Areas <- merge(Fingal_Areas, leisure_count, by = 'SMALL_AREA', all.x = TRUE)
14 Fingal_Areas <- merge(Fingal_Areas, tree_count, by = 'SMALL_AREA', all.x = TRUE)
15 Fingal_Areas <- merge(Fingal_Areas, dist_to_health, by = 'SMALL_AREA')
16 Fingal_Areas$row_id <- NULL
17 Fingal_Areas <- merge(Fingal_Areas, dist_to_schools, by = 'SMALL_AREA')
18 Fingal_Areas$row_id <- NULL
19 Fingal_Areas <- merge(Fingal_Areas, dist_to_playg, by = 'SMALL_AREA')
20 Fingal_Areas$row_id <- NULL
21 Fingal_Areas <- merge(Fingal_Areas, dist_to_leisure, by = 'SMALL_AREA')
22 Fingal_Areas$row_id <- NULL
23 Fingal_Areas <- merge(Fingal_Areas, dist_to_park, by = 'SMALL_AREA')
24 Fingal_Areas$row_id <- NULL
25 Fingal_Areas <- merge(Fingal_Areas, Fing[,c(1, 32)], by = 'SMALL_AREA', all.x = TRUE)
26 Fingal_Areas[is.na(Fingal_Areas)] <- 0
27 Fingal_Areas_2 <- merge(Fingal_Areas, HP_Factors[,c(1:3)], by = "SMALL_AREA", all.x = TRUE)
28

```

Figure 6: Merge the derived points and aggregations

4 Visualisation and Inspection

The aggregated data sets were inspected to produce visual and tabular depictions showing the dispersion of areas with below-average health determination across Fingal. The polygons must be fortified to produce the shapes and merged back to the derived data.

The plots are built using ggplot2 and colour coding polygons based on their health categorisation.

```

1 # in order to plot polygons, first fortify the data
2 Fingal_Shp_File@data$id <- rownames(Fingal_Shp_File@data)
3 Fingal_Shp_File.df <- fortify(Fingal_Shp_File, region='id')
4
5
6 #Merge to original data
7 Fingal_Shp_File.df <- merge(Fingal_Shp_File.df, as.data.frame(Fingal_Shp_File[,c('id', 'Assigned')]), by = "id")
8
9 #Make reposnse a factor
10 Fingal_Shp_File.df$Assigned <- factor(Fingal_Shp_File.df$Assigned, levels = c("Higher", "Above", "Middle", "Below", "Poor"))
11
12 # build plot
13 Fingal_Shp_File$map1 <- ggplot(data = Fingal_Shp_File.df, aes(long, lat, group=id))
14 Fingal_Shp_File$map1 <- Fingal_Shp_File$map1 + geom_polygon(aes(fill = Assigned))
15 ##Fingal_Shp_File$map1 <- Fingal_Shp_File$map1 + geom_path(colour = 'gray', linestyle = 2)
16 Fingal_Shp_File$map1 <- Fingal_Shp_File$map1 + scale_fill_brewer(palette='YlGnBu')
17 Fingal_Shp_File$map1 <- Fingal_Shp_File$map1 + labs(title="Fingal by Health Categorisation")
18

```

Figure 7: Depict visually using ggplot2

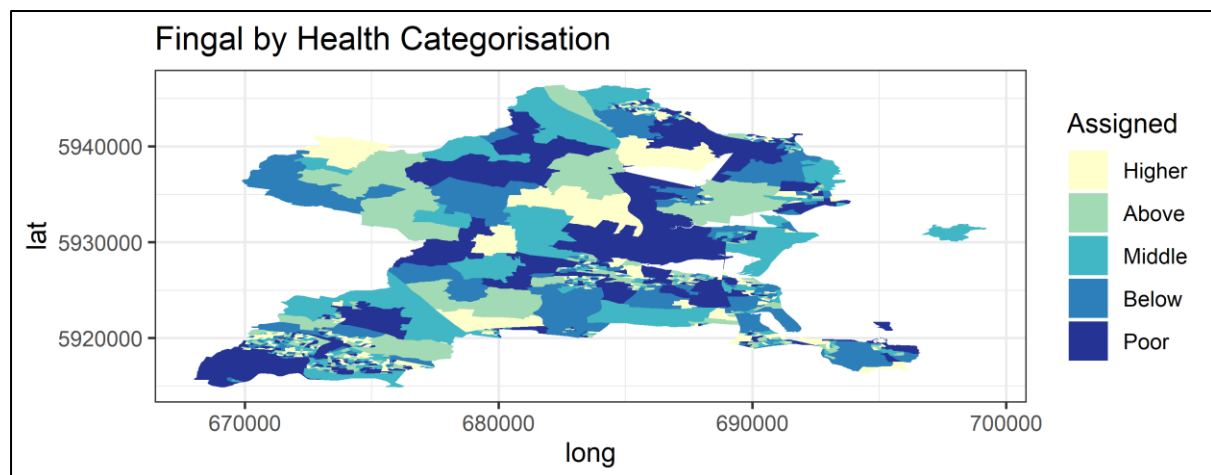


Figure 8: Dispersion of health categorisation across Fingal

The data was further aggregated by Assigned category and Electoral district to view of the underlying data points for each entity.

