



ΠΑΝΕΠΙΣΤΗΜΙΟ ΘΕΣΣΑΛΙΑΣ
ΣΧΟΛΗ ΘΕΤΙΚΩΝ ΕΠΙΣΤΗΜΩΝ
ΤΜΗΜΑ ΠΛΗΡΟΦΟΡΙΚΗΣ ΜΕ ΕΦΑΡΜΟΓΕΣ
ΣΤΗ ΒΙΟΙΑΤΡΙΚΗ

**Επεξεργασία και ανάλυση σημάτων τηλεπαρακολούθησης
ασθενών για την πρόβλεψη ανεπιθύμητων συμβάντων με
μεθόδους μηχανικής μάθησης.**

Σιούζου Μαρία

Επεξεργασία και ανάλυση σημάτων τηλεπαρακολούθησης ασθενών για την πρόβλεψη ανεπιθύμητων συμβάντων με μεθόδους μηχανικής μάθησης. | ΣΙΟΥΖΟΥ ΜΑΡΙΑ



ΠΑΝΕΠΙΣΤΗΜΙΟ ΘΕΣΣΑΛΙΑΣ
ΣΧΟΛΗ ΘΕΤΙΚΩΝ ΕΠΙΣΤΗΜΩΝ
ΤΜΗΜΑ ΠΛΗΡΟΦΟΡΙΚΗΣ ΜΕ ΕΦΑΡΜΟΓΕΣ ΣΤΗ
ΒΙΟΙΑΤΡΙΚΗ

Επεξεργασία και ανάλυση σημάτων τηλεπαρακολούθησης ασθενών για την πρόβλεψη ανεπιθύμητων συμβάντων με μεθόδους μηχανικής μάθησης.

Σιούζου Μαρία

ΠΤΥΧΙΑΚΗ ΕΡΓΑΣΙΑ

Επιβλέπων

Ιακωβίδης Δημήτριος

Καθηγητής

Λαμία, 2023

Με ατομική μου ευθύνη και γνωρίζοντας τις κυρώσεις ⁽¹⁾, που προβλέπονται από της διατάξεις της παρ. 6 του άρθρου 22 του Ν. 1599/1986, δηλώνω ότι:

1. Δεν παραθέτω κομμάτια βιβλίων ή άρθρων ή εργασιών άλλων αυτολεξεί **χωρίς να τα περικλείω σε εισαγωγικά** και χωρίς να αναφέρω το συγγραφέα, τη χρονολογία, τη σελίδα. Η αυτολεξεί παράθεση χωρίς εισαγωγικά χωρίς αναφορά στην πηγή, είναι λογοκλοπή. Πέραν της αυτολεξεί παράθεσης, λογοκλοπή θεωρείται και η παράφραση εδαφίων από έργα άλλων, συμπεριλαμβανομένων και έργων συμφοιτητών μου, καθώς και η παράθεση στοιχείων που άλλοι συνέλεξαν ή επεξεργάστηκαν, χωρίς αναφορά στην πηγή. Αναφέρω πάντοτε με πληρότητα την πηγή κάτω από τον πίνακα ή σχέδιο, όπως στα παραθέματα.
2. Δέχομαι ότι η αυτολεξεί **παράθεση χωρίς εισαγωγικά**, ακόμα κι αν συνοδεύεται από αναφορά στην πηγή σε κάποιο άλλο σημείο του κειμένου ή στο τέλος του, είναι αντιγραφή. Η αναφορά στην πηγή στο τέλος π.χ. μιας παραγράφου ή μιας σελίδας, δεν δικαιολογεί συρραφή εδαφίων έργου άλλου συγγραφέα, έστω και παραφρασμένων, και παρουσίασή τους ως δική μου εργασία.
3. Δέχομαι ότι υπάρχει επίσης περιορισμός στο μέγεθος και στη συχνότητα των παραθεμάτων που μπορώ να εντάξω στην εργασία μου εντός εισαγωγικών. Κάθε μεγάλο παράθεμα (π.χ. σε πίνακα ή πλαίσιο, κλπ), προϋποθέτει ειδικές ρυθμίσεις, και όταν δημοσιεύεται προϋποθέτει την άδεια του συγγραφέα ή του εκδότη. Το ίδιο και οι πίνακες και τα σχέδια
4. Δέχομαι όλες τις συνέπειες σε περίπτωση λογοκλοπής ή αντιγραφής.

Ημερομηνία: 19/10/2023

Η Δηλούσα

Σιούζου Μαρία

(1) «Όποιος εν γνώσει του δηλώνει ψευδή γεγονότα ή αρνείται ή αποκρύπτει τα αληθινά με έγγραφη υπεύθυνη δήλωση του άρθρου 8 παρ. 4 Ν. 1599/1986 τιμωρείται με φυλάκιση τουλάχιστον τριών μηνών. Εάν ο υπαίτιος αυτών των πράξεων σκόπευε να προσπορίσει στον εαυτόν του ή σε άλλον περιουσιακό όφελος βλάπτοντας τρίτον ή σκόπευε να βλάψει άλλον, τιμωρείται με κάθειρξη μέχρι 10 ετών.

Επεξεργασία και ανάλυση σημάτων τηλεπαρακολούθησης ασθενών για την πρόβλεψη ανεπιθύμητων συμβάντων με μεθόδους μηχανικής μάθησης.

Σιούζου Μαρία

Τριμελής Επιτροπή:

Ιακωβίδης Δημήτριος, Καθηγητής

Δελήμπασης Κωνσταντίνος, Καθηγητής

Τασουλής Σωτήριος, Επίκουρος Καθηγητής

Ευχαριστίες

Σε αυτό το σημείο, θα επιθυμούσα να ευχαριστήσω όλους τους ανθρώπους που με στήριξαν και με βοήθησαν να φέρω εις πέρας την πτυχιακή μου εργασία. Αρχικά, θα ήθελα να ευχαριστήσω τον επιβλέποντα καθηγητή, κ. Δημήτριο Ιακωβίδη, καθηγητή του τμήματος Πληροφορικής με Εφαρμογές στη Βιοϊατρική του Πανεπιστημίου Θεσσαλίας, για την καθοδήγηση του ώστε να διεκπεραιωθεί η παρούσα εργασία. Επιπλέον, θα ήθελα να ευχαριστήσω το Μεταδιδακτορικό Ερευνητή Μιχάλη Βασιλακάκη για το χρόνο του και την ανεκτίμητη βοήθειά του ώστε να βγει το καλύτερο δυνατό αποτέλεσμα. Θέλω, επίσης, να ευχαριστήσω βαθιά όλους τους πολύ καλούς μου φίλους που ήταν δίπλα μου, με στήριξαν και έδειξαν κατανόηση όλα αυτά τα χρόνια. Τέλος, θα ήθελα να πω ένα μεγάλο ευχαριστώ στην οικογένειά μου η οποία υπήρξε πάντα ένα ανεκτίμητο στήριγμα για μένα και ειδικότερα στους γονείς μου, στους οποίους οφείλω όλη τη διαδρομή των σπουδών μου, μέχρι σήμερα.

Πίνακας Περιεχομένων

Ευχαριστίες.....	6
Πίνακας Περιεχομένων Εικόνων	10
Πίνακας Περιεχομένων Πινάκων.....	11
Κεφάλαιο 1	12
<i>Εισαγωγή</i>	12
1.1 Σήματα Τηλεπαρακολούθησης	12
1.2 Καρδιακή Ανεπάρκεια	12
1.3 Στόχοι Πτυχιακής Εργασίας και Ερευνητική Συνεισφορά	13
1.4 Δομή Πτυχιακής Εργασίας.....	14
Κεφάλαιο 2	15
<i>Θεωρητικό Υπόβαθρο</i>	15
2.1 Μηχανική Μάθηση (Machine Learning)	15
2.2 Κατηγορίες Μηχανικής Μάθησης	15
2.3 Αλγόριθμοι Μηχανικής Μάθησης.....	16
2.3.1 Αλγόριθμος K- Κοντινότερων Γειτόνων	16
2.3.2 Bayes.....	17
2.3.3 Δέντρα αποφάσεων	18
2.3.4 Τυχαία Δάση	19
2.3.5 Εξαιρετικά Τυχαία Δέντρα	20
2.3.6 Μηχανές Διανυσμάτων Στήριξης	20
2.4 Ενίσχυση	25
2.4.1 Ορισμός της τεχνικής της ενίσχυσης	25
2.4.2 Αλγόριθμοι ενίσχυσης	25
2.4.3 Προσαρμοστική Ενίσχυση.....	25
2.4.4 Ενίσχυση Κλίσης	27
2.4.5 Εξαιρετική Ενίσχυση Κλίσης	28
2.5 Μεθοδολογίες Ερμηνευσιμότητας.....	29
Κεφάλαιο 3	32
<i>Ασαφείς Ταξινομητές</i>	32
3.1 Ασαφής Λογική (Fuzzy Logic).....	32
3.2 Ασαφή Σύνολα (Fuzzy Sets)	32
3.3 Ασαφείς Κανόνες Ταξινόμησης	35

3.4 Ασαφή Συστήματα.....	35
3.4.1 Ασαφείς Ταξινομητές – Takagi-Sugeno	37
3.4.2 Προσαρμοστικό Νευροασαφές Σύστημα Συμπερασμάτων βασισμένο σε Αρχιτεκτονική Takagi –Sugeno	39
Κεφάλαιο 4	42
Βιβλιογραφική Έρευνα	42
4.1 Βιβλιογραφική Ανασκόπηση Μεθόδων	42
4.1.1 Βιβλιογραφική Ανασκόπηση Μεθόδων Μηχανικής Μάθησης	42
4.1.2 Βιβλιογραφική Ανασκόπηση Ασαφών Μοντέλων	44
Κεφάλαιο 5	46
Μεθοδολογία.....	46
5.1 Μέθοδος Ασαφών Φράσεων Ομοιότητας (Fuzzy Similarity Phrases- FSP).....	46
5.1.1 Κατασκευή Μοντέλου	46
5.1.2 Δημιουργία Φράσεων για Ερμηνεύσιμη Ταξινόμηση.....	52
5.1.3 Επιλογή Χαρακτηριστικών με μείωση Φράσεων	54
5.1.4 Δημιουργία Κανόνων	56
5.1.5 Επαλήθευση Μοντέλου και Δημιουργία Βάσης Κανόνων.....	57
5.1.6 Ταξινόμηση ενός διανύσματος άγνωστης κλάσης.....	59
5.1.7 Ερμηνεία της Ταξινόμησης.....	61
5.2 Παραδείγματα Βημάτων Υλοποίησης της Μεθόδου	64
Κεφάλαιο 6	72
Αποτελέσματα και Πειράματα	72
6.1 Μετρικές Αξιολόγησης	72
6.2 Προκαταρκτική Μελέτη.....	73
6.3 Πειραματικά Αποτελέσματα Μεθοδολογίας FSP	76
Κεφάλαιο 7	80
Καρδιολογικά Δεδομένα.....	80
7.1 Δεδομένα καρδιολογικών σημάτων.....	80
7.1.1 Περιγραφή συνόλου Δεδομένων (DREAMER)	80
7.1.2 Αποτελέσματα Συνόλου Δεδομένων DREAMER.....	82
7.2 Δεδομένα νόσων της καρδιάς στη Νότια Αφρική (SAHeart-SAH).....	90
7.2.1 Περιγραφή Συνόλου Δεδομένων SAHeart	90
7.2.2 Αποτελέσματα Συνόλου Δεδομένων SAHeart	91

7.3 Δεδομένα αξονικής τομογραφίας εκπομπής μονήρους φωτονίου της καρδιάς (SPECTF heart-SPEC).....	94
7.3.1 Περιγραφή Συνόλου Δεδομένων SPECTF heart	94
7.3.2 Αποτελέσματα Συνόλου Δεδομένων SPECTF Heart	95
7.4 Δεδομένα νόσων της καρδιάς (Statlog heart).....	97
7.4.1 Περιγραφή Συνόλου Δεδομένων Statlog Heart	97
7.4.2 Αποτελέσματα Συνόλου Δεδομένων Statlog Heart	98
Κεφάλαιο 8	101
Συμπεράσματα και Προοπτικές	101
8.1 Συμπεράσματα.....	101
8.2 Μελλοντικές Προοπτικές	103
Βιβλιογραφία	104

Πίνακας Περιεχομένων Εικόνων

Εικόνα 1 Παράδειγμα εφαρμογής KNN αλγορίθμου [29].	17
Εικόνα 2 Τυχαίο Δάσος με τελική πρόβλεψη: κλάση 0. [20].	19
Εικόνα 3 Βέλτιστα υπερεπίπεδα διαχωρισμού σε διαφορετικές διαστάσεις [25].	21
Εικόνα 4 Γραμμικός διαχωρισμός υπερεπιπέδων. Τα κυκλωμένα σημεία αποτελούν τα διανύσματα στήριξης.	22
Εικόνα 5 Λειτουργία αλγορίθμου Προσαρμοστικής Ενίσχυσης (Adaboost).	26
Εικόνα 6 Λειτουργία αλγορίθμου ενίσχυσης κλίσης.	28
Εικόνα 7 Τοπική ερμηνευσιμότητα χαρακτηριστικών με τη μέθοδο SHAP.	30
Εικόνα 8 Σχηματική αναπαράσταση Συναρτήσεων συμμετοχής	34
Εικόνα 9 Δομή ενός ασαφούς συστήματος.	36
Εικόνα 10 Σχηματική αναπαράσταση πρώτου βαθμού Takagi-Sugeno. [101]	39
Εικόνα 11 Αρχιτεκτονική Takagi-Sugeno ANFIS.	40
Εικόνα 12 Συναρτήσεις συμμετοχής των εξαγόμενων ασαφών συνόλων.	51
Εικόνα 13 Διαγραμματική απεικόνιση κατασκευής μοντέλου.	52
Εικόνα 14 Διαγραμματική απεικόνιση δημιουργίας φράσεων.	54
Εικόνα 15 Διαγραμματική απεικόνιση μείωσης φράσεων	56
Εικόνα 16 Διαγραμματική απεικόνιση για την κατασκευή της ανανεωμένης βάσης κανόνων.	59
Εικόνα 17 Παράδειγμα οπτικοποίησης της μεθοδολογίας.	63
Εικόνα 18 Αρχική μορφή συνόλου δεδομένων καρδιακής ανεπάρκειας	64
Εικόνα 19 1 8-διαστατό κέντρο συστάδας.	65
Εικόνα 20 Παραγόμενα διανύσματα ομοιοτήτων.	67
Εικόνα 21 Συναρτήσεις συμμετοχής για την ανάπτυξη των ασαφών συνόλων.	68
Εικόνα 22 Τριπλέτες εξαγόμενες από το βήμα 5.	69
Εικόνα 23 Μορφή διανυσμάτων μετά τη φάση της πρόσθεσης βαρών και τη μείωση των φράσεων.	69
Εικόνα 24 Μορφή βάσης μοναδικών κανόνων (RB).	70
Εικόνα 25 Βάση Κανόνων έπειτα από την εφαρμογή του αλγορίθμου GD	71
Εικόνα 26 Τελικό αποτέλεσμα μοντέλου.	71
Εικόνα 27 Ποσοστά ακρίβειας σε σετ δεδομένων σχετικά με παθήσεις καρδιάς.	74
Εικόνα 28 Ποσοστά ακρίβειας σε βιολογικό σετ δεδομένων.	75
Εικόνα 29 Ποσοστά ακρίβειας σε σετ δεδομένων σχετικά με την ασθένεια του διαβήτη.	75
Εικόνα 30 Το μοντέλο VAD στο χώρο ως προς τα 6 πιο βασικά συναισθήματα.	81
Εικόνα 31 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.	84
Εικόνα 32 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.	87
Εικόνα 33 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.	89
Εικόνα 34 Σύγκριση τελικών αποτελεσμάτων μεθοδολογιών με τα καλύτερα αποτελέσματα.	89

Εικόνα 35 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.....	93
Εικόνα 36 Σύγκριση αποτελεσμάτων	94
Εικόνα 37 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.....	96
Εικόνα 38 Σύγκριση διάφορων εκδόσεων.	97
Εικόνα 39 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.....	100
Εικόνα 40 Αποτελέσματα διάφορων εκδόσεων.....	100

Πίνακας Περιεχομένων Πινάκων

Πίνακας 1 Αλγόριθμοι συσταδοποίησης που δοκιμάστηκαν	47
Πίνακας 2 Εκδοχές με διάφορες ομοιότητες που δοκιμάστηκαν	48
Πίνακας 3 Εκδοχές που δοκιμάστηκαν με συναρτήσεις συμμετοχής	50
Πίνακας 4 Σύγκριση αποτελεσμάτων παρούσας μεθοδολογίας με άλλους ασαφείς ταξινομητές [104].	77
Πίνακας 5 Σύγκριση αποτελεσμάτων με άλλους ταξινομητές [105].....	78
Πίνακας 6 Μέση τιμή και τυπική απόκλιση από όλους του συμμετέχοντες για κάθε συναίσθημα και για κάθε κλίμακα βαθμολόγησης.	81
Πίνακας 7 Πληροφορίες του συνόλου δεδομένων.	82
Πίνακας 8 Αποτελέσματα ακρίβειας και AUC.....	83
Πίνακας 9 Αποτελέσματα ακρίβειας και AUC.....	86
Πίνακας 10 Αποτελέσματα ακρίβειας και AUC.....	88
Πίνακας 11 Χαρακτηριστικά συνόλου Δεδομένων SAHeart.	91
Πίνακας 12 Αποτελέσματα ακρίβειας και AUC.....	92
Πίνακας 13 Αναλυτικά αποτελέσματα όλων των εκδόσεων.	95
Πίνακας 14 Χαρακτηριστικά Συνόλου Δεδομένων Statlog Heart.....	98
Πίνακας 15 Αναλυτικά αποτελέσματα όλων των εκδόσεων.	99

Κεφάλαιο 1

Εισαγωγή

1.1 Σήματα Τηλεπαρακολούθησης

Ένα μέσο πρόληψης αλλά και πρόβλεψης συμβάντων καρδιακής ανεπάρκειας αποτελούν τα σήματα τηλεπαρακολούθησης ασθενών. Η τηλεπαρακολούθηση ασθενών είναι μία μέθοδος που αποσκοπεί στη βελτίωση της κατάστασης υγείας των ασθενών με καλύτερη ποιότητα ζωής [8]. Τα σήματα αυτά δίνουν πληροφορία στους παρόχους υγείας για την πορεία του ασθενούς. Στα σήματα τηλεπαρακολούθησης ανήκουν ο καρδιακός παλμός, η αρτηριακή πίεση (συστολική/διαστολική), το σωματικό βάρος αλλά και ο κορεσμός οξυγόνου. Είναι δηλαδή μετρήσεις οι οποίες μπορούν να παρακολουθούνται μέσω της τεχνολογίας και να μεταδίδονται στον πάροχο φροντίδας. Οι χρησιμοποιούμενες τεχνολογίες είναι τηλεφωνικές γραμμές, ευρυζωνικά, δορυφορικά ή ασύρματα δίκτυα. Ο εκάστοτε πάροχος φροντίδας είναι πλέον σε θέση να αξιολογήσει την πορεία της υγείας του ασθενούς αλλά και να χορηγήσει πιθανή φαρμακευτική συνταγή.

Τα δεδομένα που συλλέγονται από τα σήματα αυτά αποτελούν θεμέλια για την ανάπτυξη μοντέλων πρόβλεψης. Μοντέλα, δηλαδή, που χρησιμοποιούν αλγορίθμους μηχανικής μάθησης για την πρόβλεψη ανεπιθύμητων συμβάντων συμπεριλαμβανομένης και της καρδιακής ανεπάρκειας. Τα τελευταία χρόνια έχουν αναπτυχθεί διάφορα τέτοια μοντέλα με σκοπό την πρόβλεψη της [9],[10].