```

159
160 summed_tables <- aggregate(~Assigned, Fingal_Areas, sum)
161 summed_ED_tables <- aggregate(~EDNAME, Fingal_Areas, sum)
162 mean_tables <- aggregate(~Assigned, Fingal_Areas, mean)
163 mean_ED_tables <- aggregate(~EDNAME, Fingal_Areas, mean)
164
165

```

Figure 9: Aggregate to ED and Category

Table 1: Lowest 5 by categorisation

<i>Assigned</i>	<i>EDNAME</i>	<i>n</i>	<i>Total</i>	<i>Percentage</i>
Poor	Blanchardstown-Corduff	10	12	83.33
Poor	Balbriggan Urban	15	28	53.57
Poor	Blanchardstown-Roselawn	3	6	50.00
Poor	Blanchardstown-Tyrrelstown	3	6	50.00
Poor	Kilsallaghan	3	6	50.00

Visually comparing each ED to its overall mean in ggplot2 by transposing the wide data sets to long format and generating group means to add as points.

```

1 library(ggplot2)
2 library(data.table)
3
4 #Change mean ED tables from wide to long
5 df_melted2 <-
6   melt(mean_ED_tables[,c(1, 21:31)]
7         , id.vars="EDNAME", variable.name="Resource", value.name="Count")
8
9 # View lowest areas by mean assigned against overall mean
10 df_melted2 %>%
11   group_by(Resource) %>%
12   summarise(Count = mean(Count))
13 df_melted2 <- df_melted2[order(-df_melted2$Count),]
14
15 # Plot with ggplot
16 p <- ggplot(df_melted2 %>%
17   filter(EDNAME %in% c('Blanchardstown-Corduff', 'Blanchardstown-Roselawn',
18     'Swords Village', 'Balbriggan Urban', 'Blanchardstown-Tyrrelstown', 'Kilsallaghan')) , aes(Resource, Count, fill=Resource)) +
19   geom_bar(stat="identity", position = position_dodge()) +
20   facet_wrap(~ EDNAME, scales = "free") +
21   coord_flip() +
22   theme(panel.background = element_blank()) +
23   geom_point(data = df_melted2 %>%
24     group_by(Resource) %>%
25     summarise(Count = mean(Count))) +
26   labs(
27     title = "Lowest 6 Areas by Health Determination",
28     x = "Count",
29     y = "Resource"
30   )
31
32

```

Figure 10: Plotting ED and overall means

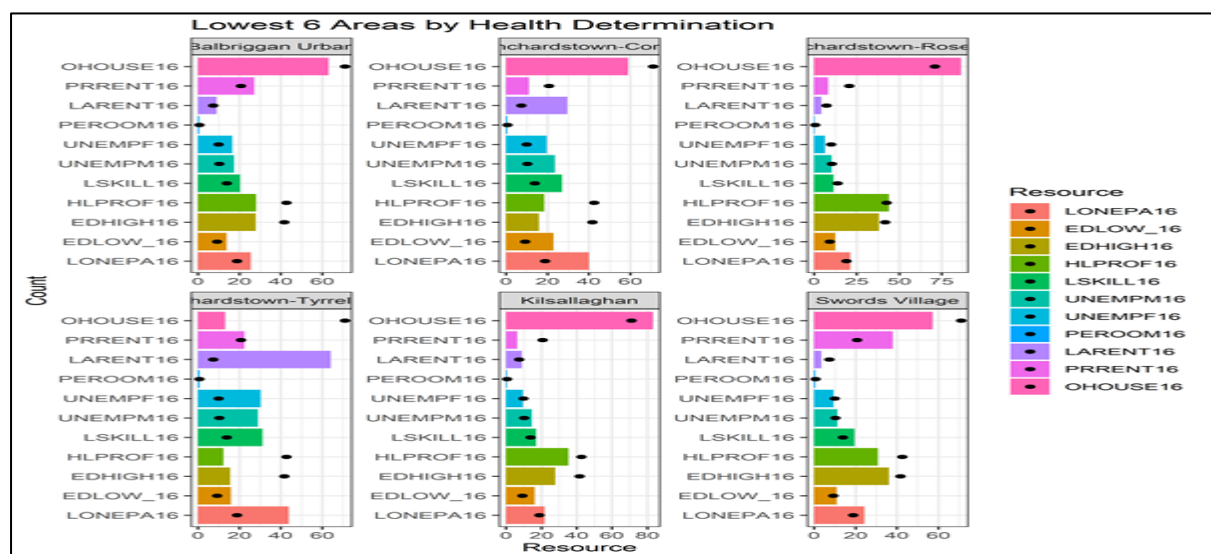


Figure 11: Lowest 6 by self-determined health

5 Modelling

The package Caret in R was used to implement all the models (excluding Auto ML) and to develop optimised parameters for XGBoost and Random Forest. Classification accuracy, precision, and recall were used as measures of a model's performance on a random 25% holdout set.