1.2 Καρδιακή Ανεπάρκεια

Οι καρδιαγγειακές παθήσεις αποτελούν την κυριότερη αιτία θανάτου παγκοσμίως με τον αριθμό των θανάτων να ξεπερνούν περίπου τα 17.9 εκατομμύρια ετησίως [1]. Αυτές οι παθήσεις αφορούν επιπλοκές που σχετίζονται με την καρδιά και τα αιμοφόρα αγγεία και περιλαμβάνουν πολλών ειδών ασθένειες μία εκ των οποίων είναι και η καρδιακή ανεπάρκεια. Η πάθηση της καρδιακής ανεπάρκειας ορίζεται ως το κλινικό σύνδρομο κατά το οποίο η καρδιά αδυνατεί να εξωθήσει την απαραίτητη ποσότητα αίματος στο σώμα για να εκτελέσει τις φυσιολογικές του λειτουργίες, με αποτέλεσμα να εμφανίζονται συμπτώματα όπως δύσπνοια, κόπωση και κατακράτηση υγρών [2]. Σήμερα ο αριθμός των ασθενών που πάσχουν από καρδιακή ανεπάρκεια ξεπερνά τα 26 εκατομμύρια παγκοσμίως και η ανησυχία των γιατρών συνεχώς αυξάνεται [3]. Για αυτό πραγματοποιούνται αρκετές κλινικές προσπάθειες με σκοπό την πρόβλεψή της. Ωστόσο, οι

προσπάθειες αυτές έχουν αποτύχει, καθώς τα ποσοστά ακρίβειας κυμαίνονται σε χαμηλό επίπεδο [4].

Με δεδομένη τη συχνότητα εμφάνισης των καρδιαγγειακών παθήσεων και την ανησυχία των γιατρών συνεχώς συλλέγονται δεδομένα από ασθενείς, παθήσεις, νοσηλείες, νοσοκομεία ακόμη και έξοδα ιατρικού εξοπλισμού με σκοπό την ανάλυσή τους. Αυτή η συλλογή πραγματοποιείται είτε με παρατηρήσεις από ερευνητές είτε με περισυλλογή από τον ηλεκτρονικό φάκελο υγείας είτε από εγγραφές νοσοκομείων [5]. Η επεξεργασία και ανάλυση αυτών των δεδομένων εκτελείται με μεθόδους μηχανικής μάθησης και εξόρυξης δεδομένων με σκοπό την ανάπτυξη συστημάτων στήριξης ιατρικών αποφάσεων με υψηλή ακρίβεια [6],[7]. Οι μέθοδοι αυτοί λοιπόν έχουν την ικανότητα να προβλέπουν τον κίνδυνο προσβολής από όλες αυτές τις παθήσεις συμπεριλαμβάνοντας και την καρδιακή ανεπάρκεια. Ένα τέτοιο σύστημα στήριξης ιατρικής απόφασης αναπτύχθηκε και στην παρούσα εργασία.

1.3 Στόχοι Πτυχιακής Εργασίας και Ερευνητική Συνεισφορά

Στην παρούσα ερευνητική εργασία βασικός στόχος είναι η διερεύνηση ανάλυσης δεδομένων και η εφαρμογή τους για την ανάλυση σημάτων στο πλαίσιο συστημάτων στήριξης ιατρικών αποφάσεων. Η έρευνα εστιάζει σε καρδιολογικά σήματα που μπορούν να συλλεχθούν με τη μέθοδο της τηλεπαρακολούθησης, με σκοπό την πρόβλεψη ανεπιθύμητων συμβάντων, όπως είναι η πρόβλεψη της καρδιακής ανεπάρκειας. Για το σκοπό αυτό, μελετήθηκαν και δοκιμάστηκαν πειραματικά αλγόριθμοι μηχανικής μάθησης που έχουν εφαρμογή για τη λήψη αποφάσεων στα εξεταζόμενα σήματα. Επιπροσθέτως, στους εξεταζόμενους αλγορίθμους μηχανικής μάθησης εφαρμόστηκαν μεθοδολογίες που είναι ερμηνεύσιμες (interpretable) με σκοπό την επεξήγηση του εξαγόμενου αποτελέσματος.

Με βάση την προσπάθεια για την πραγματοποίηση των αναφερόμενων στόχων, στα πλαίσια της παρούσας πτυχιακής εργασίας μελετήθηκε η ανάπτυξη ενός ερμηνεύσιμου μοντέλου με μεγάλη ακρίβεια στις προβλέψεις σε άγνωστα δεδομένα, βάση της ασαφούς λογική.

Η παρούσα πτυχιακή εργασία, βασίζεται στη μεθοδολογία που αναπτύχθηκε από τους M. D. Vasilakakis & D. K. Iakovidis [105]. Η ερευνητική συνεισφορά της παρούσας πτυχιακής εργασίας συνοψίζεται:

- Στη βελτίωση της υλοποίησης της παραπάνω μεθοδολογίας και μετατροπή της από τη γλώσσα προγραμματισμού MATLAB σε γλώσσα προγραμματισμού Python, εστιάζοντας στη δυνατότητα επαναχρησιμοποίησης του κώδικα
- Στη διερεύνηση εναλλακτικών συστατικών και παραμετροποιήσεων της μεθοδολογίας με σκοπό τη βελτίωση των αποτελεσμάτων σε δημοσίως διαθέσιμα δεδομένα, με έμφαση σε μεθόδους τηλεπαρακολούθησης
- Στη συστηματική και λεπτομερή πειραματική αξιολόγηση της υλοποίησης της μεθοδολογίας συγκριτικά με τα αποτελέσματα άλλων σύγχρονων σχετικών μεθοδολογιών

1.4 Δομή Πτυχιακής Εργασίας

Η παρούσα εργασία εκτείνεται σε επτά κεφάλαια. Στο δεύτερο κεφάλαιο περιγράφονται θεωρητικά και με αναλυτικά παραδείγματα γνωστοί αλγόριθμοι μηχανικής μάθησης. Στο τρίτο κεφάλαιο περιγράφονται οι ασαφείς ταξινομητές (Fuzzy Classifiers), η λειτουργία τους και δίνεται βάση στους ασαφείς ταξινομητές Takagi Sugeno που χρησιμοποιήθηκαν περισσότερο στην πειραματική αξιολόγηση. Στο τέταρτο κεφάλαιο παρουσιάζεται η βιβλιογραφική έρευνα για μεθόδους μηχανικής μάθησης, ασαφών μοντέλων, καθώς και τα αντίστοιχα πειραματικά αποτελέσματα σε μεθοδολογίες όπου η υλοποίηση ήταν διαθέσιμη. Στο πέμπτο κεφάλαιο παρουσιάζεται η προτεινόμενη μεθοδολογία καθώς και τα βήματα υλοποίησης της. Στο έκτο κεφάλαιο παρουσιάζονται όλα τα αποτελέσματα που εξήχθησαν από τη μεθοδολογία. Στο έβδομο κεφάλαιο παρουσιάζονται τα αποτελέσματα σε καρδιολογικά δεδομένα και τέλος στο όγδοο κεφάλαιο αναφέρονται τα συμπεράσματα από όλη τη μελέτη και πιθανές μελλοντικές επεκτάσεις της.

Κεφάλαιο 2

Θεωρητικό Υπόβαθρο

2.1 Μηχανική Μάθηση (Machine Learning)

Η μηχανική μάθηση (Machine Learning) αποτελεί κλάδο της Τεχνητής Νοημοσύνης που συνδυάζει την επιστήμη της Πληροφορικής και της Στατιστικής [11]. Ορίζεται ως το πεδίο της επιστήμης που δίνει τη δυνατότητα στους υπολογιστές να «μαθαίνουν» χωρίς να χρειάζεται να προγραμματιστούν [12]. Η μηχανική μάθηση χρησιμοποιείται για να διδάξει στις μηχανές πως να διαχειρίζονται τα δεδομένα πιο αποδοτικά. Έτσι, προγραμματίζονται υπολογιστές με στόχο τη βελτιστοποίηση ενός κριτηρίου χρησιμοποιώντας δείγμα δεδομένων ή εμπειρία από το παρελθόν [13]. Στόχος της μηχανικής μάθησης είναι να μαθαίνει από τα δεδομένα. Έτσι αναπτύχθηκαν διάφοροι αλγόριθμοι που πρεσβεύουν τους σκοπούς της. Η επιλογή του κατάλληλου αλγορίθμου κάθε φορά εξαρτάται από το πρόβλημα που πρέπει να αντιμετωπιστεί και τα δεδομένα που παρέχονται για την επίλυση του προβλήματος [11].

Οι εφαρμογές της μηχανικής μάθησης είναι ποικίλες. Σε αυτές συμπεριλαμβάνονται η υποβοήθηση για τη διάγνωση ασθενειών, η αναγνώριση εικόνων, οικονομικές προβλέψεις και διάφορες άλλες που αφορούν την καθημερινότητα [14].

2.2 Κατηγορίες Μηχανικής Μάθησης

Σύμφωνα με τον τρόπο που το σύστημα επεξεργάζεται τα δεδομένα εισόδου κατά την εκπαίδευση υπάρχουν διαφορετικές προσεγγίσεις μηχανικής μάθησης, μεταξύ των οποίων οι κυριότερες είναι οι εξής:

1. **Επιβλεπόμενη Μάθηση (Supervised Learning):** Η επιβλεπόμενη μάθηση προσεγγίζει μία συνάρτηση η οποία αντιστοιχίζει μια είσοδο σε μία έξοδο βάσει προκαθορισμένων ζευγών εισόδου-εξόδου. Η διαδικασία αυτή της προσέγγισης ονομάζεται εκπαίδευση. Σκοπός όλων των αλγορίθμων που ανήκουν σε αυτήν την κατηγορία είναι να μάθουν κάποια μοτίβα από το σύνολο εκπαίδευσης ώστε μετά να το εφαρμόσουν στο σύνολο δοκιμής για πρόβλεψη ή ταξινόμηση [11].

2. **Μη επιβλεπόμενη Μάθηση (Unsupervised Learning):** Στη μη επιβλεπόμενη μάθηση δεν υπάρχουν σωστές απαντήσεις ούτε διδασκαλία. Οι αλγόριθμοι πλέον μόνοι τους έρχονται αντιμέτωποι να ανακαλύψουν τη δομή των δεδομένων εισόδου. Στην κατηγορία αυτή ανήκει η τεχνική της συσταδοποίησης (clustering) με την οποία δημιουργούνται συστάδες δεδομένων εισόδου, των οποίων τα χαρακτηριστικά είναι παρόμοια και διαφέρουν από αυτά άλλων συστάδων [11].
3. **Ενισχυτική Μάθηση (Reinforcement Learning):** Στην ενισχυτική μάθηση το σύστημα προσπαθεί να μάθει με την αλληλεπίδραση με το περιβάλλον. Δέχεται «επιβραβεύσεις» ή «τιμωρίες» κατά τη διάρκεια της εκπαίδευσης ώστε να οδηγηθεί στην επιθυμητή κατεύθυνση [15].

Στην παρούσα εργασία οι αλγόριθμοι που υλοποιήθηκαν αφορούν αποκλειστικά την κατηγορία της επιβλεπόμενης μάθησης.

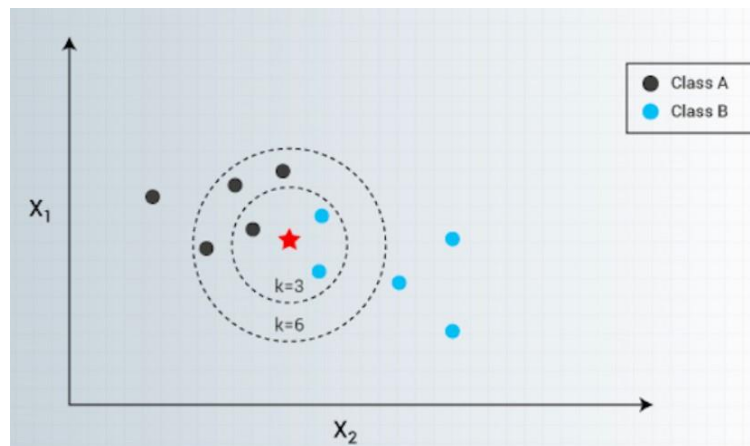
2.3 Αλγόριθμοι Μηχανικής Μάθησης

Παρακάτω αναλύονται οι αλγόριθμοι μηχανικής μάθησης που χρησιμοποιήθηκαν για την παρούσα εργασία.

2.3.1 Αλγόριθμος K- Κοντινότερων Γειτόνων

Ο αλγόριθμος των K-κοντινότερων γειτόνων (K-Nearest Neighbors-KNN) είναι ένας απλός και εύκολος στην εφαρμογή αλγόριθμος που ανήκει στην κατηγορία εποπτευόμενης μηχανικής μάθησης. Ο συγκεκριμένος αλγόριθμος βασίζεται στην θεωρία πως τα χαρακτηριστικά σε ένα σύνολο δεδομένων θα υπάρχουν σε κοντινή απόσταση από άλλα χαρακτηριστικά που έχουν παρόμοιες ιδιότητες [27].

Έστω λοιπόν πως είναι γνωστή η κλάση στην οποία ανήκουν N δείγματα (διανύσματα χαρακτηριστικών) και ένα δείγμα x άγνωστης κλάσης. Στόχος είναι να ταξινομηθεί το x σύμφωνα με τον αλγόριθμο K κοντινότερων γειτόνων. Υπολογίζονται όλες οι αποστάσεις του x από όλα τα υπόλοιπα δείγματα. Συλλέγονται όλες αυτές οι αποστάσεις και ταξινομούνται σε αύξουσα σειρά. Το δείγμα x θα ταξινομηθεί στην κλάση A αν η πλειοψηφία των k κοντινότερων γειτόνων του x ανήκουν στην κλάση A [28]. Ακολουθεί μία εικόνα με ένα παράδειγμα εφαρμογής του αλγορίθμου (Εικόνα 1).



Εικόνα 1 Παράδειγμα εφαρμογής KNN αλγορίθμου [29].

Υπάρχουν δύο διαφορετικές κλάσεις: η κλάση Α (μαύρο χρώμα) και η κλάση Β (μπλε χρώμα). Στόχος είναι να κατηγοριοποιηθεί το αστερί σε μία από τις δύο κλάσεις. Ορίζεται το $k=3$ από το χρήστη με στόχο την εύρεση των 3 κοντινότερων γειτόνων του αστεριού. Έχοντας υπολογίσει τις αποστάσεις του αστεριού από όλα τα σημεία όλων των κλάσεων φαίνονται από την Εικόνα 1 οι 3 κοντινότεροι γείτονες. Το αστερί λοιπόν θα ταξινομηθεί στην κλάση Β (μπλε χρώμα) καθώς η πλειοψηφία των 3 κοντινότερων γειτόνων (δηλ. 2 σημεία) ανήκει στην κλάση Β. Αντίστοιχα, αν ο χρήστης επιλέξει το $k=6$ το αστερί θα ταξινομηθεί στην κλάση Α.

2.3.2 Bayes

Ένας ακόμη ταξινομητής απλός στην εφαρμογή του είναι ο ταξινομητής που βασίζεται στο θεώρημα του Bayes (Naïve Bayes). Ο Bayes ανήκει στους στατιστικούς ταξινομητές οι οποίοι προβλέπουν το αποτέλεσμα αποκλειστικά με τη χρήση πιθανοτήτων. Στηρίζεται, επίσης, στην υπόθεση πως οι τιμές των χαρακτηριστικών που ανήκουν σε μία κλάση είναι ανεξάρτητες από εκείνες που ανήκουν σε άλλη κλάση [29]. Αυτή η υπόθεση ονομάζεται ανεξαρτησία κλάσεων υπό συνθήκη. Έχει κατασκευαστεί με σκοπό να απλοποιήσει τον εμπλεκόμενο υπολογισμό και υπό αυτήν την έννοια θεωρείται «αφελής».

Το θεώρημα του Bayes συσχετίζει την πιθανότητα του αρχικού γεγονότος που έχει συμβεί με ένα γεγονός που έχει συμβεί νωρίτερα. Υπολογίζει την εκ των υστέρων πιθανότητα για κάθε κατηγορία

c_i και επιλέγει την κατηγορία που έχει τη μεγαλύτερη πιθανότητα. Σύμφωνα, λοιπόν με το θεώρημα αυτό ισχύει:

$$P(C_i|X_i) = \frac{P(X_i \cap C_i)}{P(X_i)} = \frac{P(X_i|C_i) \cdot P(C_i)}{P(X_i)} [30]$$

Όπου X_i είναι τα αποδεικτικά στοιχεία, C_i είναι οι υποθέσεις, $P(C_i)$ είναι η εκ των προτέρων πιθανότητα των υποθέσεων, $P(X_i|C_i)$ είναι υπό συνθήκη πιθανότητα και $P(C_i|X_i)$ είναι η εκ των υστέρων πιθανότητα που πρόκειται να υπολογιστεί.

Για παράδειγμα έστω ένα σύνολο συμπτωμάτων για ένα συγκεκριμένο περιστατικό. Επομένως όπου $X_1, X_2, \dots, X_n$ τα συμπτώματα. Στόχος είναι να βρεθεί η υπόθεση C_i (π.χ. ασθένεια) με την μεγαλύτερη εκ των υστέρων πιθανότητα: $\text{Max}(P(C_i|X_i))$.

Ο συγκεκριμένος ταξινομητής έχει προταθεί και για την επίλυση προβλημάτων που αφορούν τις καρδιαγγειακές ασθένειες [31]. Ανταποκρίνεται γρήγορα στη διάγνωσή τους χωρίς μεγάλο υπολογιστικό κόστος.

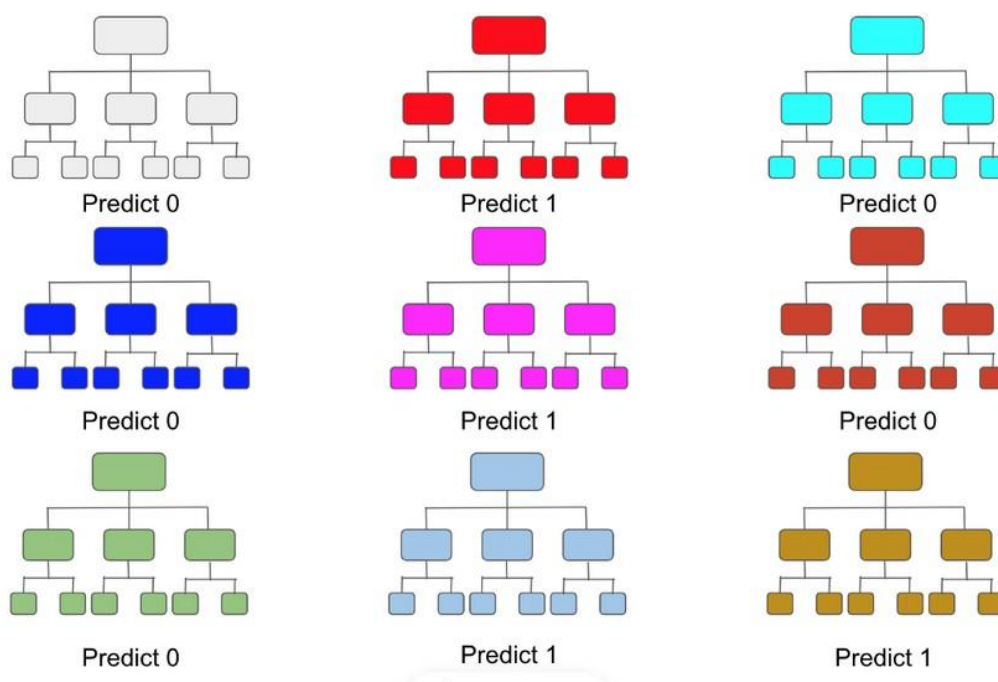
2.3.3 Δέντρα αποφάσεων

Τα δέντρα αποφάσεων (decision trees) είναι μία μέθοδος ταξινόμησης που χρησιμοποιείται σε εργασίες εξόρυξης δεδομένων. Η έξοδος από ένα δέντρο αποφάσεων είναι ένα γράφημα με κλαδιά και κόμβους και αναπαριστά μια πιθανή απόφαση για μία πρόβλεψη. Το γράφημα αυτό είναι ένα καθοδηγούμενο δέντρο το οποίο αποτελείται από μία μοναδική ρίζα η οποία δεν έχει εισερχόμενες ακμές. Εκτός από τη ρίζα υπάρχουν και οι κόμβοι. Ένας κόμβος με εξερχόμενες ακμές ονομάζεται εσωτερικός κόμβος. Όλοι οι υπόλοιποι κόμβοι ονομάζονται φύλα του δέντρου [16]. Σε ένα τέτοιο δέντρο αποφάσεων κάθε εσωτερικός κόμβος χωρίζει το δέντρο σε δύο ή περισσότερα υποδέντρα σύμφωνα με μία διακριτή συνάρτηση της τιμής των χαρακτηριστικών εισόδου. Συνήθως ένας εσωτερικός κόμβος αναφέρεται σε ένα μοναδικό χαρακτηριστικό τέτοιο ώστε ο χώρος περιπτώσεων να είναι χωρισμένος σύμφωνα με την τιμή του.

Μια δοκιμή σε ένα χαρακτηριστικό αναπαρίσταται με έναν εσωτερικό κόμβο του δέντρου, το αποτέλεσμα της δοκιμής αναπαρίσταται από τα υποδέντρα και τέλος η κλάση παρουσιάζεται με τα φύλα του δέντρου. Τα μονοπάτια από τη ρίζα στα φύλα του δέντρου αποτελούν τους κανόνες της ταξινόμησης [16].

2.3.4 Τυχαία Δάση

Ο ταξινομητής τυχαίου δάσους (random forest) αποτελεί ένα σύνολο από δέντρα αποφάσεων που ονομάζεται δάσος. Ο ταξινομητής τυχαίου δάσους ακολουθεί τη μέθοδο bootstrap για την εκπαίδευση των παρατηρήσεων. Σύμφωνα, λοιπόν, με αυτή τη μέθοδο κάθε δέντρο στα τυχαία δάση αντί να εκπαιδεύεται πάνω σε όλες τις παρατηρήσεις, εκπαιδεύεται σε ένα υποσύνολο των παρατηρήσεων. Το κάθε δέντρο δημιουργείται από ένα τυχαία επιλεγμένο διάνυσμα δειγματοληπτημένο ανεξάρτητα από το διάνυσμα εισόδου. Όσο περισσότερα δέντρα σχηματίζονται τόσο πιο ισχυρό είναι το σύνολο που μελετάται [17]. Κάθε δέντρο αποφασίζει για την πιο δημοφιλή κλάση για να ταξινομηθεί ένα επερχόμενο διάνυσμα εισόδου [18]. Τελικά η κλάση που έχει μαζέψει τις περισσότερες «ψήφους» από τα δέντρα δηλαδή αυτή που έχει προβλεφθεί τις περισσότερες φορές αποτελεί και το τελικό αποτέλεσμα του αλγορίθμου [19]. Ένα παράδειγμα φαίνεται στην παρακάτω Εικόνα 2 όπου φαίνονται 9 δέντρα αποφάσεων εκ των οποίων 4 προέβλεψαν την κατηγορία 1 και 5 την κατηγορία 0. Συνεπώς η τελική πρόβλεψη είναι η κατηγορία 0.



Εικόνα 2 Τυχαίο Δάσος με τελική πρόβλεψη: κλάση 0. [20]

2.3.5 Εξαιρετικά Τυχαία Δέντρα

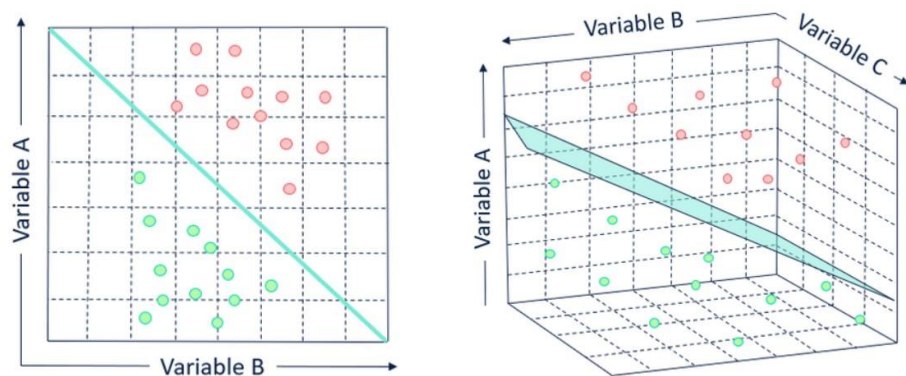
Ο ταξινομητής εξαιρετικά τυχαίων δέντρων (eXtremely Randomized Trees- Extra Trees) παρουσιάζει αρκετές ομοιότητες με τον προαναφερθείσα, αυτόν των τυχαίων δασών. Πιο συγκεκριμένα αποτελείται κι αυτός από πολλαπλά δέντρα αποφάσεων τα οποία δημιουργούνται από τυχαία επιλεγμένα δεδομένα. Παρουσιάζει, ωστόσο, και κάποιες σημαντικές διαφορές συγκριτικά με τα τυχαία δάση. Δεν ακολουθεί, αρχικά, τη μέθοδο bootstrap για τις παρατηρήσεις, γεγονός το οποίο αποκαλύπτει πως ο αλγόριθμος χρησιμοποιεί δείγματα χωρίς αντικατάσταση. Χρησιμοποιεί δηλαδή ολόκληρο το δείγμα της εκπαίδευσης κάθε φορά για να αναπτύξει το κάθε δέντρο [21]. Επίσης, οι κόμβοι σε κάθε δέντρο χωρίζονται βασιζόμενοι σε τελείως τυχαίους διαχωρισμούς ενός τυχαία επιλεγμένου υποσυνόλου των χαρακτηριστικών. Αυτή η τυχειότητα, λοιπόν, δεν οφείλεται στη μέθοδο bootstrap αλλά πηγάζει από τους τυχαίους διαχωρισμούς των δειγμάτων. Ο αλγόριθμος των Εξαιρετικά Τυχαίων Δέντρων εν ολίγοις προσπαθεί αντί να βρει ένα βέλτιστο σημείο κοπής για καθένα από τα τυχαία επιλεγμένα χαρακτηριστικά K κάθε κόμβου, να επιλέξει ένα σημείο κοπής τελείως τυχαία [22].

Αυτή η μέθοδος συχνά οδηγεί σε μεγάλη ακρίβεια χάρη στην εξομάλυνση, όπως αποδεικνύεται στη συνέχεια και από τα πειράματα της παρούσας έρευνας. Μειώνει ταυτόχρονα σημαντικά την υπολογιστική επιβάρυνση η οποία οφείλεται στον προσδιορισμό του βέλτιστου σημείου κοπής στα δέντρα. Έχει αποδειχθεί από μελέτες πως η συγκεκριμένη μέθοδος αποφέρει εξαιρετικά αποτελέσματα σε πολύπλοκα προβλήματα [21].

2.3.6 Μηχανές Διανυσμάτων Στήριξης

Η βασική ιδέα γύρω από τον ταξινομητή μηχανών διανυσμάτων στήριξης (Support Vector Machines-SVM) είναι ακριβώς ίδια με αυτή των πολλαπλών στρώσεων νευρωνικών δικτύων. Βασικό χαρακτηριστικό των μηχανών διανυσμάτων στήριξης αποτελούν οι συναρτήσεις πυρήνα (kernel functions) οι οποίες μπορεί να είναι πολυωνυμικές, εκθετικές ή γραμμικές. Σκοπός των συναρτήσεων αυτών είναι να μετασχηματίσουν τα δεδομένα από έναν αρχικό χώρο χαρακτηριστικών χαμηλής διάστασης σε έναν νέο χώρο υψηλότερης διάστασης. Με αυτόν τον τρόπο καθίσταται εφικτός ο γραμμικός διαχωρισμός των δεδομένων ώστε να βρεθεί το βέλτιστο υπερεπίπεδο (hyperplane) διαχωρισμού τους [23]. Στη γεωμετρία το υπερεπίπεδο είναι ένας

υπόχωρος μίας διάστασης λιγότερης από τον περιβάλλοντα χώρο [24]. Ο διαχωρισμός των κλάσεων και η ταξινόμηση νέων δεδομένων ορίζεται από τα υπερεπίπεδα. Τα υπερεπίπεδα καθορίζουν τα όρια απόφασης για την ταξινόμηση νέων σημείων. Οι υπόχωροι που προσδιορίζονται από τα υπερεπίπεδα ανήκουν σε διαφορετικές κατηγορίες όπως φαίνεται παρακάτω στην Εικόνα 3.



Εικόνα 3 Βέλτιστα υπερεπίπεδα διαχωρισμού σε διαφορετικές διαστάσεις [25].

Δοθέντος ενός συνόλου δεδομένων εκπαίδευσης αποτελούμενο από n σημεία της μορφής:

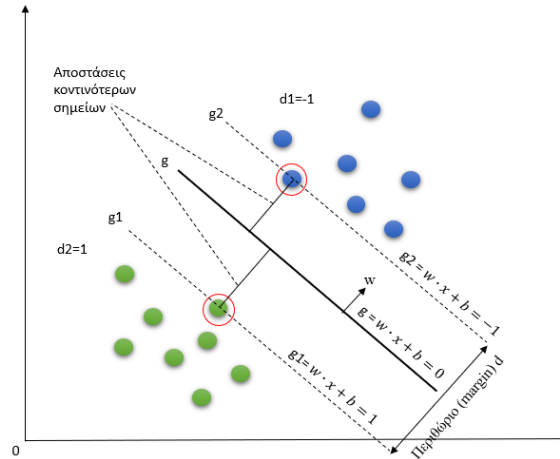
$$(x_1, y_1), \dots, (x_n, y_n)$$

Όπου $y_i, i=1, \dots, n$ είναι είτε 1 είτε -1 μαρτυρώντας την κλάση στην οποία ανήκει το σημείο x_i και $x_i \in \mathbb{R}$. Έστω, λοιπόν, σύμφωνα με την Εικόνα 4 υπάρχει ένα υπερεπίπεδο g το οποίο χωρίζει τα αρνητικά από τα θετικά διανύσματα. Τα σημεία x τα οποία ανήκουν πάνω στο υπερεπίπεδο g ικανοποιούν την παρακάτω εξίσωση:

$$w \cdot x^T + b = 0 \quad (2.1)$$

Όπου w είναι ένα κάθετο διάνυσμα στο υπερεπίπεδο g και η παράμετρος $\frac{|b|}{||w||}$ αποτελεί την κάθετη απόσταση του υπερεπιπέδου g από την αρχή των αξόνων. Σε αυτό το σημείο αξίζει να σημειωθεί πως στη γεωμετρία η απόσταση ενός σημείου (έστω σημείο $\{x_0, y_0\}$) από ευθεία (έστω μία εξίσωση ευθείας $ax+by+c=0$) υπολογίζεται σύμφωνα με τη σχέση:

$$dis = \frac{|ax_0+by_0+c|}{\sqrt{a^2+b^2}} \quad (2.2)$$



Εικόνα 4 Γραμμικός διαχωρισμός υπερεπιπέδων. Τα κυκλωμένα σημεία αποτελούν τα διανύσματα στήριξης.

Στόχος, λοιπόν, είναι να βρεθεί το μέγιστο δυνατό περιθώριο (margin) υπερεπιπέδου το οποίο χωρίζει την ομάδα των σημείων x_i για τα οποία ισχύει $y_i=1$ (πράσινα σημεία σύμφωνα με την Εικόνα 4) από την ομάδα των σημείων x_i για τα οποία ισχύει $y_i=-1$ (μπλε σημεία σύμφωνα με την Εικόνα 4). Για τον υπολογισμό του περιθωρίου είναι αναγκαίο να υπολογιστούν οι αποστάσεις των υπερεπιπέδων g, g_1, g_2 από την αρχή των αξόνων δηλαδή από το σημείο $(0,0)$. Συνεπώς θα είναι:

$$g(x, w) = wx^T + b = 0 \quad , \quad dis = \frac{|b|}{||w||} \quad (2.3)$$

$$g_1(x, w) = wx^T + b - 1 = 0 \quad , \quad dis1 = \frac{|b+1|}{||w||} \quad (2.4)$$

$$g_2(x, w) = wx^T + b + 1 = 0 \quad , \quad dis2 = \frac{|b-1|}{||w||} \quad (2.5)$$

Ο υπολογισμός του περιθωρίου d προκύπτει από τη διαφορά του $dis2-dis1$. Τελικά, το περιθώριο d θα είναι ίσο με:

$$d = dis2 - dis1 = \frac{2}{||w||} \quad (2.6)$$

Πρέπει, λοιπόν, για κάθε διάνυσμα x_i το οποίο ανήκει στην κλάση $y_i \forall i \in [1, n]$ να βρίσκεται από την εξωτερική πλευρά του περιθωρίου d . Για να εντοπίζεται η κάθε κλάση στην περιοχή της (σύμφωνα με την Εικόνα 4 πράσινα-μπλε σημεία) πρέπει να ισχύουν οι παρακάτω σχέσεις:

$$g_1(x, w) \geq 0 \Leftrightarrow wx^T + b - 1 \geq 0, \text{ αν } y_i = 1 \quad (2.7)$$

$$g_2(x, w) \leq 0 \Leftrightarrow wx^T + b + 1 \leq 0, \text{ αν } y_i = -1 \quad (2.8)$$

Αυτές οι σχέσεις μπορούν να συνδυαστούν σε μία ως εξής:

$$y_i(wx_i^T + b) - 1 \geq 0, \forall i \quad (2.9)$$

Τελικά η ταξινόμηση των διανυσμάτων σε κλάσεις είναι ένα πρόβλημα το οποίο αναζητεί το υπερεπίπεδο $g(x,w)=0$ ώστε να ικανοποιεί τους παρακάτω περιορισμούς:

$$d_{max} = \frac{2}{||w||}, \quad (2.10)$$

$$y_i(wx_i^T + b) - 1 \geq 0, \forall i \quad (2.11)$$

Αυτό είναι ένα πρόβλημα τετραγωνικού προγραμματισμού [127] το οποίο επιλύεται με τη μέθοδο Lagrange. Έπειτα από την εφαρμογή της μεθόδου όπως περιγράφεται στο [126] η μορφή των περιορισμών θα είναι ως εξής:

$$L = \frac{1}{2} ||w||^2 - \sum_{i=1}^N \lambda_i y_i (wx_i^T + b) + \sum_{i=1}^N \lambda_i \quad (2.12)$$

Όπου λ_i ένας πολλαπλασιαστής Lagrange $\forall i \in [1,n]$. Υπολογίζονται, λοιπόν, οι τιμές για τα λ, w, b . Ο υπολογισμός του διανύσματος w πραγματοποιείται συναρτήσει των σημείων εκπαίδευσης x_i και των πολλαπλασιαστών λ_i . Αν $\lambda_i > 0$ και όχι ίσο με 0 τότε το αντίστοιχο σημείο x_i συμμετέχει στον υπολογισμό του w και ονομάζεται διάνυσμα στήριξης (support vector). Αυτή η περίπτωση που περιεγράφηκε παραπάνω αφορούσε προβλήματα στα οποία οι κλάσεις ήταν διαχωρίσιμες μεταξύ τους. Σε αντίθετη περίπτωση που οι κλάσεις δεν είναι διαχωρίσιμες μεταξύ τους εισάγονται επιπλέον μεταβλητές οι οποίες ονομάζονται μεταβλητές χαλάρωσης (έστω $\xi_i \forall i \in [1,n]$ για τις οποίες πρέπει να ισχύει $\xi_i \geq 0$). Επομένως, η σχέση (2.11) μετατρέπεται ως εξής:

$$y_i(wx_i^T + b) - 1 + \xi_i \geq 0, \forall i \quad (2.13)$$

Και για να υπάρξει σφάλμα θα πρέπει $\xi_i > 1$. Το σύνολο των σφαλμάτων θα είναι $\sum_{i=1}^n \xi_i$. Το κριτήριο που πρέπει να ελαχιστοποιηθεί είναι το εξής:

$$\frac{ww^T}{2} + C \sum_{i=1}^n \xi_i \quad (2.14)$$

Όπου C μία παράμετρος κόστους που επιλέγεται από το χρήστη και καθορίζει το βαθμό εσφαλμένης ταξινόμησης των δειγμάτων εκπαίδευσης. Η συνάρτηση έπειτα από την εφαρμογή της μεθόδου Lagrange είναι πολυπλοκότερη της προηγούμενης, βέβαια η τελική μορφή της εξαρτάται μόνο από τους πολλαπλασιαστές λ_i που αντιστοιχούν στα σημεία x_i . Η διαφορά σε αυτήν την

περίπτωση είναι πως για διανύσματα στήριξης επιλέγονται εκείνα τα σημεία για τα οποία οι πολλαπλασιαστές Lagrange είναι θετικοί και μικρότεροι από την παράμετρο κόστους C .

Τα διανύσματα στήριξης, γενικά, αντιστοιχούν σε θετικούς πολλαπλασιαστές λ άρα όσο αυτοί πλησιάζουν την παράμετρο κόστους C τόσο πιο σημαντικό είναι το στοιχείο x_i και επομένως θα έχει και τη μεγαλύτερη επιρροή στο υπερεπίπεδο διαχωρισμού. Μια μεγάλη τιμή στην παράμετρο κόστους C μπορεί να καθυστερήσει την εκπαίδευση από τη στιγμή που το εύρος τιμών των πολλαπλασιαστών λ που διερευνώνται αυξάνεται.

Μέχρι τώρα οι περιπτώσεις αφορούσαν γραμμικό διαχωρισμό γραμμικών κλάσεων. Υπάρχει και μία τελευταία περίπτωση, αυτή στην οποία οι κλάσεις είναι μη γραμμικά διαχωρίσιμες. Για να επιτευχθεί, λοιπόν, ο διαχωρισμός τους χρειάζεται η κατασκευή μιας μη γραμμικής συνάρτησης διαχωρισμού. Η κατασκευή αυτής της συνάρτησης πραγματοποιείται αν αντικατασταθούν τα εσωτερικά γινόμενα, που εμφανίζονται στις σχέσεις των μηχανών διανυσμάτων στήριξης, με κατάλληλες συναρτήσεις. Οι συναρτήσεις αυτές ονομάζονται συναρτήσεις πυρήνα (kernel functions) [126]. Υπάρχουν διάφορων ειδών, κάποιες από αυτές παρατίθενται παρακάτω και συμβολίζονται ως εξής:

$$K(x, y) = x \cdot y \quad (2.15)$$

$$K(x, y) = (x \cdot y + c_0)^p \quad (2.16)$$

$$K(x, y) = e^{-\|x-y\|^2/\gamma} \quad (2.17)$$

Η πρώτη ονομάζεται γραμμική, η δεύτερη πολυωνυμική και η Τρίτη γκαουσιανή. Συνεπώς, εφόσον η συνάρτηση πυρήνα είναι σταθερή, η μόνη παράμετρος που ορίζεται από το χρήστη είναι η παράμετρος του κόστους. Μπορούν να υλοποιηθούν διαφορετικές συναρτήσεις διαχωρισμού αλλάζοντας μόνο τη συνάρτηση πυρήνα.

Γενικά, ο ταξινομητής μηχανών διανυσμάτων στήριξης έχει και πλεονεκτήματα και μειονεκτήματα. Δεν επηρεάζεται, αρχικά, εύκολα από τη διάσταση του χώρου των χαρακτηριστικών εισόδου. Βέβαια, αργεί να εκπαιδευτεί σε μεγάλα σύνολα δειγμάτων εκπαίδευσης. Επίσης, επιλύει προβλήματα ταξινόμησης μόνο δύο κλάσεων και όχι παραπάνω.

2.4 Ενίσχυση

2.4.1 Ορισμός της τεχνικής της ενίσχυσης

Γενικά, στις μεθόδους μηχανικής μάθησης χρησιμοποιείται ένα μόνο μοντέλο (single learner) για την πρόβλεψη ή ταξινόμηση όπως ακριβώς τα μοντέλα που αναλύθηκαν παραπάνω. Η ανάγκη, ωστόσο, για ανάπτυξη μοντέλων με μεγαλύτερη απόδοση γέννησε την ιδέα για τεχνικές που χρησιμοποιούν σύνολο μοντέλων (Ensemble Methods) [32]. Οι μέθοδοι αυτοί χρησιμοποιούν σύνολο ταξινομητών για να βελτιώσουν τα αποτελέσματα. Οι ταξινομητές αυτοί καλούνται αδύναμοι ταξινομητές (weak classifiers) ο συνδυασμός των οποίων παράγει τον τελικό ισχυρό ταξινομητή (strong classifier).