5.1 Auto ML

Specifying the max number of model's limits run time and ensures model selection can be completed.

A screenshot of an R script editor window showing the implementation of AutoML using the h2o package. The script includes the following code:

```
1 # Load library
2 library(h2o)
3 library(dplyr)
4 library(caret)
5
6
7 cols.num <- c("POPCHG16", "AGEDEP16", "LONEPA16", "EDLOW_16", "EDHIGH16", "HLPROF16",
8             "LSKILL16", "UNEMP16", "UNEMPF16", "PEROOM16", "LARENT16", "PRRENT16", "OHOUSE16" )
9 Fingal_Areas_2[cols.num] <- sapply(Fingal_Areas_2[cols.num], as.numeric)
10
11
12 Fingal_Areas_2$Assigned <- recode(Fingal_Areas_2$Assigned, 'Higher'=5, 'Above'=4, 'Middle'=3, 'Below'=2, 'Poor' = 1)
13 Fingal_Areas_2$Assigned <- as.factor(Fingal_Areas_2$Assigned)
14
15 TrainxG2 <- createDataPartition(Fingal_Areas_2$Assigned, p=0.75, list=FALSE)
16 trainingxG2 <- Fingal_Areas_2[ TrainxG2, ]
17 testingxG2 <- Fingal_Areas_2[ -TrainxG2, ]
18
19
20 # start h2o cluster
21 invisible(h2o.init())
22
23 # convert data as h2o type
24 train_h = as.h2o(trainingxG2)
25 test_h = as.h2o(testingxG2)
26
27 # set label type
28 y = 'Assigned'
29 pred = setdiff(names(trainingxG2), y)
30
31 #convert variables to factors
32 trainingxG2[,y] = as.factor(trainingxG2[,y])
33 testingxG2[,y] = as.factor(testingxG2[,y])
34
35 # Run AutoML for 20 base models
36 aml = h2o.automl(x = pred, y = y,
37               training_frame = train_h,
38               max_models = 20,
39               seed = 1,
40               max_runtime_secs = 3600
41 )
42
43
```

Figure 12: AutoML implementation

The output includes a set of models as a part of a leader board and allows the user to select which model to implement or allow default selection. Predictions of the testing data set can be made within the code, and a confusion matrix demonstrated using caret's in-built function.

AutoML using the package h2o requires pre-processing of the data to execute. This involves recoding the textual response variable to numeric. For this problem, all other variables were set as numeric with a 75%/25% split between training and testing data. The h2o cluster is initialised, and the data sets uploaded.

```

2
3
4 # AutoML Leaderboard
5 lb = aml@leaderboard
6
7 aml@leader
8
9 # prediction result on test data
0 prediction = h2o.predict(aml@leader, test_h[,-5]) %>%
1   as.data.frame()
2
3 prediction$predict
4 #prediction$predict <- as.factor(round(prediction$predict,0))
5 # create a confusion matrix
6 confusion_matrix_h2o <- caret::confusionMatrix(testingXG2$Assigned, prediction$predict)
7
8 # close h2o connection
9 h2o.shutdown(prompt = F)
0

```

Figure 13: AutoML output

5.2 Naïve Bayes

Naive Bayes was implemented using caret and its in-built functions. The default grid search and parameter optimisation were selected by caret. As with AutoML, the training set was split into a test and train set with a 25% and 75% split respectively.

```

207
208 #### Naïve Bayes Basic
209
210 start.time <- Sys.time()
211 Assigned_nb <- train(Assigned ~., data = trainingXG2, method = "nb", metric = "Accuracy", importance = TRUE, preProc = c("center", "scale"))
212 system.time()
213
214 library(caret)
215
216 nb2_VapIMP <- varImp(Assigned_nb)
217
218 test_predict_nb2 <- predict(Assigned_nb, testingXG2)
219
220 summary(test_predict_nb2)
221
222 confusion_matrix_nb2 <- confusionMatrix(test_predict_nb2, as.factor(testingXG2$Assigned))
223 confusion_matrix_nb2
224
225 end.time <- Sys.time()
226
227 time.taken_nb2 <- end.time - start.time
228
229

```

Figure 14: Naïve Bayes implementation

5.3 Random Forest

Additional parameters were specified for the Random Forest model to expand the tuning grid and parameter selection. This was implemented using caret on the same test and train data sets.

```

230 ##### Random Forest Tune
231
232 start.time <- Sys.time()
233
234 # Specifying cross validation - 10 folds.
235 trctrl <- trainControl(method = "cv", number = 10)
236
237 #Specify randomly selected predictors
238 mtry <- sqrt(ncol(trainingXG2))
239
240 tuneGrid <- expand.grid(mtry=mtry)
241
242 # fit model
243
244 rf_fit <- train(Assigned ~., data = trainingXG2, method = "rf",
245               metric="Accuracy",
246               trcontrol=trctrl,
247               tuneLength = 13,
248               )
249
250 # variable importance
251
252 rf_vapIMP <- varImp(rf_fit)
253
254 # predict on new data
255
256 test_predict_rf <- predict(rf_fit, testingXG2)
257
258 #confusion matrix
259
260 confusion_matrix_rf <- confusionMatrix(test_predict_rf, as.factor(testingXG2$Assigned))
261
262 end.time <- Sys.time()
263
264 time.taken_rf <- end.time - start.time
265
266

```

Figure 15: Random Forest implementation

5.4 XGBoost

XGBoost (gradient boosted decision trees implementation) was developed using caret's manual tuning options to create an expanded grid search for the optimised parameters.