Η τεχνική της ενίσχυσης (Boosting) είναι μια συνδυασμού πολλαπλών ταξινομητών (ensemble) με κυριότερο στόχο τη μείωση τη μεροληψίας (bias) και της διακύμανσης (variance). Σε αυτό το σημείο αξίζει να σημειωθεί η διαφορά της τεχνικής του σακουλιάσματος (bagging) με αυτή της ενίσχυσης (boosting). Στην πρώτη οι αδύναμοι ταξινομητές εκπαιδεύονται παράλληλα ενώ στη δεύτερη διαδοχικά για να επιτύχουν το στόχο τους [33].

2.4.2 Αλγόριθμοι ενίσχυσης

Στην παρούσα εργασία αναλύονται οι τρεις παρακάτω αλγόριθμοι ενίσχυσης:

1. Προσαρμοστικής Ενίσχυσης (Adaptive Boosting-AdaBoost)
2. Ενίσχυσης Κλίσης (Gradient Boosting)
3. Εξαιρετικής Ενίσχυσης Κλίσης (eXtreme Gradient Boosting-XGBoost)

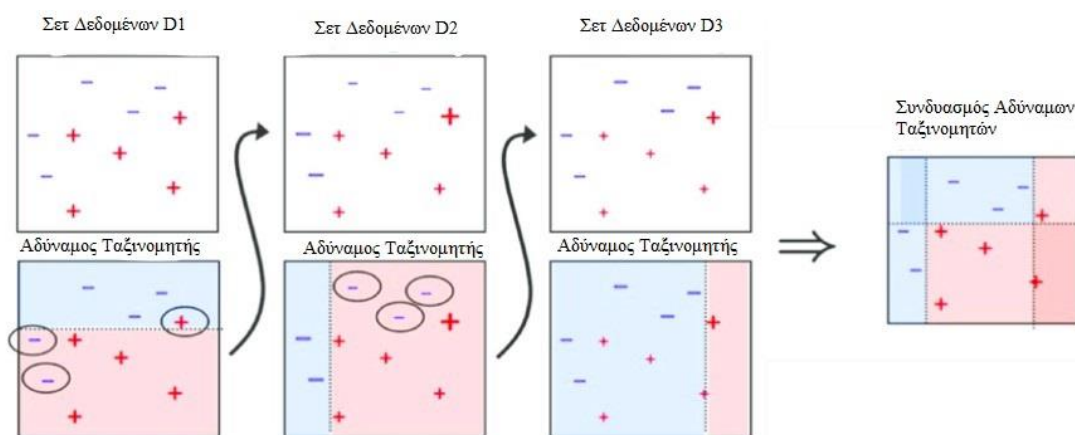
2.4.3 Προσαρμοστική Ενίσχυση

Η έννοια της ενίσχυσης (boosting) εφευρέθηκε ώστε να βελτιώσει την ακρίβεια και συνεπώς την απόδοση των ταξινομητών μηχανικής μάθησης. Ο αλγόριθμος προσαρμοστικής ενίσχυσης (Adaptive Boosting-Adaboost) προτάθηκε από τους Yoav Freund & Robert Shapire [34,35] το 1995 με σκοπό τη δημιουργία ενός ισχυρού ταξινομητή. Συγκεκριμένα, ο αλγόριθμος αυτός βασίζεται στην τεχνική της ενίσχυσης και προσπαθεί να δημιουργήσει έναν ισχυρό ταξινομητή από ένα σύνολο «αδύναμων» ταξινομητών (weak classifiers) [34]. Οι αδύναμοι αυτοί ταξινομητές που χρησιμοποιούνται μπορούν να θεωρηθούν ως δέντρα απόφασης ενός κόμβου [36].

Το σύνολο εκπαίδευσης, λοιπόν, επιλέγεται βασιζόμενο στις προηγούμενες επιδόσεις του ταξινομητή. Πιο συγκεκριμένα, ο αλγόριθμος προσαρμοστικής ενίσχυσης εκπαιδεύει τον αδύναμο ταξινομητή σε πολλούς κύκλους εκπαίδευσης χρησιμοποιώντας ίσα βάρη στο αρχικό σύνολο εκπαίδευσης. Αυτά τα βάρη προσαρμόζονται σε κάθε κύκλο εκπαίδευσης δημιουργώντας έτσι αυτούς τους αδύναμους ταξινομητές [36]. Τα βάρη των δειγμάτων εκπαίδευσης που ταξινομήθηκαν εσφαλμένα από τον τρέχον αδύναμο ταξινομητή αυξάνονται, ενώ τα βάρη εκείνων των δειγμάτων που ταξινομήθηκαν σωστά μειώνονται [36].

Ο αλγόριθμος αυτός έχει γενικά καλή απόδοση λόγω της ικανότητας του να δημιουργεί διευρυμένη ποικιλομορφία. Για να βελτιωθεί ακόμη περισσότερο η απόδοση του τελικού μοντέλου είναι σημαντικό να αποτελείται από διαφοροποιημένους αδύναμους ταξινομητές [37].

Παρακάτω φαίνεται ένα σχήμα που αναπαριστά τη λειτουργία του αλγορίθμου η οποία εξηγείται παρακάτω.



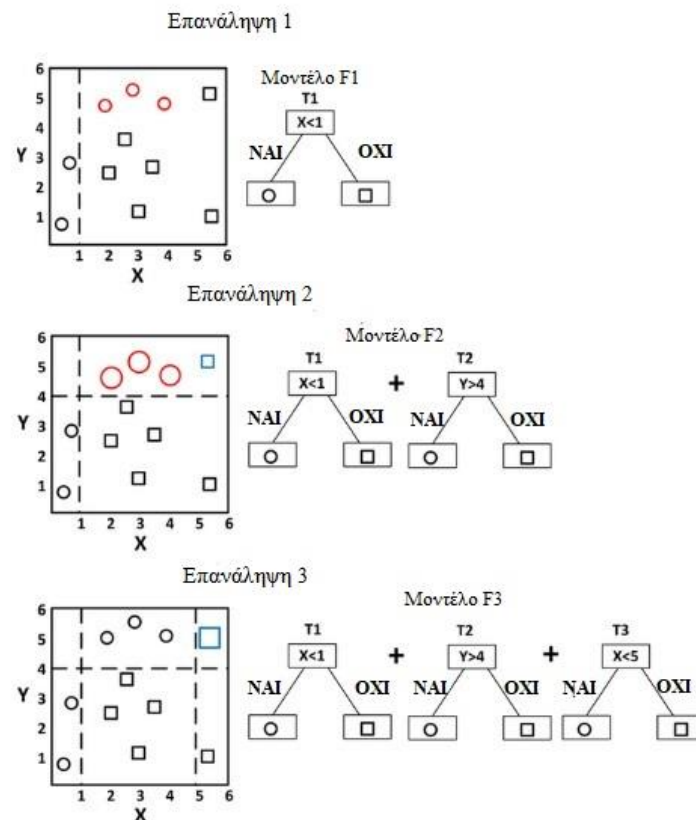
Εικόνα 5 Λειτουργία αλγορίθμου Προσαρμοστικής Ενίσχυσης (Adaboost).

Αρχικά όλα τα δείγματα από το σετ δεδομένων D1 λαμβάνουν μια τιμή βάρους που είναι ίδια για όλα (θετικά και αρνητικά). Στην πρώτη επανάληψη ο αδύναμος ταξινομητής ταξινομεί τα δείγματα σύμφωνα με την οριζόντια γραμμή (αν είναι επάνω ως αρνητικά ενώ αν είναι κάτω ως θετικά). Οπότε υπάρχει αποτυχία στην ταξινόμηση τριών δειγμάτων (1 θετικό έπρεπε να ταξινομηθεί ως αρνητικό και 2 αρνητικά έπρεπε να ταξινομηθούν ως θετικά). Σε αυτό το σημείο υπολογίζεται το σφάλμα ταξινόμησης σύμφωνα με τα βάρη των δειγμάτων. Στη συνέχεια τα βάρη των λάθος ταξινομημένων δειγμάτων αυξάνονται ενώ τα υπόλοιπα μειώνονται, όπως φαίνεται στην αναπαράσταση του σετ δεδομένων D2. Στην επόμενη επανάληψη δημιουργείται μία άλλη γραμμή

ταξινόμησης με βάση τα νέα βάρη. Και εδώ υπάρχει λάθος ταξινόμηση τριών δειγμάτων. Ωστόσο, το σφάλμα της ταξινόμησης είναι μικρότερο από αυτό της πρώτης επανάληψης καθώς το βάρος αυτών των τριών δειγμάτων έχει μειωθεί. Ανανεώνονται και πάλι τα βάρη και ακολουθεί η Τρίτη επανάληψη. Σε αυτήν την επανάληψη έχουν ταξινομηθεί σωστά όλα τα δείγματα. Ο τελικός ισχυρός ταξινομητής, επομένως, προκύπτει από το συνδυασμό των τριών αδύναμων ταξινομητών που επιλέχθηκαν και ένα κατώφλι. Τα αποτελέσματα των τριών αδύναμων ταξινομητών αθροίζονται, και αν το άθροισμα ξεπερνά το κατώφλι του ταξινομητή, το υπό εξέταση δείγμα ταξινομείται ως θετικό, αλλιώς ως αρνητικό. Παρόλο που κάθε ταξινομητής ξεχωριστά δεν καταφέρνει να ταξινομήσει σωστά όλα τα δείγματα του, ο τελικός ισχυρός ταξινομητής καταφέρνει και ταξινομεί σωστά όλα τα δείγματα όπως φαίνεται και από την Εικόνα 5.

2.4.4 Ενίσχυση Κλίσης

Η μέθοδος της ενίσχυσης κλίσης (Gradient Boosting) ανήκει στις τεχνικές ενίσχυσης (boosting) και χρησιμοποιείται για την ανάπτυξη μοντέλων ταξινόμησης και παλινδρόμησης [38,39]. Ως μέθοδος ενίσχυσης αποτελείται από αδύναμους ταξινομητές όπως ακριβώς και ο αλγόριθμος προσαρμοστικής ενίσχυσης. Αυτοί οι αδύναμοι ταξινομητές είναι μη γραμμικά μοντέλα γνωστά ως δέντρα απόφασης. Αυτός ο αλγόριθμος λειτουργεί εκπαιδεύοντας πολλά μοντέλα με βαθμιαίο, προσθετικό και διαδοχικό τρόπο. Η μεγάλη διαφορά του αλγορίθμου ενίσχυσης κλίσης και του αλγορίθμου προσαρμοστικής ενίσχυσης έγκειται στον τρόπο με τον οποίο οι δύο αλγόριθμοι εντοπίζουν τις ατέλειες των αδύναμων ταξινομητών. Το μοντέλο προσαρμοστικής ενίσχυσης, όπως προαναφέρθηκε, αναγνωρίζει τις αδυναμίες χρησιμοποιώντας υψηλού βάρους δεδομένα, ενώ το μοντέλο ενίσχυσης κλίσης κάνει το ίδιο τοποθετώντας κλίσεις (gradients) στη συνάρτηση απώλειας (loss function) [40]. Ο αλγόριθμος ενίσχυσης κλίσης λειτουργεί επαναληπτικά έως ότου δεν υπάρχει κάποιο πρότυπο για να διαμορφωθεί. Σκοπός είναι να ελαχιστοποιηθεί όσο το δυνατόν περισσότερο η συνάρτηση απώλειας. Σε αυτό το σημείο τερματίζει και ο αλγόριθμος όπως φαίνεται και παρακάτω από την Εικόνα 6.



Εικόνα 6 Λειτουργία αλγορίθμου ενίσχυσης κλίσης.

2.4.5 Εξαιρετική Ενίσχυση Κλίσης

Ο αλγόριθμος εξαιρετικής ενίσχυσης κλίσης (eXtreme Gradient Boosting-XGBoost) βασίζεται κι αυτός στα δέντρα απόφασης και χρησιμοποιεί την τεχνική της ενίσχυσης κλίσης η οποία αναλύθηκε παραπάνω για τη βελτίωση της επίδοσης του μοντέλου. Παρουσιάζει αρκετές ομοιότητες με τους δύο παραπάνω αλγορίθμους που βασίζονται στην τεχνική της ενίσχυσης. Δημιουργούνται αδύναμοι ταξινομητές σε κάθε επανάληψη εκπαίδευσης με στόχο τον τελικό ισχυρό ταξινομητή βάσει αυτών.

Σε κάθε επανάληψη γίνονται προβλέψεις σε όλο το σύνολο εκπαίδευσης και κατόπιν υπολογίζεται το σφάλμα της ταξινόμησης. Ο αδύναμος ταξινομητής παράγει συνήθως ανεπαρκείς προβλέψεις οπότε είναι αναγκαίο να εισαχθεί σε επιπλέον επαναλήψεις. Έτσι, ο αδύναμος ταξινομητής της επόμενης επανάληψης λαμβάνει υπόψιν το σφάλμα της προηγούμενης και επιχειρεί να το εξαλείψει

[41]. Το τελικό αποτέλεσμα του ισχυρού ταξινομητή παράγεται από το συνδυασμό των αδύναμων ταξινομητών όπως εξηγήθηκε προηγουμένως.

Η διαφορά της εξαιρετικής ενίσχυσης κλίσης με την προσαρμοστική ενίσχυση έγκειται στο γεγονός ότι στο συγκεκριμένο αλγόριθμο αντί να αλλαχθούν τα βάρη των δειγμάτων που ταξινομηθήκαν εσφαλμένα, γίνεται εκπαίδευση κάθε νέου μοντέλου αξιοποιώντας τα σφάλματα του προηγούμενου. Διαφέρει, επίσης, από τους υπόλοιπους ταξινομητές καθώς συρρικνώνει αναλογικά τα φύλα των δέντρων και τυχαιοποιεί επιπλέον τις παραμέτρους του [41].

2.5 Μεθοδολογίες Ερμηνευσιμότητας

Οι μεθοδολογίες ερμηνευσιμότητας (interpretability methods) υποβοηθούν την επεξήγηση του εξερχόμενου αποτελέσματος από όλα τα παραπάνω μοντέλα. Όλα αυτά τα μοντέλα μηχανικής μάθησης που αναλύθηκαν παραπάνω από μόνα τους δεν αποτελούν ερμηνεύσιμα μοντέλα καθώς το αποτέλεσμα της πρόβλεψης δεν είναι ερμηνεύσιμο και κατανοητό από την ανθρώπινη λογική [42]. Βέβαια υπάρχουν τεχνικές ερμηνευσιμότητας και για αυτά τα μοντέλα με σκοπό να απαντήσουν στην ερώτηση «γιατί έχει εξαχθεί το συγκεκριμένο αποτέλεσμα;». Οι μέθοδοι αυτές έρχονται να δώσουν και μία ερμηνεία που κρύβεται πίσω από το «μαύρο κουτί» (black box) των μοντέλων μηχανικής μάθησης. Βοηθούν λοιπόν τους ανθρώπους να καταλάβουν πλήρως αλλά και να εμπιστευτούν την αναδυόμενη γενιά της τεχνητής νοημοσύνης. Η ανάπτυξή τους κρίθηκε αναγκαία τα τελευταία χρόνια καθώς υπήρξαν διάφορες αποτυχίες ερευνών βασισμένων σε περίπλοκα μοντέλα πρόβλεψης, σε διάφορους τομείς, εξαιτίας της έλλειψης αυτών των μεθόδων ερμηνευσιμότητας [43]. Έχει αποδειχθεί ότι αυτά τα ερμηνεύσιμα μοντέλα πετυχαίνουν υψηλά ποσοστά ακρίβειας [44] συγκριτικά και χρησιμοποιούνται ιδιαίτερα στον τομέα της υγείας καθώς δεν απαιτούν από τους γιατρούς και τους ειδικούς να έχουν κάποια τεχνική ή στατιστική γνώση [45].

Παρακάτω λοιπόν παρουσιάζονται κάποιες μέθοδοι ερμηνευσιμότητας που χρησιμοποιήθηκαν και στην παρούσα πτυχιακή εργασία.

1) Διαγράμματα Μερικής Εξάρτησης (Partial Dependence Plots-PDP)

Τα διαγράμματα αυτά δείχνουν την οριακή επίδραση ενός ή περισσότερων χαρακτηριστικών στο τελικό αποτέλεσμα της πρόβλεψης [46]. Ένα τέτοιο διάγραμμα μπορεί να δώσει πληροφορίες για το αν η σχέση μεταξύ ενός ή περισσότερων χαρακτηριστικών με την κλάση είναι γραμμική, μονότονη ή πιο περίπλοκη. Για παράδειγμα όταν αυτό το διάγραμμα υλοποιείται για ένα μοντέλο γραμμικής παλινδρόμησης, πάντα η σχέση θα φαίνεται ότι είναι γραμμική.

2) Τιμή Shapley (Shapley Value-SHAP)

Η μέθοδος αυτή αναδεικνύει τη σημαντικότητα του κάθε χαρακτηριστικού ως προς την τελική πρόβλεψη [47]. Η ιδέα της μεθόδου αυτή προέρχεται από τη θεωρία των παιγνίων [48]. Σε αυτήν την ιδέα η πρόβλεψη μπορεί να επεξηγηθεί υποθέτοντας ότι κάθε τιμή χαρακτηριστικού αποτελεί έναν «παίκτη» ενός παιχνιδιού. Η συνεισφορά κάθε παίκτη μετριέται προσθέτοντας και αφαιρώντας τον παίκτη από όλα τα υποσύνολα των υπολοίπων παικτών. Η τιμή SHapley για έναν παίκτη είναι το σταθμισμένο άθροισμα όλων των συνεισφορών του. Η τιμή αυτή είναι προσθετική και τοπικά ακριβής. Εάν προστεθούν οι τιμές SHapley όλων των χαρακτηριστικών, συν τη βασική τιμή, που είναι ο μέσος όρος πρόβλεψης, θα ληφθεί η ακριβής τιμή πρόβλεψης. Αυτό είναι ένα χαρακτηριστικό που πολλές άλλες μέθοδοι δεν έχουν. Παρακάτω παρουσιάζεται ένα παράδειγμα της μεθόδου αυτής από πείραμα της παρούσας εργασίας.



Εικόνα 7 Τοπική ερμηνευσιμότητα χαρακτηριστικών με τη μέθοδο SHAP.

Η παραπάνω Εικόνα 7 δείχνει τη τιμή Shapley για κάθε χαρακτηριστικό που συμμετέχει στην πρόβλεψη. Το ροζ χρώμα δείχνει θετική συνεισφορά στο αποτέλεσμα ενώ το μπλε χρώμα αρνητική.

3) Τοπικά Ερμηνεύσιμα Μοντέλα Αγνωστικιστικών Εξηγήσεων (Local Interpretable Model-agnostic Explanations- LIME)

Οι μέθοδοι τοπικών ερμηνεύσιμων μοντέλων διαφέρουν από αυτές των ολικών. Συγκεκριμένα, οι τοπικές μέθοδοι δεν προσπαθούν να επεξηγήσουν ολόκληρο το μοντέλο αλλά εκπαιδεύει ερμηνεύσιμα μοντέλα για την προσέγγιση των μεμονωμένων προβλέψεων [49]. Το LIME προσπαθεί να καταλάβει πώς αλλάζουν οι προβλέψεις όταν ανακατανέμονται τα δείγματα δεδομένων. Κάθε χαρακτηριστικό λαμβάνει ένας βάρος είτε θετικό είτε αρνητικό. Τελικά τα χαρακτηριστικά με τα μεγαλύτερα θετικά βάρη παρουσιάζονται ως επεξήγηση.