```

288 ##### XGB Tuned|
289
290 # Specify Train Control
291
292 trctrl <- trainControl(method = "cv", number = 10)
293
294 # Specify expanded tuning grid
295
296 tune_grid <- expand.grid(nrounds = 500,
297                       max_depth = 30,
298                       eta = 0.01,
299                       gamma = 1,
300                       colsample_bytree = 1,
301                       min_child_weight = 100,
302                       subsample = 1)
303
304 # fit model
305
306
307 xgb_fit <- train(Assigned ~., data = trainingXG2, method = "xgbTree",
308               trControl=trctrl,
309               tuneGrid = tune_grid,
310               tuneLength = 5,
311               )
312
313 # Variable importance
314
315 xgb_var_imp <- varImp(xgb_fit)
316
317 # predict
318
319 xgb_test_predict <- predict(xgb_fit, testingXG2)
320
321 # Confusion Matrix
322
323 confusion_matrix_xgb <- confusionMatrix(xgb_test_predict, as.factor(testingXG2$Assigned))
324

```

Figure 16: XGBoost implementation

5.5 Multinomial Regression

Multinomial regression was implemented utilising the same parameter optimisation features. However, decay is the critical parameters, and an appropriate search grid was selected. The goal of the decay parameter is to discourage overfitting on data sets and increase reproducibility.

```

466 ### MLM
467
468 tuneGrid_mml <- expand.grid(decay = seq(0, 1, by = 0.1))
469
470 mml_fit <- train(Assigned ~., data = trainingXG2, method = "multinom",
471                 trControl=trctrl,
472                 ,tuneGrid = tuneGrid_mml
473                 ,tuneLength = 13
474 )
475
476
477 mml_var_imp <- varImp(mml_fit)
478
479 mml_test_predict <- predict(mml_fit, testingXG2)
480
481 summary(mml_test_predict)
482
483 confuson_matrix_mml <- confusionMatrix(mml_test_predict, as.factor(testingXG2$Assigned))
484 confuson_matrix_mml
485

```

Figure 17: Multinomial Regression Implementation

6 Model Results

The variable importance (where available) is computed for each model in the processing steps. They are aggregated and visualised. Each model produces predictions on the test data set and the creation of a confusion matrix to allow comparisons of the accuracy, sensitivity, and specificity of each model.

```

1 # Variable Importance
2 library(data.table)
3 # VarImpo Output to Dataframe
4 mmlVarImp <- as.data.frame(setDT(mmlVarImp, keep.rownames = TRUE)[,])
5 xgbVarImp <- as.data.frame(setDT(xgbVarImp, keep.rownames = TRUE)[,])
6 rfVarImp <- as.data.frame(setDT(rfVarImp, keep.rownames = TRUE)[,])
7 colnames(rfVarImp)[1] <- 'Predictor'
8 colnames(xgbVarImp)[1] <- 'Predictor'
9 colnames(mmlVarImp)[1] <- 'Predictor'
10 colnames(rfVarImp)[2] <- 'Overall'
11
12 # Merge to single tabular view
13 VarImps <- merge(rfVarImp, mmlVarImp, by = 'Predictor')
14 VarImps <- merge(VarImps, xgbVarImp, by = 'Predictor')
15

```

Figure 18: Variable Importance

The confusion matrices were amalgamated to show the differential in performance across the accuracy, sensitivity, specificity, and kappa.

```

660
661 results <- rbind(confusion_matrix_rf$overall,
662                 confusion_matrix_mml$overall,
663                 confusion_matrix_xgb$overall,
664                 confusion_matrix_nb$overall,
665                 confusion_matrix_h2o$overall)
666 results <- as.data.frame(results)
667
668 colnames(results)[1] <- "Model"
669 Names <- as.data.frame(c(as.character('Random Forest'),
670                         as.character('Multinomial Regression'),
671                         as.character('XGBoost'),
672                         as.character('Naive Bayes'),
673                         as.character('H2O')))
674
675 Results <- cbind(Names, results)
676

```

Figure 19: Confusion Matrices