Παραπάνω αναφέρθηκαν μεθοδολογίες ερμηνευσιμότητας για μοντέλα μηχανικής μάθησης τα οποία δεν είναι ερμηνεύσιμα από μόνα τους. Ωστόσο υπάρχουν και ταξινομητές οι οποίοι είναι από μόνοι τους ερμηνεύσιμοι, ακολουθούν, δηλαδή, την ανθρώπινη λογική ώστε να εξαχθεί η τελική πρόβλεψη. Αυτοί οι ταξινομητές ονομάζονται ασαφείς και χρησιμοποιούν κανόνες κατανοητούς από τους ανθρώπους με σκοπό να κάνουν μια ακριβή πρόβλεψη. Στην παρούσα πτυχιακή εργασία αναπτύχθηκε ένας τέτοιος ταξινομητής. Παρακάτω, λοιπόν, αναλύεται η θεωρητική βάση των ασαφών ταξινομητών καθώς και ο τρόπος με τον οποίο προβλέπουν ένα αποτέλεσμα μέσα από λογικούς κανόνες.

Κεφάλαιο 3

Ασαφείς Ταξινομητές

3.1 Ασαφής Λογική (Fuzzy Logic)

Η ασαφής λογική πηγάζει από τη θεωρία των ασαφών συνόλων, η οποία εφευρέθηκε το 1965 από το Lofti Zadeh [50] και εκ τότε φαίνεται να έχει μεγάλη απήχηση στην επιστημονική κοινότητα σε διάφορες εφαρμογές. Η θεωρία της ασαφούς λογικής έρχεται να καλύψει την αβεβαιότητα και την ανακρίβεια που κυριαρχεί στην ανθρώπινη πραγματικότητα. Έχει τη δυνατότητα να μιμηθεί την ανθρώπινη κρίση με μια μορφή κανόνων που προσεγγίζουν τη συλλογιστική πορεία της ανθρώπινης σκέψης. Η κλασσική θεωρία των συνόλων βασίζεται στην αρχή ότι αν ένα στοιχείο ανήκει σε ένα σύνολο δεν ανήκει σε άλλο σύνολο [51]. Με λίγα λόγια η απάντηση στο ερώτημα αν ένα στοιχείο ανήκει σε ένα σύνολο μπορεί να είναι είτε αληθής είτε ψευδής. Η θεωρία των ασαφών συνόλων αποτελεί επέκταση της κλασσικής θεωρίας συνόλων αφού δέχεται και το γεγονός ένα στοιχείο και να ανήκει σε ένα σύνολο αλλά και να μην ανήκει, ανάλογα με το βαθμό συμμετοχής του σε αυτό [52].

3.2 Ασαφή Σύνολα (Fuzzy Sets)

Τα ασαφή σύνολα όπως προαναφέρθηκε αποτελούν μια επέκταση της θεωρίας των συνόλων. Εξετάζουν το βαθμό συμμετοχής μίας μεταβλητής στο κάθε σύνολο. Η έννοια της μεταβλητής αυτής στην ασαφή λογική μεταφράζεται ως λεκτική μεταβλητή καθώς εισάγονται λεκτικοί όροι (linguistic terms) οι οποίοι την περιγράφουν. Κάθε λεκτικός όρος περιγράφεται από ένα ασαφές σύνολο με συγκεκριμένο νόημα όπως για παράδειγμα «χαμηλός», «μέτριος», «υψηλός».

Μαθηματικά, αυτός ο βαθμός συμμετοχής κάθε μεταβλητής στο ασαφές σύνολο αλλά και το ίδιο το σύνολο καθορίζεται από τις συναρτήσεις συμμετοχής (membership functions) [53]. Έστω, λοιπόν, ένα ασαφές σύνολο F . Αυτό το σύνολο έχει μία συνάρτηση συμμετοχής μ_F ορισμένη στο X για την οποία ισχύει: $\mu_F : X \rightarrow [0,1]$. Υπάρχουν διάφορων ειδών συναρτήσεις συμμετοχής παρακάτω αναφέρονται αυτές που χρησιμοποιούνται συχνότερα στη βιβλιογραφία.

- a) Τριγωνική συνάρτηση συμμετοχής (Triangular Membership Function):** Η τριγωνική συνάρτηση συμμετοχής ορίζεται από τρία σημεία (a,b,c) και περιγράφεται από την παρακάτω κλαδωτή μαθηματική συνάρτηση. Τα τρία αυτά σημεία ορίζονται από τον χρήστη ανάλογα με την τοποθεσία που επιθυμεί πάνω στον άξονα x να έχει η τριγωνική συνάρτηση.

$$\mu_{tri}(x) = \begin{cases} \frac{x-a}{b-a}, & \text{αν } a < x \leq b \\ \frac{c-x}{c-b}, & \text{αν } b < x < c \\ 0, & \text{διαφορετικά} \end{cases} \quad (3.1)$$

- b) Τραπεζοειδής Συνάρτηση Συμμετοχής (Trapezoid Membership Function):** Η τραπεζοειδής συνάρτηση συμμετοχής περιγράφεται από τέσσερις παραμέτρους και μαθηματικά ορίζεται από την παρακάτω κλαδωτή συνάρτηση.

$$\mu_{trap}(x) = \begin{cases} \frac{x-a}{b-a}, & \text{αν } a < x < b \\ 1, & \text{αν } b \leq x \leq c \\ \frac{d-x}{d-c}, & \text{αν } c < x < d \\ 0, & \text{διαφορετικά} \end{cases} \quad (3.2)$$

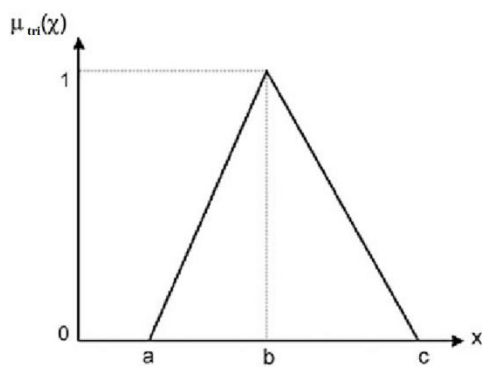
- c) Γκαουσιανή Συνάρτηση Συμμετοχής (Gaussian Membership Function):** Η γκαουσιανή συνάρτηση συμμετοχής περιγράφεται από το κέντρο m και το εύρος σ και μαθηματικά ορίζεται από την παρακάτω συνάρτηση. Συγκεκριμένα, όπου m είναι η μέση τιμή και όπου σ η τυπική απόκλιση της γκαουσιανής συνάρτησης

$$\mu_{gauss}(x) = \exp \left\{ - \frac{(x-m)^2}{\sigma^2} \right\} \quad (3.3)$$

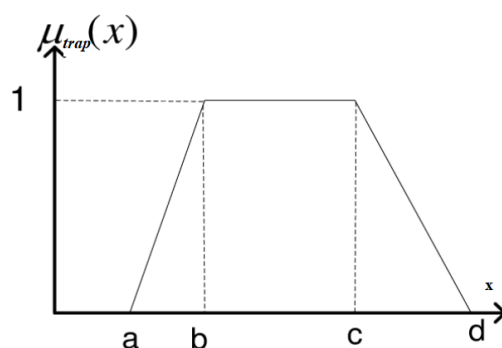
- d) Δίπλευρες Γκαουσιανές Συναρτήσεις Συμμετοχής:** Οι δίπλευρες Γκαουσιανές συναρτήσεις συμμετοχής αποτελούν δίπλευρη επέκταση των γκαουσιανών συναρτήσεων που αναφέρθηκαν παραπάνω. Μαθηματικά ορίζονται από την παρακάτω σχέση:

$$\mu_{gauss2}(x) = \begin{cases} \exp \left\{ - \frac{(x-ml)^2}{\sigma_l^2} \right\}, & \text{αν } x < ml \\ 1, & \text{αν } ml \leq x \leq mr \\ \exp \left\{ - \frac{(x-mr)^2}{\sigma_r^2} \right\}, & \text{αν } x > mr \end{cases} \quad (3.4)$$

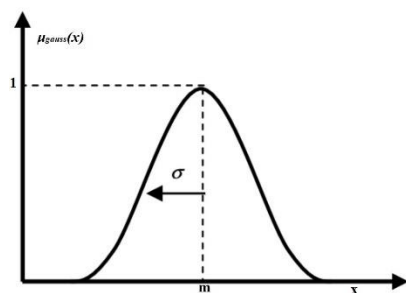
Οι παραπάνω τέσσερις συναρτήσεις συμμετοχής αποτελούν τις πιο συχνές μορφές συναρτήσεων που χρησιμοποιούνται σε μεθόδους ασαφούς λογικής. Παρακάτω φαίνονται και σχηματικά οι συναρτήσεις που παραπάνω εξηγήθηκε ο μαθηματικός τους τύπος. Όλες οι μεταβλητές στην μαθηματική αναπαράσταση αντιστοιχίζονται και στις σχηματικές αναπαραστάσεις.



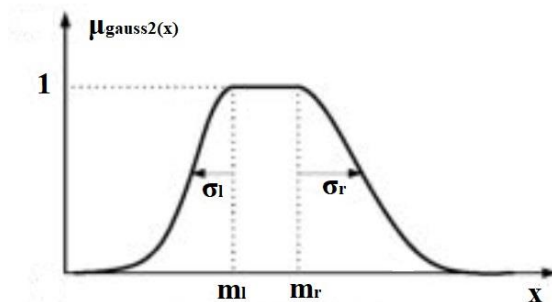
A) Τριγωνική Συνάρτηση Συμμετοχής



B) Τραπεζοειδής Συνάρτηση Συμμετοχής



Γ) Γκαουσιανή Συνάρτηση Συμμετοχής



Δ) Δίπλευρη Γκαουσιανή Συνάρτηση Συμμετοχής

Εικόνα 8 Σχηματική αναπαράσταση Συναρτήσεων συμμετοχής

Για την παρούσα μελέτη χρησιμοποιήθηκαν τριγωνικές αλλά και τραπεζοειδείς συναρτήσεις συμμετοχής συνδυαστικά με σκοπό την εξαγωγή των ασαφών συνόλων.

3.3 Ασαφείς Κανόνες Ταξινόμησης

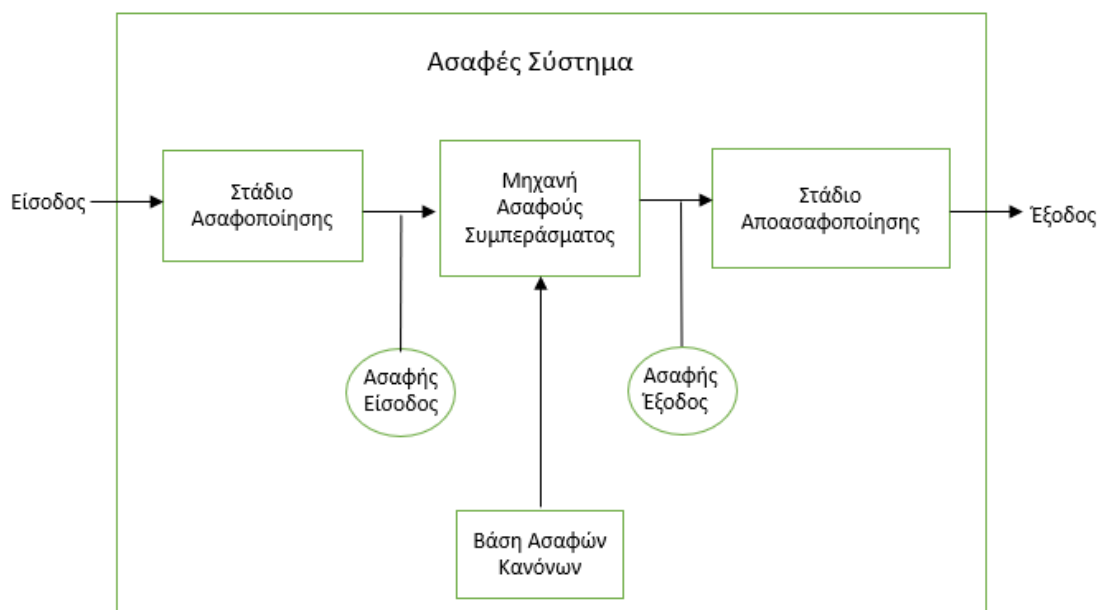
Βασικό δομικό στοιχείο των ασαφών ταξινομητών αποτελούν οι ασαφείς κανόνες (fuzzy rules). Είναι απλές υποθέσεις της μορφής (IF/THEN) με στόχο το συνδυασμό του χώρου των χαρακτηριστικών εισόδου με αυτό των χαρακτηριστικών εξόδου (κλάσεων) ακολουθώντας την ανθρώπινη λογική. Έστω, λοιπόν, ένα πρόβλημα με n χαρακτηριστικά τα διανύσματα των οποίων ορίζονται στο χώρο $X = \{x_1, x_2, \dots, x_n\} \in \mathbb{R}^n$ και ένα σύνολο m κλάσεων $C = \{C_1, C_2, \dots, C_m\}$. Ένας ασαφής ταξινομητής χρησιμοποιεί κανόνες ταξινόμησης που είναι της παρακάτω μορφής:

$$IF X_1 \text{ is } A_1 \text{ AND } X_2 \text{ is } A_2 \text{ AND } \dots \text{ AND } X_n \text{ is } A_n \text{ THEN class is } C_j \text{ [54]}$$

,όπου $A_i, i=1,2,\dots,n$ είναι ασαφή σύνολα ορισμένα στους αντίστοιχους χώρους χαρακτηριστικών και κλάσεων. Το τμήμα IF αποτελεί το τμήμα της υπόθεσης ενώ το υπόλοιπο το τμήμα του συμπεράσματος. Το μέρος της υπόθεσης συνδυάζει τα επιμέρους ασαφή σύνολα σε κάθε διάσταση με τον τελεστή της τομής AND. Μπορεί αυτοί οι κανόνες να είναι πολλαπλοί για την περιγραφή μιας κλάσης. Οι κανόνες αυτοί επίσης δεν απαιτούν στο τμήμα της υπόθεσης να υπάρχουν όλα τα διανύσματα των χαρακτηριστικών. Για παράδειγμα για ένα πρόβλημα δύο διανυσμάτων και δύο κλάσεων σύμφωνα με τον παραπάνω κανόνα μπορεί να δημιουργηθεί ασαφής κανόνας και χωρίς τη μεταβλητή X_2 και θα είναι: «**IF X_1 is A_1** » [55].

3.4 Ασαφή Συστήματα

Τα πρώτα ασαφή μοντέλα που αναπτύχθηκαν στις δεκαετίες του 80 και του 90 δεν ήταν ικανά να βρουν λύση ακόμη και σε απλά ζητήματα. Για αυτό το λόγο γεννήθηκε η ανάγκη να αναπτυχθούν συστήματα στήριξης αποφάσεων τα οποία με αυτοματοποιημένες διαδικασίες θα έπαιρναν μία απόφαση και θα έδιναν τη σωστή λύση. Στόχος ήταν να επεξεργαστούν πολλαπλοί κανόνες. Τα ασαφή συστήματα δημιουργούν και επεξεργάζονται ασαφείς κανόνες ώστε να εξάγουν ένα συμπέρασμα. Το ασαφές συμπέρασμα είναι μια διαδικασία απόκτησης νέας γνώσης μέσω της υπάρχουσας γνώσης με χρήση ασαφούς λογικής. Ασαφή αποκαλούνται αυτά τα συστήματα απόφασης και ελέγχου που μοντελοποιούν ασαφείς μεταβλητές. Παρακάτω λοιπόν παρουσιάζονται τα στάδια που εκτελούνται σε ένα ασαφές σύστημα όπως φαίνονται και στην Εικόνα 9:



Εικόνα 9 Δομή ενός ασαφούς συστήματος

- 1) **Στάδιο Ασαφοποίησης:** Το πρώτο βήμα αφορά την ασαφοποίηση των χαρακτηριστικών εισόδου. Η έννοια της ασαφοποίησης ορίζεται ως τη μετατροπή των χαρακτηριστικών εισόδου σε ασαφείς τιμές [60]. Εφόσον λοιπόν δημιουργηθούν τα ασαφή σύνολα στο στάδιο αυτό καθορίζεται ο βαθμός συμμετοχής κάθε χαρακτηριστικού εισόδου σε καθένα από αυτά τα σύνολα. Πρακτικά σε αυτό το βήμα παράγονται οι συναρτήσεις συμμετοχής παραδείγματα των οποίων αναλύθηκαν παραπάνω για κάθε ασαφές σύνολο. Μέσω των συναρτήσεων συμμετοχής, λοιπόν, οι εισοδοί που είναι αριθμητικές τιμές στο όριο του χώρου αναφοράς μετατρέπονται σε εξόδους στο προσδιορισμένο ασαφές σύνολο. Οι αριθμητικές τιμές τελικά, μετατρέπονται σε λεκτικές όπως για παράδειγμα «χαμηλός/ή/ό», «μεσαίος/α/ο», «υψηλός/ή/ό» ώστε στη συνέχεια να δημιουργηθούν οι ασαφείς κανόνες από τις δημιουργηθείσες ασαφείς εισόδους.
- 2) **Στάδιο Εκτίμησης Κανόνων:** Σε αυτό το στάδιο δημιουργούνται οι ασαφείς κανόνες ώστε να χτιστεί η γνωσιακή βάση κανόνων. Για να σχηματιστεί η μορφή κανόνων όπως παρουσιάστηκε παραπάνω θα πρέπει να εφαρμοστούν κάποιοι τελεστές όπως για παράδειγμα AND, OR μεταξύ των υποθέσεων σε περίπτωση που υπάρχουν παραπάνω από μία. Έτσι λοιπόν δημιουργείται η βάση κανόνων η οποία παίρνει σαν είσοδο τις λεκτικές μεταβλητές έπειτα από το στάδιο της ασαφοποίησης και δίνει σαν έξοδο ένα βαθμό αληθείας.
- 3) **Στάδιο Συμπερασμάτων:** Σε αυτό το στάδιο εφαρμόζεται η έννοια της συνεπαγωγής. Προϋπόθεση αποτελεί η γνώση του βαθμού (degree) κάθε κανόνα (τιμές στο εύρος 0-1).

Παίρνονται λοιπόν σαν είσοδοι οι ασαφείς τιμές που έχουν δημιουργηθεί από το στάδιο της ασαφοποίησης και η γνωσιακή βάση κανόνων. Σύμφωνα με αυτή είναι εφικτή η επεξεργασία των κανόνων σύμφωνα με το βαθμό συμμετοχής των εισόδων. Με αυτόν τον τρόπο εξάγονται οι νέες ασαφείς έξοδοι για κάθε κανόνα.

- 4) **Στάδιο Συνάθροισης Κανόνων:** Στην συνέχεια συγκεντρώνονται οι έξοδοι όλων των κανόνων. Η συνάθροιση είναι μια διαδικασία κατά την οποία όλα τα ασαφή σύνολα που προκύπτουν ως έξοδοι των κανόνων συνδυάζονται ώστε να προκύψει μόνο ένα ασαφές σύνολο.
- 5) **Στάδιο Αποασαφοποίησης:** Το τελικό αυτό στάδιο εφαρμόζεται μόνο στα ασαφή μοντέλα Mamdani. Η αποασαφοποίηση (defuzzification) αποτελεί μια μαθηματική διαδικασία κατά την οποία ένα ασαφές σύνολο μετατρέπεται σε πραγματικό αριθμό [61].

3.4.1 Ασαφείς Ταξινομητές – Takagi-Sugeno

Οι ασαφείς ταξινομητές Takagi-Sugeno αποτελούν από τα πιο γνωστά ασαφή συστήματα μαζί με τα Mamdani. Οι ταξινομητές Takagi-Sugeno, όπως ονομάζονται, έχουν προταθεί από τους Takagi, Sugeno, και Kang σε μία προσπάθεια να αναπτύξουν μια συστηματική μεθοδολογία για γενίκευση των ασαφών κανόνων, από ένα δοθέν σετ δεδομένων, με προκαθορισμένες εισόδους και εξόδους χαρακτηριστικών [55]. Ακολουθούν μια συγκεκριμένη μέθοδο έως ότου καταλήξουν στο τελικό συμπέρασμα. Λόγω της υψηλής ερμηνευσιμότητας και ακρίβειας που έχουν σαν ταξινομητές έχει βρεθεί ότι πετυχαίνουν υψηλά ποσοστά σε διάφορες πρακτικές εφαρμογές όπως για παράδειγμα στις οικονομικές προβλέψεις, στην ανάλυση εικόνας, σε ιατρική διάγνωση, αναγνώριση ανθρώπινης δραστηριότητας καθώς και ανίχνευση και διάγνωση σφαλμάτων [56],[57],[58].

Σε έναν ταξινομητή Takagi-Sugeno ένας τυπικός ασαφής κανόνας είναι της ακόλουθης γενικής μορφής:

$$\text{IF } x_1 \text{ is } A_1^k \wedge x_2 \text{ is } A_2^k \wedge \dots \wedge x_d \text{ is } A_d^k \text{ THEN } f^k(x) = p_0^k + p_1^k x_1 + \dots + p_d^k x_d, \quad k=1,2,\dots,K$$

Όπου A_i^k είναι ένα ασαφές σύνολο το οποίο περιγράφεται από την i -οστή είσοδο x_i για τον k -οστό ασαφή κανόνα και τα p_0 αποτελούν σταθερές του πολωνύμου f . Όπου K το σύνολο των ασαφών κανόνων και d η διάσταση του χώρου χαρακτηριστικών εισόδου. Σε αυτό το σημείο αξίζει να σημειωθεί πως στους συγκεκριμένους ασαφείς ταξινομητές οι είσοδοι x_i αποτελούν όλα τα χαρακτηριστικά του αρχικού συνόλου δεδομένων που ορίζονται στο χώρο χαρακτηριστικών d . Το

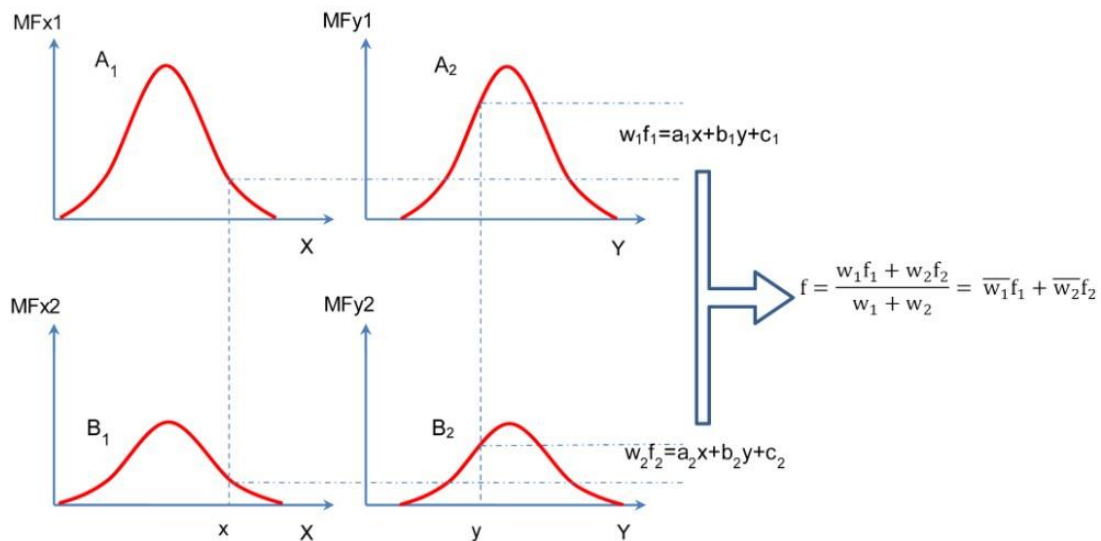
σύμβολο $\wedge$ αποτελεί το σύμβολο της τομής των υποθέσεων και y είναι η συνάρτηση εξόδου συναρτήσει των εισόδων. Αυτή μπορεί να είναι οποιαδήποτε πολυωνυμική συνάρτηση των εισόδων. Όταν αυτή η συνάρτηση είναι πολυωνυμική πρώτου βαθμού τότε και ο ταξινομητής αποκαλείται πρώτου βαθμού Takagi-Sugeno (first-order Takagi-Sugeno) όπως ακριβώς προτάθηκε και από τους δημιουργούς του [55], [59]. Ενώ όταν η συνάρτηση αυτή είναι μία σταθερά τότε ο ταξινομητής αποκαλείται μηδενικού βαθμού Takagi-Sugeno (zero-order Takagi-Sugeno) ο οποίος αποτελεί είτε μία ειδική περίπτωση ασαφούς μοντέλου Mamdani είτε μία ειδική περίπτωση ασαφούς μοντέλου Tsukamoto. Στην περίπτωση αυτή η μορφή του κανόνα μετατρέπεται ως εξής:

$$R_k: \text{ IF } x_1 \text{ is } A_1^k \wedge x_2 \text{ is } A_2^k \wedge \dots \wedge x_d \text{ is } A_d^k \text{ THEN } f^k(x) = p_0^k, k=1,2,\dots,K [59]$$

Η διαφορά τους, λοιπόν, έγκειται στη μορφή που θα έχει η πολυωνυμική συνάρτηση των εισόδων του μοντέλου. Τα ασαφή μοντέλα Takagi-Sugeno ακολουθούν τα στάδια όπως αναλύθηκαν σε προηγούμενο κεφάλαιο εκτός από το τελικό στάδιο της ασαφοποίησης. Η τελική έξοδος από έναν ταξινομητή Takagi-Sugeno είτε πρώτου είτε μηδενικού βαθμού είναι ο σταθμισμένος μέσος όρος όλων των εξόδων των ασαφών κανόνων και υπολογίζεται ως εξής:

$$\text{Τελική Έξοδος} = \frac{\sum_{i=1}^K w_i \cdot f_i}{\sum_{i=1}^K w_i} \quad (3.5)$$

Όπου w αποτελούν τα βάρη (firing strength) του κάθε ασαφούς κανόνα. Παρακάτω λοιπόν στην Εικόνα 10 παρουσιάζεται ένα παράδειγμα δύο εισόδων x, y και δύο ασαφών κανόνων. Όπως φαίνεται και από την εικόνα και όπως περιγράφηκε παραπάνω η τελική έξοδος αποτελεί ο σταθμισμένος μέσος όρος των εξόδων των ασαφών κανόνων. Στη συγκεκριμένη περίπτωση του f_1, f_2 εφόσον αποκτήθηκαν από τις συναρτήσεις συμμετοχής της κάθε εισόδου.

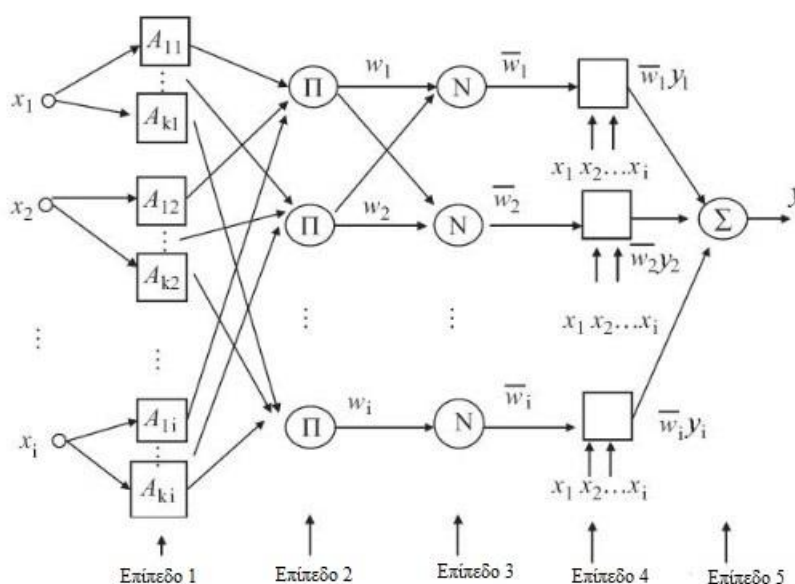


Εικόνα 10 Σχηματική αναπαράσταση πρώτου βαθμού Takagi-Sugeno. [101]

3.4.2 Προσαρμοστικό Νευροασαφές Σύστημα Συμπερασμάτων βασισμένο σε Αρχιτεκτονική Takagi – Sugeno

Ένα προσαρμοστικό δίκτυο είναι ένα πολλαπλών επιπέδων δίκτυο τροφοδοσίας (feed forward) στο οποίο κάθε κόμβος εκτελεί μια συγκεκριμένη συνάρτηση κόμβου με βάση τα εισερχόμενα σήματα, καθώς και το σύνολο των παραμέτρων που σχετίζονται με αυτόν. Δηλαδή, κάθε συνάρτηση κόμβου εξαρτάται από το σύνολο των παραμέτρων που σχετίζονται με αυτόν. Στη διαδικασία εκπαίδευσης αυτές οι παράμετροι διορθώνονται με στόχο τη συνολική ελαχιστοποίηση των σφαλμάτων στην έξοδο του δικτύου [102].

Το ANFIS είναι ένα υβριδικό νευροασαφές σύστημα. Σε αυτό το σύστημα οι ασαφείς κανόνες αναπαρίστανται σε δομή δικτύου όπως φαίνεται και παρακάτω από την Εικόνα 11. Οι τεχνικές μάθησης των νευρωνικών δικτύων μπορούν να εφαρμοστούν με σκοπό τη ρύθμιση των παραμέτρων των ασαφών μοντέλων [103]. Έτσι, εάν μια άγνωστη είσοδος έχει αναπαρασταθεί στο ANFIS, μετά την υβριδική διαδικασία μάθησης, μπορεί να προβλέψει με επιτυχία την έξοδο για τη δεδομένη είσοδο.



Εικόνα 11 Αρχιτεκτονική Takagi-Sugeno ANFIS.

Γενικά, οι ταξινομητές αυτού του τύπου μπορούν να εκπαιδευτούν ως νευρωνικά δίκτυα πέντε επιπέδων όπως ακριβώς φαίνεται στη σχηματική αναπαράσταση της Εικόνας 11. Περιγράφονται, λοιπόν, παρακάτω τα πέντε αυτά επίπεδα το καθένα ξεχωριστά.

- 1) Επίπεδο 1:** Στο επίπεδο 1 πραγματοποιείται η διαδικασία της ασαφοποίησης. Η διαδικασία αυτή ουσιαστικά είναι η μετατροπή των κλασσικών εισόδων σε ασαφείς ή διαφορετικά λεκτικές μεταβλητές. Κάθε κόμβος σε αυτό το επίπεδο αποτελεί μία λεκτική μεταβλητή αποτέλεσμα το οποίο έχει προέλθει από τις αντίστοιχες συναρτήσεις συμμετοχής παραδείγματα των οποίων περιγράφηκαν στο κεφάλαιο 3.2. Στα Takagi-Sugeno η συνηθέστερη συνάρτηση συμμετοχής αποτελεί η γκαουσιανή. Έστω, λοιπόν, για ένα χαρακτηριστικό εισόδου x_i υπολογίζεται η γκαουσιανή συνάρτηση συμμετοχής, όπως ακριβώς αποδίδεται και στη σχέση (3.3), η οποία ορίζεται από το κέντρο v και το εύρος σ :

$$\mu_{A_i^k}(x_i) = \exp \left\{ - \frac{(x_i - v_i)^2}{\sigma_i} \right\} \quad , i=1,2,\dots,d \text{ and } k=1,2,\dots,K \quad (3.6)$$

Όπου A_i είναι το ασαφές σύνολο, K ο συνολικός αριθμός των ασαφών κανόνων και d ο αριθμός των χαρακτηριστικών εισόδου x . Ο αριθμός των ασαφών κανόνων είναι $K \times d$.

- 2) Επίπεδο 2:** Ο αριθμός των κόμβων σε αυτό το επίπεδο είναι όσος είναι και ο αριθμός των ασαφών κανόνων. Κάθε κόμβος εκτελεί μια πράξη στις εξόδους του επιπέδου της

ασαφοποίησης με σκοπό να αποκτήσει ένα βαθμό (degree) των εισόδων με τα ασαφή σύνολα. Αυτό υπολογίζεται συνδυάζοντας όλες τις εισερχόμενες τιμές από το προηγούμενο επίπεδο με την ακόλουθη πράξη:

$$w_i = \prod_{k=1}^d \mu_{A_i^k}(x_i) \quad (3.7)$$

Όπου w_i είναι ο βαθμός του κόμβου i και $\mu_{A_i^k}$ η συνάρτηση συμμετοχής για το ασαφές σύνολο A_i^k όπως υπολογίστηκε από το επίπεδο 1.

- 3) **Επίπεδο 3:** Ανάλογα με τον κάθε εξαγόμενο βαθμό από το επίπεδο κανόνων η συνάρτηση εξόδου του i κόμβου υπολογίζεται σύμφωνα με το μέσο όρο των βαθμών w και συμβολίζεται ως εξής:

$$\bar{w}_i = w_i / \sum_{i=1}^c w_i \quad (3.8)$$

Όπου c είναι ο αριθμός των συναρτήσεων συμμετοχής κάθε χαρακτηριστικού εισόδου x_i .

- 4) **Επίπεδο 4:** Το συμπέρασμα κάθε κανόνα υπολογίζεται πολλαπλασιάζοντας τον αντίστοιχο κανόνα με το βαθμό του ως εξής:

$$\bar{y}_i = y_i \cdot \bar{w}_i \quad (3.9)$$

Όπου y_i είναι το συμπέρασμα του κόμβου i

- 5) **Επίπεδο 5:** Η τελική έξοδος του δικτύου ή του ασαφούς συστήματος υπολογίζεται προσθέτοντας όλα τα εισερχόμενα σταθμισμένα συμπεράσματα όπως αναλύθηκε και παραπάνω.

$$y = \sum_{i=1}^c \bar{y}_i \quad (3.10)$$

Κεφάλαιο 4

Βιβλιογραφική Έρευνα

4.1 Βιβλιογραφική Ανασκόπηση Μεθόδων

Στο συγκεκριμένο κεφάλαιο γίνεται αναφορά στη βιβλιογραφία και στις σχετικές μεθόδους που έχουν προταθεί είτε χρησιμοποιώντας μοντέλα μηχανικής μάθησης είτε ερμηνεύσιμες μεθόδους των μηχανών μάθησης είτε ασαφή μοντέλα, για πρόβλεψη ανεπιθύμητων συμβάντων στον ιατρικό τομέα και ειδικότερα των καρδιαγγειακών παθήσεων. Παρουσιάζονται οι εφαρμογές στις οποίες βρήκαν ανταπόκριση οι συγκεκριμένες μέθοδοι καθώς και η συνεισφορά τους στην υποβοήθηση της στήριξης μίας ιατρικής απόφασης. Από την ανασκόπηση αυτή απαντώνται διάφορα ερωτήματα όπως για παράδειγμα ποιος είναι ο ταξινομητής μηχανικής μάθησης που φέρει τα μεγαλύτερα ποσοστά ακρίβειας στο συγκεκριμένο πεδίο ή ποια είναι η πιο αποδοτική μέθοδος ερμηνευσιμότητας. Κάποιες από αυτές τις μεθόδους, που προτάθηκαν και είχαν διαθέσιμη την υλοποιημένη μεθοδολογία τους, δοκιμάστηκαν με διάφορες παραλλαγές.

4.1.1 Βιβλιογραφική Ανασκόπηση Μεθόδων Μηχανικής Μάθησης

Τα τελευταία χρόνια χρησιμοποιούνται συχνά τα ερμηνεύσιμα μοντέλα μηχανικής μάθησης ειδικά στον τομέα της υγείας. Έχουν αναπτυχθεί διάφορα ερμηνεύσιμα μοντέλα για ασθένειες ματιών όπως για παράδειγμα μοντέλο για τη διάγνωση του γλαυκώματος [62] και της διαβητικής αμφιβληστροειδοπάθειας [63], για διάφορα είδη καρκίνων όπως την πρόβλεψη του καρκίνου του μαστού [64] και της θνησιμότητας του καρκίνου του πνεύμονα [65], για ασθένειες γρίπης και λοίμωξης και πιο συγκεκριμένα μοντέλο για την πρόβλεψη θανάτων των ασθενών με γρίπη [66], για τη νόσο COVID-19 και ειδικότερα ερμηνεύσιμα μοντέλα σχετικά με τη θνησιμότητα της νόσου [67] αλλά και τον εντοπισμό της [68]. Από τις παραπάνω μεθοδολογίες ο ταξινομητής XGBoost αποδείχθηκε στις περισσότερες ο καλύτερος με υψηλό ποσοστό ακρίβειας και συγκεκριμένα στην περίπτωση του γλαυκώματος, του καρκίνου του πνεύμονα, της γρίπης και της θνησιμότητας του COVID-19. Σε αυτές τις περιπτώσεις χρησιμοποιήθηκε ως επί το πλείστον η βιβλιοθήκη SHAP (Shapley Additive exPlanations) για να επεξηγηθούν οι προβλέψεις.

Όσον αφορά τις καρδιαγγειακές παθήσεις, αποτελούν επίσης ένα πεδίο στο οποίο έχουν αναπτυχθεί διάφορα ερμηνεύσιμα μοντέλα πρόβλεψης. Οι Duarte et al. [69] ακολούθησαν μία μέθοδο bootstrap στην οποία συνδυάστηκε η λογιστική παλινδρόμηση, η ανάλυση γραμμικής διάκρισης και η μέθοδος ερμηνευσιμότητας Out-Of-Bag για την πρόβλεψη ασθενών με καρδιακή ανεπάρκεια. Οι Ghorbani et al. [70] εκπαίδευσαν ένα συνελκτικό νευρωνικό δίκτυο για να εντοπίσουν την ανατομία και τις δομές της καρδιάς από υπερηχοκαρδιογραφικές εικόνες. Στο κομμάτι της ερμηνευσιμότητας χρησιμοποιήθηκε ένας χάρτης προβολής (Saliency map) για να εξαχθεί ένα σκορ σημαντικότητας για κάθε χαρακτηριστικό. Σε αρκετές ερευνητικές εργασίες που αφορούν επίσης τις καρδιαγγειακές παθήσεις ο ταξινομητής με τα καλύτερα αποτελέσματα αποδείχθηκε ο XGBoost [71],[74]. Στην πλειονότητα αυτών ,επίσης, χρησιμοποιούνται κυρίως οι βιβλιοθήκες SHAP [72] αλλά και η PDP (Partial Dependence Plots) [73] για τη σημαντικότητα των χαρακτηριστικών.

Πρόσφατα, οι Moreno-Sanchez et al. [75] ανέπτυξαν ένα ερμηνεύσιμο μοντέλο μηχανικής μάθησης για την πρόβλεψη ασθενών με καρδιακή ανεπάρκεια χρησιμοποιώντας ένα δημόσια διαθέσιμο σετ δεδομένων το οποίο δημοσιεύθηκε από τους Ahmad et al. [76] στο UCI repository [77]. Οι Moreno-Sanchez et al. [75] πέτυχαν με τον XGBoost ακρίβεια 0.83 με τέσσερα σημαντικότερα χαρακτηριστικά τα χαρακτηριστικά του χρόνου (time), του επιπέδου της κρεατινίνης (serum creatinine) , του κλάσματος εξώθησης υγρού (ejection fraction) και της αναιμίας (anaemia). Η παρούσα εργασία η οποία ακολουθεί τη μεθοδολογία αυτή [75], καλύπτει το κενό που υπάρχει για τη διαχείριση του μη ισορροπημένου δείγματος και προσπαθεί να εξάγει καλύτερα αποτελέσματα. Χρησιμοποιείται, λοιπόν, το δημόσια διαθέσιμο σετ δεδομένων των Ahmad et al [76] και αναπτύσσεται ένα ερμηνεύσιμο μοντέλο μηχανικής μάθησης με διάφορους ταξινομητές αλλά και αλγορίθμους για τη σημαντικότητα των χαρακτηριστικών ώστε να βελτιώσει τα αποτελέσματα που υπάρχουν στη συγκρινόμενη μεθοδολογία για την πρόβλεψη ασθενών με καρδιακή ανεπάρκεια. Ακολουθήθηκε η συγκεκριμένη μέθοδος σε προηγούμενη εργασία και υλοποιήθηκε με κάποιους επιπλέον ταξινομητές (με όσους αναφέρθηκαν στο κεφάλαιο 2) ώστε να εξαχθούν κάποια αποτελέσματα. Πραγματοποιήθηκαν πειράματα με αλλαγές στις παραμέτρους του κάθε ταξινομητή με στόχο την εξαγωγή καλύτερων αποτελεσμάτων αλλά και πειράματα με διαφορετικά σετ δεδομένων για αξιολόγηση της μεθόδου. Σε επόμενη ενότητα αναλύονται τα αποτελέσματα που παρήχθησαν από αυτές τις δοκιμές καθώς και σημαντικά συμπεράσματα που εξάγονται.

4.1.2 Βιβλιογραφική Ανασκόπηση Ασαφών Μοντέλων

Τα τελευταία χρόνια τα ασαφή μοντέλα αναπτύσσονται ολοένα και περισσότερο καθώς όπως προαναφέρθηκε αποτελούν από μόνα τους ερμηνεύσιμα και κατανοητά από την ανθρώπινη λογική. Βρίσκουν εφαρμογή σε διάφορα ερευνητικά πεδία όπως για παράδειγμα σε επεξεργασία κειμένου, τόσο στον ιατρικό αλλά και στον εκπαιδευτικό τομέα, σε προβλήματα ιατρικής διάγνωσης, στην ψυχολογία [92], σε ρομποτικές εφαρμογές [93] αλλά και στην τεχνολογία ξήρανσης [94]. Για την εξόρυξη κειμένου (text mining) για παράδειγμα έχει προταθεί ασαφής μέθοδος κατά την οποία πραγματοποιείται μια προσπάθεια για την επίλυση του προβλήματος του θορύβου και της αραιότητας σε σύντομα κείμενα [78]. Ένα ασαφές μοντέλο, επίσης, έχει αναπτυχθεί με στόχο τη δημιουργία μίας προσωπικής web αναζήτησης ώστε να φέρει ακριβή αποτελέσματα σύμφωνα με την αναζήτηση του χρήστη [79]. Η ασαφής λογική σε αυτή τη μέθοδο έρχεται να διαχειριστεί την αβεβαιότητα σε κάθε τίτλο ενός άρθρου και τελικά αποδεικνύεται πως το πετυχαίνει με μεγάλο ποσοστό ακρίβειας σε σύγκριση με άλλες μεθόδους [79]. Αναφορικά με τον ιατρικό τομέα στην επεξεργασία κειμένου έχουν γίνει προσπάθειες για μοντελοποίηση των ιατρικών κειμένων με ασαφή λογική με λιγότερο υπολογιστικό κόστος και υψηλά ποσοστά ακρίβειας [80].

Αναφορικά με τα προβλήματα ιατρικής διάγνωσης τα ασαφή συστήματα φέρουν κι εκεί υψηλά ποσοστά επιτυχίας. Έχει εφαρμοστεί για ταξινόμηση διάφορων ασθενειών όπως για παράδειγμα για το Αλτσχάιμερ [81], για τη χρόνια νεφρική νόσο [82-83], για τη χολέρα [84], για τον COVID-19 [85], για την ηπατίτιδα [86], για το διαβήτη [87], για το άσθμα [88] αλλά και φυσικά για ασθένειες που αφορούν την καρδιά [89-91] καθώς και διάφορες άλλες. Η χρήση της ασαφούς λογικής σε όλες τις παραπάνω μεθοδολογίες διάγνωσης αποδεικνύεται από τους ερευνητές ιδιαίτερα χρήσιμη καθώς αντιμετωπίζει την ανακρίβεια των στοιχείων αλλά και είναι ικανή να υπολογίσει την πιθανότητα κατά πόσο κοντά είναι η υπολογιζόμενη διάγνωση από την πραγματικότητα. Πολλοί ερευνητές στρέφονται προς τα υβριδικά ασαφή μοντέλα τα οποία συνδυάζουν τα ασαφή συστήματα με τα νευρωνικά δίκτυα σχετικά με το πρόβλημα των ιατρικών διαγνώσεων (adaptive neuro fuzzy systems- ANFIS), όπως αναλύθηκε παραπάνω το συγκεκριμένο δίκτυο, και φαίνεται να εξάγουν μεγαλύτερα ποσοστά ακρίβειας από τα αυθεντικά ασαφή συστήματα [83], [86-88]. Για αυτό τον τελευταίο καιρό γίνονται προσπάθειες για ανάπτυξη πολλαπλών επιπέδων νευρωνικών δικτύων ασαφών συστημάτων για ιατρικές διαγνώσεις με σκοπό την εξαγωγή καλύτερων αποτελεσμάτων.

Τα Takagi-Sugeno, μία ειδική κατηγορία ασαφών συστημάτων που μελετάται στην παρούσα εργασία, βρίσκουν επίσης διάφορες εφαρμογές στην διάγνωση ασθενειών αλλά συνδυάζονται κι αυτά με νευρωνικά δίκτυα ώστε να πετύχουν μεγαλύτερα ποσοστά ακρίβειας ταξινόμησης. Έχουν εφαρμοστεί σε ασθένειες όπως για παράδειγμα του αυτισμού [95], μελέτη των σημάτων του εγκεφάλου για διάφορα συμπεράσματα όπως διάγνωση της κόπωσης οδήγησης [96], για πρόβλεψη θεραπευτικών πεπτιδίων [97], για διάγνωση της λευχαιμίας μέσω τμηματοποίησης μικροσκοπικών εικόνων κυττάρων [98], για εντοπισμό άνοιας μέσω εικόνων εγκεφάλου [99] και διάφορες άλλες. Στη συγκεκριμένη εργασία για να εφαρμοστούν τα Takagi-Sugeno και να συγκριθούν με τους υπόλοιπους ταξινομητές μηχανικής μάθησης μελετήθηκε μία μεθοδολογία των Y. Cui, D. Wu, X. Jiang & Y. Xu [100] η οποία παρέχει ένα πακέτο στην Python το **PyTSK**. Είναι ένα πακέτο για δοκιμές και πειράματα με τους ταξινομητές αυτούς. Με αυτό το πακέτο, λοιπόν, πραγματοποιήθηκαν πειράματα με διαφορετικά σετ δεδομένων, διαφορετικό αριθμό ασαφών κανόνων, διαφορετικό τύπο Takagi-Sugeno (zero & first order) αλλά και διαφορετικές μεθοδολογίες βελτιστοποίησης.

Κεφάλαιο 5

Μεθοδολογία

5.1 Μέθοδος Ασαφών Φράσεων Ομοιότητας (Fuzzy Similarity Phrases-FSP)

Βασιζόμενοι στα ενθαρρυντικά αποτελέσματα από τους ταξινομητές Takagi-Sugeno σε προηγούμενο κεφάλαιο στην παρούσα πτυχιακή εργασία πραγματοποιήθηκε η αναπαραγωγή υλοποίησης ενός συστήματος ταξινόμησης βασισμένο στην ασαφή λογική και στους ασαφείς ταξινομητές Takagi-Sugeno [100]. Πιο συγκεκριμένα, η ανάγκη για ερμηνεύσιμες μεθόδους ταξινόμησης κατανοητές από την ανθρώπινη λογική οδήγησαν στην ανάπτυξη της παρούσας ιδέας. Ο κύριος άξονας πάνω στον οποίο αναπτύχθηκε η ιδέα του παρόντος ασαφή ταξινομητή αποτελεί η μετρική της ομοιότητας (similarity). Σύμφωνα, λοιπόν, με τις ομοιότητες μεταξύ των κλάσεων αναπτύσσονται ασαφείς κανόνες με στόχο την ακριβή ταξινόμηση άγνωστων δειγμάτων. Η μεθοδολογία, λοιπόν, αποκαλείται ασαφείς φράσεις ομοιότητας (fuzzy similarity phrases-FSP) [105]. Παρακάτω αναλύονται τα βήματα της μεθοδολογίας των ασαφών φράσεων ομοιότητας.

5.1.1 Κατασκευή Μοντέλου

Η μέθοδος των ασαφών φράσεων ομοιότητας ανήκει στην κατηγορία της επιβλεπόμενης μηχανικής μάθησης. Βασίζεται σε D-διάστατα διανύσματα χαρακτηριστικών εκπαίδευσης τα οποία ανήκουν σε C διαφορετικές κλάσεις. Έστω λοιπόν $X^C = \{\mathbf{x}_1^C, \mathbf{x}_2^C, \dots, \mathbf{x}_{N^C}^C\}$ το οποίο είναι ένα σύνολο από N_C διανύσματα χαρακτηριστικών εκπαίδευσης $\mathbf{x}_n^C = (x_{n,1}^C, x_{n,2}^C, \dots, x_{n,D}^C)$, $n = 1, \dots, N$ τα οποία ανήκουν στην κλάση $c=1,2,\dots,C$. Αρχικά χωρίζεται το σύνολο δεδομένων σε C υποσύνολα σύμφωνα με την κλάση στην οποία ανήκουν τα χαρακτηριστικά. Για παράδειγμα όταν $C=2$, χωρίζονται τα χαρακτηριστικά της κλάσης 1 και τα χαρακτηριστικά της κλάσης 0 σε δύο διαφορετικά σύνολα. Σε κάθε σύνολο από αυτά εφαρμόζεται ο αλγόριθμος μη επιβλεπόμενης ομαδοποίησης k-means με σκοπό την ομαδοποίηση του κάθε χαρακτηριστικού διανυσμάτων X^C , $c=1,2,\dots,C$ σε $M^C < N^C$ συστάδες. Στα πλαίσια αυτής της πτυχιακής εργασίας, εκτός από τον αλγόριθμο kmeans δοκιμάστηκε ο αλγόριθμος συσταδοποίησης ο οποίος είναι ένα μοντέλο Μίγματος Γκαουσιανών (gauss mixture model-GMM) και η έκδοση του μοντέλου για αυτήν την παραλλαγή αποδίδεται ως **FSP-GMM** και **FSP-GD-GMM**. Ο αλγόριθμος αυτός χρησιμοποιεί τον

αλγόριθμο αναμονής-μεγιστοποίησης (expectation – maximization). Ένα μοντέλο Gaussian Mixture παρουσιάζει μια σύνθετη κατανομή της οποίας τα σημεία προέρχονται από k Gaussian υπό-κατανομές. Ειδικότερα, κάθε ομάδα μπορεί να αναπαρασταθεί μαθηματικά από μια παραμετρική Gaussian κατανομή με κέντρο τα βαρύκεντρα τους και ολόκληρο το σύνολο δεδομένων από μια μίξη από αυτές τις κατανομές.

Αλγόριθμοι Συσταδοποίησης			
K-means	Δημιουργία Συστάδων βασισμένων στην Ευκλείδεια Απόσταση	Σφαιρικές Συστάδες	$d_{sq} = \sum_{i=1}^D (x_i - y_i)^2$
Gaussian Mixture Model	Εκτιμήσεις πυκνότητας πιθανότητας χρησιμοποιώντας τον αλγόριθμο Αναμονής-Μέγιστοποίησης (Expectation-	Συστάδες βασισμένες στην Γκαουσιανή Κατανομή	$\mathcal{N}(X \mu, \Sigma) = \frac{1}{(2\pi)^{\frac{D}{2}} \sqrt{ \Sigma }} \exp \left\{ -\frac{(X - \mu)^T \Sigma^{-1} (X - \mu)}{2} \right\}$

Πίνακας 1 Αλγόριθμοι συσταδοποίησης που δοκιμάστηκαν

Με αυτόν τον τρόπο εξάγονται τα διανύσματα των κέντρων των συστάδων που δημιουργήθηκαν, έστω $\mathbf{w}_m^C = (w_{m,1}^C, w_{m,2}^C, \dots, w_{m,D}^C)$, $m = 1, \dots, M$. Στόχος σε αυτό το βήμα αποτελεί να δημιουργηθεί ένα λεξικό (codebook) με λέξεις $\bar{W}^C = \{\mathbf{w}_1^C, \mathbf{w}_2^C, \dots, \mathbf{w}_M^C\}$ για κάθε κλάση C το οποίο τελικά θα αντικατοπτρίζει το περιεχόμενο και τα χαρακτηριστικά των αρχικών δεδομένων. Επόμενο βήμα αποτελεί η αξιοποίηση αυτού του λεξικού. Σύμφωνα, λοιπόν, με το λεξικό που έχει δημιουργηθεί, σε αυτό το βήμα στόχος είναι να κωδικοποιηθούν όλα τα αρχικά δεδομένα x_n^C , $n=1, \dots, N^C$, $c=1, 2, \dots, C$ σε διανύσματα ομοιοτήτων με τις λέξεις του λεξικού $\bar{W}^C$ της κάθε κλάσης. Ο τύπος της ομοιότητας $S_{n,m}^{C,C'}$ που υπολογίζεται μεταξύ των διανυσμάτων χαρακτηριστικών x_n^C και των λέξεων $w_m^{C'}$, $m = 1, \dots, M$, $c' = 1, \dots, C$ σε αυτό το βήμα είναι ο παρακάτω:

$$S_{n,m}^{C,C'} = 1 - \frac{\|x_n^C - w_m^{C'}\|}{\sum_{j=1}^M \sum_{i=1}^{M_j} \|x_n^C - w_i^j\|} \quad (5.1)$$

Όπου $S_{n,m}^{C,C'} \in [0,1]$ και $\|\cdot\|$ αποτελεί μία προτιμώμενη απόσταση η οποία στην παρούσα εργασία είναι η Ευκλείδεια Απόσταση ως πρώτη εκδοχή που δοκιμάστηκε. Η λογική του συγκεκριμένου τύπου παρέχει μια κανονικοποιημένη εκτίμηση της ομοιότητας, λαμβάνοντας υπόψη ότι ένας άνθρωπος αναγνωρίζει και ταξινομεί αντικείμενα βασιζόμενος στις συγκρίσεις με προηγούμενα γνωστά πρότυπα. Συγκεκριμένα, θεωρεί ότι για την επίλυση ενός προβλήματος ταξινόμησης, ένα διάνυσμα χαρακτηριστικών x_n^C πρέπει να συγκριθεί όχι μόνο με έναν όρο w_m^C , αλλά και με τους υπόλοιπους όρους που αποτελούν τον κώδικα. Με αυτό τον τρόπο, η ομοιότητα του x_n^C προς τον w_m^C θα είναι υψηλότερη εάν η συνολική απόσταση του x_n^C από τους υπόλοιπους όρους είναι υψηλή.

Επίσης, αυτή η μέτρηση ομοιότητας λειτουργεί ως δεσμευτικός παράγοντας, διότι εάν ένα διάνυσμα χαρακτηριστικών x_n^C έχει την ίδια απόσταση από διάφορους όρους, η ομοιότητα του διανύσματος χαρακτηριστικών x_n^C προς κάθε όρο θα είναι διαφορετική, διότι λαμβάνει υπόψη τις αποστάσεις του x_n^C από τους υπόλοιπους όρους στον κώδικα. Η ιδέα σε αυτό το βήμα βασίζεται στην υπόθεση πως τα δεδομένα και οι λέξεις που ανήκουν στην ίδια κλάση θα έχουν μεγαλύτερη τιμή ομοιότητας από αυτά που δεν ανήκουν στην ίδια κλάση. Έτσι, κάθε διάνυσμα χαρακτηριστικών x_n^C μετά τον υπολογισμό των ομοιοτήτων θα μετατραπεί στο εξής διάνυσμα: $S_n^{C,C'} = (S_{n,1}^{C,C'}, S_{n,2}^{C,C'}, \dots, S_{n,M}^{C,C'})$. Στα πλαίσια αυτής της πτυχιακής εργασίας δοκιμάστηκε μία μικρή παραλλαγή με εφαρμογή και της ομοιότητας συνημίτονου (**FSP-COS-SIM**, **FSP-GD-COS-SIM**) σύμφωνα με τον παρακάτω τύπο:

$$S_{n,m}^{C,C'} = \frac{x_n^C \cdot w_m^{C'}}{\|x_n^C\| \|w_m^{C'}\|} \quad (5.2)$$

Στη συνέχεια για τη μελέτη και δοκιμή διάφορων πειραμάτων δοκιμάστηκε ο τύπος 5.1 με επιπλέον αποστάσεις, τις αποστάσεις Mahalanobis (**FSP-SIM-MAH**, **FSP-GD-SIM-MAH**) και Manhattan (**FSP-SIM-MAN**, **FSP-GD-SIM-MAN**), όπως αντίστοιχα φαίνονται και στον παρακάτω πίνακα.

Ομοιότητες	Τύπος	Τύπος Απόστασης
Ομοιότητα όπου $\ \cdot\ $ Ευκλείδεια Απόσταση	$S_{n,m}^{C,C'} = 1 - \frac{\ x_n^C - w_m^{C'}\ }{\sum_{j=1}^C \sum_{i=1}^{M_j} \ x_n^C - w_i^j\ }$	$\ \cdot\ = \ x - y\ = d_{EUC} (x - y) = \sqrt{(x - y)^2}$
Ομοιότητα Συνημίτονου	$S_{n,m}^{C,C'} = \frac{x_n^C * w_m^{C'}}{\ x_n^C\ \ w_m^{C'}\ }$	
Ομοιότητα όπου $\ \cdot\ $ Απόσταση Manhattan	$S_{n,m}^{C,C'} = 1 - \frac{\ x_n^C - w_m^{C'}\ }{\sum_{j=1}^C \sum_{i=1}^{M_j} \ x_n^C - w_i^j\ }$	$\ \cdot\ = \ x - y\ = d_{MAN} (x, y) = \sum x - y $
Ομοιότητα όπου $\ \cdot\ $ Απόσταση Mahalanobis	$S_{n,m}^{C,C'} = 1 - \frac{\ x_n^C - w_m^{C'}\ }{\sum_{j=1}^C \sum_{i=1}^{M_j} \ x_n^C - w_i^j\ }$	$\ \cdot\ = \ x - y\ = d_{MAH} (x, y) = \sqrt{(x - y)^T S^{-1} (x - y)}$

Πίνακας 2 Εκδοχές με διάφορες ομοιότητες που δοκιμάστηκαν

Η λογική για την εξέταση των ομοιοτήτων και έπειτα στη μοντελοποίηση του χώρου χαρακτηριστικών είναι ότι όσο πιο όμοιο είναι ένα διάνυσμα χαρακτηριστικών με μία λέξη, τόσο πιο αντιπροσωπευτική θα είναι η λέξη για αυτό το διάνυσμα χαρακτηριστικών. Τα διανύσματα εκπαίδευσης ενός ταξινομητή c δημιουργούν τον παρακάτω πίνακα ομοιότητας:

$$S^{C,C'} = \begin{pmatrix} S_1^{C,C'} \\ \vdots \\ S_{Nc}^{C,C'} \end{pmatrix} = \begin{pmatrix} S_{1,1}^{C,C'} & \cdots & S_{1,Mc}^{C,C'} \\ \vdots & \ddots & \vdots \\ S_{Nc,1}^{C,C'} & \cdots & S_{Nc,Mc}^{C,C'} \end{pmatrix}$$

Αφού υπολογιστεί ο παραπάνω πίνακας, εφαρμόζεται ο αλγόριθμος ομαδοποίησης kmeans στα στοιχεία του πίνακα ανά στήλη ώστε να ομαδοποιηθούν το αντίστοιχο σύνολο ομοιοτήτων $S_m = \{S_{n,m}^c | n = 1, \dots, Nc\}$ σε ένα σύνολο από L^c μονοδιάστατες συστάδες. Στα πλαίσια αυτής της πτυχιακής εργασίας, εκτός από τον αλγόριθμο kmeans δοκιμάστηκε ο αλγόριθμος συσταδοποίησης ο οποίος είναι ένα μοντέλο Μίγματος Γκαουσιανών (gauss mixture model-GMM) (**FSP-FUZZ-GMM, FSP-GD-FUZZ-GMM**) όπως ακριβώς και στην αρχική ομαδοποίηση. Το αποτέλεσμα από αυτή την ομαδοποίηση ανά στήλη είναι ο παρακάτω πίνακας.

$$A^C = \begin{pmatrix} a_{1,1}^C & \cdots & a_{1,Mc}^C \\ \vdots & \ddots & \vdots \\ a_{Lc,1}^C & \cdots & a_{Lc,Mc}^C \end{pmatrix}, \text{ όπου } c \text{ αποτελεί μία συντομογραφία του } (C, C') \text{ και } a_{l,m}^C, m=1, \dots, Mc \text{ τα}$$

διανύσματα των κέντρων των συστάδων που ταξινομήθηκαν κατά αύξουσα σειρά.

Για κάθε συστάδα των στοιχείων του S_m με κέντρο $a_{l,m}^C$, ορίζεται ένα ασαφές σύνολο $A_{l,m}^C$. Κάθε σύνολο χαρακτηρίζεται από μία συνάρτηση συμμετοχής $\mu_{l,m}^C \in [0,1]$, $s \in [0,1]$ η οποία για λόγους απλότητας είναι η τριγωνική και η τραπεζοειδής όπως αναλύθηκε σε προηγούμενο κεφάλαιο. Σε αυτό το σημείο της μεθόδου στα πλαίσια της παρούσας πτυχιακής εργασίας δοκιμάστηκαν παραπάνω και οι γκαουσιανές συναρτήσεις συμμετοχής (**FSP-GAUSS, FSP-GD-GAUSS**). Η κορυφή της τριγωνικής συνάρτησης αντιστοιχεί στο κέντρο της συστάδας $a_{l,m}^C$ και οι βάσεις της στη χαμηλότερη και υψηλότερη τιμή των ομοιοτήτων που ανήκουν στην αντίστοιχη συστάδα στο S_m . Αυτά τα ασαφή σύνολα ορίζονται ανά κλάση και αναπαριστούν ασαφή επίθετα (fuzzy adjectives) για το χαρακτηρισμό του βαθμού ομοιότητας των διανυσμάτων εκπαίδευσης με τις αντίστοιχες λέξεις. Ένα ασαφές επίθετο που εκφράζει μια ελάχιστη επιτρεπόμενη ομοιότητα όπως «Χαμηλή» αποδίδεται στο ασαφές σύνολο $A_{1,m}^C$ το οποίο αντιστοιχεί στο κέντρο συστάδας $a_{1,m}^C$ που είναι πιο κοντά στο μηδέν. Αντίστοιχα ασαφές επίθετο που εκφράζει μια μέγιστη ομοιότητα όπως «Υψηλή» αποδίδεται στο ασαφές σύνολο $A_{L,m}^C$ το οποίο αντιστοιχεί στο κέντρο συστάδας $a_{L,m}^C$ που είναι πιο κοντά στο 1 (όπως αναπαρίσταται στην Εικόνα 12). Ωστόσο, ένα πρόβλημα με αυτήν τη διαδικασία είναι ότι τα διαστήματα των τριγωνικών συναρτήσεων ενδέχεται να μην καλύπτουν το σύνολο των δεδομένων. Για να αντιμετωπιστεί αυτό το πρόβλημα, οι συναρτήσεις συμμετοχής προσαρμόζονται ως εξής: α) τα διαστήματα των ενδιάμεσων τριγώνων επεκτείνονται

προς το πλησιέστερο κέντρο β) τα διαστήματα των αριστερότερων και δεξιότερων συναρτήσεων συμμετοχής επεκτείνονται στο μηδέν και το ένα, αντίστοιχα γ) το δεξιότερο τρίγωνο μετατρέπεται σε τραπεζοειδή συνάρτηση με τη δεξιά πλευρά της να φτάνει μέχρι το πλησιέστερο κέντρο, με τιμή ένα, για ομοιότητες που είναι μεγαλύτερες από το $a_{L,m}^c$, αντικατοπτρίζοντας την έννοια του "Υψηλότερου", όπως φαίνεται στην Εικόνα 12 και διαγραμματικά στην Εικόνα 13· δ) η δεξιά πλευρά της αριστερότερης συνάρτησης μέλους επεκτείνεται προς το πλησιέστερο κέντρο, και η αριστερή πλευρά της φτάνει μέχρι το μηδέν για ομοιότητες που είναι μικρότερες από το $a_{1,m}^c$, λαμβάνοντας υπόψη ότι μικρότερες τιμές ομοιότητας μπορεί να είναι λιγότερο αξιόπιστες λόγω της αβεβαιότητας των δεδομένων. Πιο συγκεκριμένα, το αριστερό μέρος της αριστερής συνάρτησης συμμετοχής του τριγωνικού σχήματος ορίζεται στο διάστημα 0 έως $a_{1,m}^c$, για να αναπαραστήσει τις τιμές που είναι πιο κοντά στο μηδέν ανάλογα με την αντίστοιχη τιμή ομοιότητας. Έτσι, στις περιπτώσεις όπου η ομοιότητα είναι ίση με $a_{1,m}^c$, η τιμή συμμετοχής είναι ίση με 1 και μειώνεται γραμμικά στο 0 όταν η ομοιότητα έχει τιμή μικρότερη από $a_{1,m}^c$. Με βάση τα παραπάνω, δίνονται οι ακόλουθοι ορισμοί για να περιγράψουν τον έννοια ενός μοντέλου FSP που έχει οριστεί για ένα ή περισσότερα κλάσματα δεδομένων.

Εκδοχές Μεθόδου	
Συναρτήσεις Συμμετοχής	Τύπος
Τριγωνικές + Τραπεζοειδής	$\mu_{tri}(x) = \begin{cases} \frac{x-a}{b-a}, & \text{αν } a < x \leq b \\ \frac{b-x}{c-b}, & \text{αν } b < x < c \\ 0, & \text{διαφορετικά} \end{cases}$ $\mu_{trap}(x) = \begin{cases} \frac{x-a}{b-a}, & \text{αν } a < x < b \\ 1, & \text{αν } b \leq x \leq c \\ \frac{d-x}{d-c}, & \text{αν } c < x < d \\ 0, & \text{διαφορετικά} \end{cases}$
Γκαουσιανές	$\mu_{gauss}(x) = \exp \left\{ -\frac{(x-m)^2}{\sigma^2} \right\}$

Πίνακας 3 Εκδοχές που δοκιμάστηκαν με συναρτήσεις συμμετοχής

Ορισμός 1: Δεδομένου ενός συνόλου διανυσμάτων εκπαίδευσης x_n^c στο X^c , όπου $n = 1, 2, \dots, N^c$, από μια κλάση c , ορίζεται ένα μοντέλο FSP F^c για να περιγράψει τον χώρο χαρακτηριστικών της κλάσης c , με έναν λεξικό W^c που περιέχει M^c λέξεις και ένα σύνολο L^c ασαφών συνόλων $A_{l,m}^c$, όπου $l = 1, \dots, L^c$, για κάθε λέξη w_m^c στο W^c , με $m = 1, 2, \dots, M^c$. Τα ασαφή σύνολα αντιπροσωπεύουν ασαφή επίθετα που εκφράζουν το βαθμό ομοιότητας κάθε διανύσματος εκπαίδευσης προς κάθε λέξη.

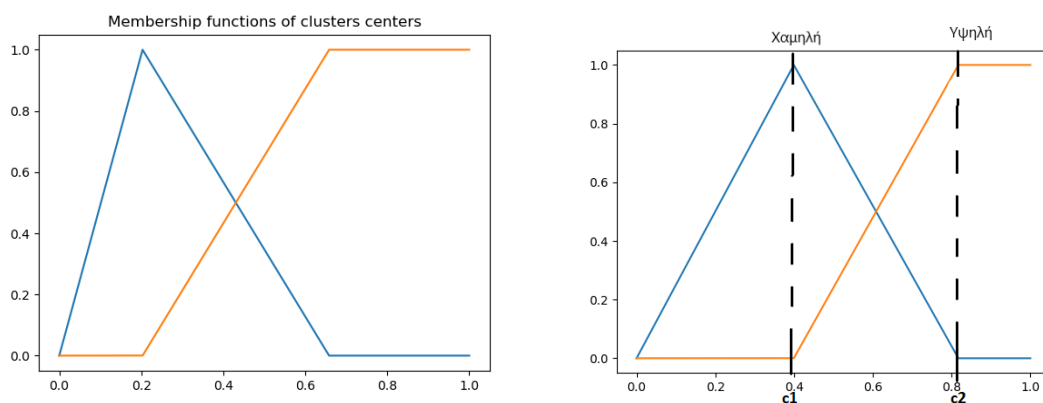
$$F^C(M^C, L^C) = (W^C, A^C) \quad (5.3),$$

Όπου $A^C = \{A_{l,m}^c | l = 1, \dots, lc, m = 1, \dots, mc\}$

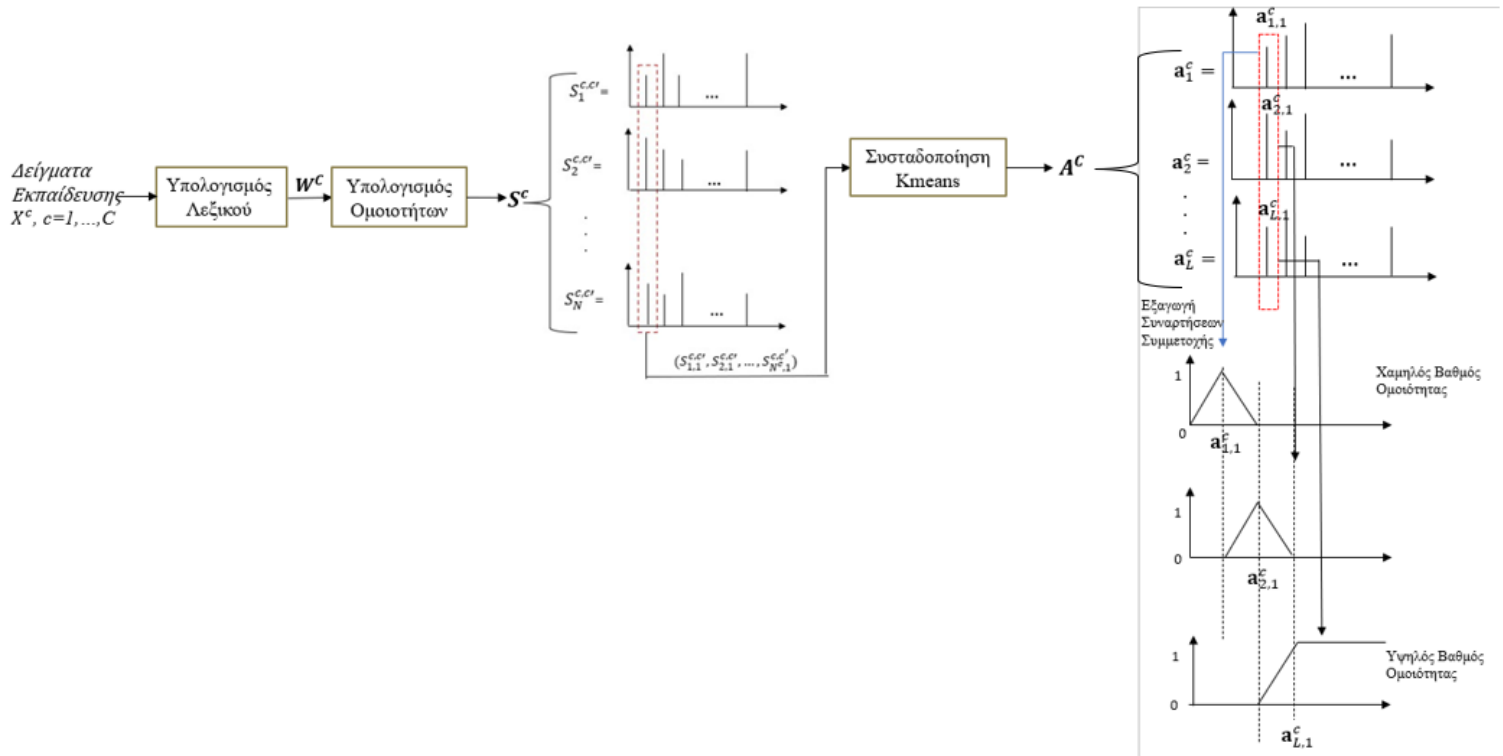
Σε αυτό το μοντέλο FSP, τα ασαφή σύνολα κάθε λέξης αντιπροσωπεύουν τον βαθμό ομοιότητας κάθε διάνυσματος εκπαίδευσης προς την συγκεκριμένη λέξη. Με άλλα λόγια, κάθε διάνυσμα εκπαίδευσης x_n^c έχει μια τιμή ομοιότητας για κάθε λέξη w_m^c στο λεξικό W^c , η οποία περιγράφεται από το αντίστοιχο ασαφές σύνολο $A_{1,m}^c$. Έτσι, μπορούμε να αξιολογήσουμε πόσο κοντά είναι κάθε διάνυσμα εκπαίδευσης σε κάθε λέξη του λεξικού, λαμβάνοντας υπόψη το βαθμό ομοιότητας που έχει με τις ασαφείς συναρτήσεις του κάθε w_m^c .

Ορισμός 2: Ένα μοντέλο FSP F ορίζεται για να περιγράψει τους χώρους χαρακτηριστικών των C κλάσεων, ως ένα σύνολο από C μίας-κλάσης FSP μοντέλα.

$$F(C) = \{F^C(M^C, L^C) | c=1, \dots, C\} \quad (5.4)$$



Εικόνα 12 Συναρτήσεις συμμετοχής των εξαγόμενων ασαφών συνόλων.



Εικόνα 13 Διαγραμματική απεικόνιση κατασκευής μοντέλου.

5.1.2 Δημιουργία Φράσεων για Ερμηνεύσιμη Ταξινόμηση

Έστω λοιπόν ένα μοντέλο FSP πολλαπλών κλάσεων $F(C)$ και x^u ένα διάνυσμα άγνωστης κλάσης. Για να ταξινομηθεί αυτό το διάνυσμα (x^u) σε μια κλάση c ($c = 1, \dots, C$), υπολογίζονται οι ομοιότητες ($s_m^{u,c}$) του x^u ως προς κάθε λέξη (w_m) του λεξικού W^c και οι τιμές συμμετοχής $\mu_{l,m}^c$ ($s_m^{u,c}$) των αντίστοιχων ασαφών συνόλων $A_{l,m}^c$, για $l = 1, \dots, L^c$. Για κάθε $s_m^{u,c}$, επιλέγουμε το ασαφές σύνολο στο οποίο η $s_m^{u,c}$ έχει τη μέγιστη τιμή συμμετοχής για να χαρακτηρίσει τον αντίστοιχο βαθμό ομοιότητας.

Ορισμός 3: Δοθέντος ενός διανύσματος εισόδου x_u και ενός μοντέλου FSP $F(C)$, ένα FSP ορίζεται ως ένα σύνολο $\theta^{u,c} \subseteq A^c$ από ασαφή σύνολα που εξηγούν αν το x_u ανήκει σε μια κλάση c ($c = 1, \dots, C$), σε σχέση με την ομοιότητά του με ένα σύνολο λέξεων $M^c \subseteq W^c$ από το μοντέλο $F(C)$.

$$\theta^{u,c} = \{ A_{l,m}^c \in A^c \mid l = \arg\max(\mu_{l,m}^c(s_m^{u,c})), \forall m: w_m^c \in M^c \} \quad (5.5)$$

Κάθε FSP μπορεί να μεταφραστεί σε μια γλωσσική έκφραση της μορφής: "Το διάνυσμα εισόδου x_u μπορεί να ανήκει στην κλάση c επειδή [η ομοιότητά του με τη λέξη w_m^c είναι $A_{l,m}^c$]", με τη

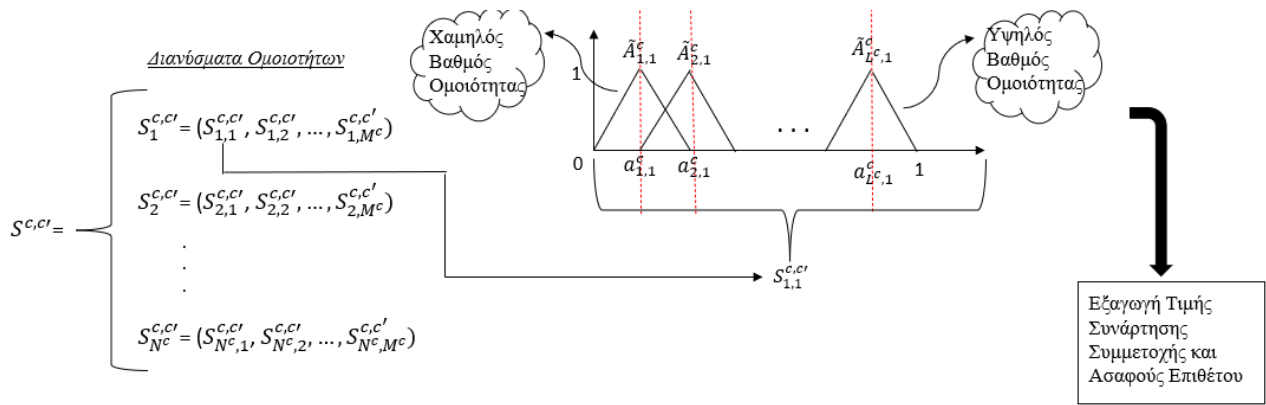
φράση μέσα στις αγκύλες να επαναλαμβάνεται για κάθε m : $w_m^c \in M^c$. Είναι σημαντικό να σημειωθεί ότι ένα FSP, όπως δημιουργείται από την παραπάνω εξίσωση, είναι ένα σύνολο από ασαφή σύνολα· συνεπώς, η ισότητα δύο FSP ορίζεται από την ισότητα των αντίστοιχων ασαφών συνόλων.

Ορισμός 4: Για οποιοδήποτε FSP, ένας τελεστής συνάφειας $\sigma(\cdot)$ ορίζεται ως το εσωτερικό γινόμενο ενός διανύσματος $s^{u,c}$, το οποίο αποτελείται από όλες τις ομοιότητες $s_m^{u,c}$ για κάθε $w_m^c \in M^c$, ταξινομημένες κατά m και με βάρη q^c , με ένα διάνυσμα $\mu_{l,m}^c$, το οποίο αποτελείται από όλες τις αντίστοιχες τιμές συμμετοχής $\mu_{l,m}^c(s_m^{u,c})$ των ασαφών συνόλων $A_{l,m}^c$. Με άλλα λόγια, ο τελεστής συνάφειας $\sigma(\cdot)$ υπολογίζεται ως εξής:

$$\sigma(\theta^{u,c}, q^c) = \frac{s^{u,c} \cdot (e^u \odot \mu^{u,c})}{\Lambda_{\theta^{u,c}}} \quad (5.6)$$

όπου $\odot$ αντιπροσωπεύει το γινόμενο Hadamard και το διάνυσμα βαρών $q^c = (q^c, \dots, q^c)$ το οποίο κυμαίνεται μέσα στο διάστημα $[0,1]$ και χρησιμοποιείται για να ρυθμίσει τις συνεισφορές των αντίστοιχων συναρτήσεων συμμετοχής, δηλαδή αντιπροσωπεύει ένα παράγοντα σημαντικότητας για κάθε ασαφές σύνολο στην κλάση c . Από προεπιλογή, μπορεί να θεωρηθεί ίσο με το διάνυσμα μονάδας $q^c = 1$.

Ο τελεστής που ορίζεται από την παραπάνω συνάρτηση υπολογίζει το $\theta^{u,c}$ σε μία κλίμακα με τιμή μέσα στο διάστημα $[0,1]$, που εκφράζει τη βεβαιότητα του x_u να ανήκει στην κλάση c . Συνεπώς, η γλωσσική έκφραση του $\theta^{u,c}$ μπορεί να λάβει τη μορφή: "Το διάνυσμα εισόδου x_u ανήκει στην κλάση c με μια βεβαιότητα $\sigma(\theta^{u,c}, q^c)$ επειδή [η ομοιότητά του με τη λέξη w_m^c είναι $A_{l,m}^c$]," με τη φράση μέσα στις αγκύλες να επαναλαμβάνεται για κάθε $w_m^c \in M^c$.



Εικόνα 14 Διαγραμματική απεικόνιση δημιουργίας φράσεων.

5.1.3 Επιλογή Χαρακτηριστικών με μείωση Φράσεων

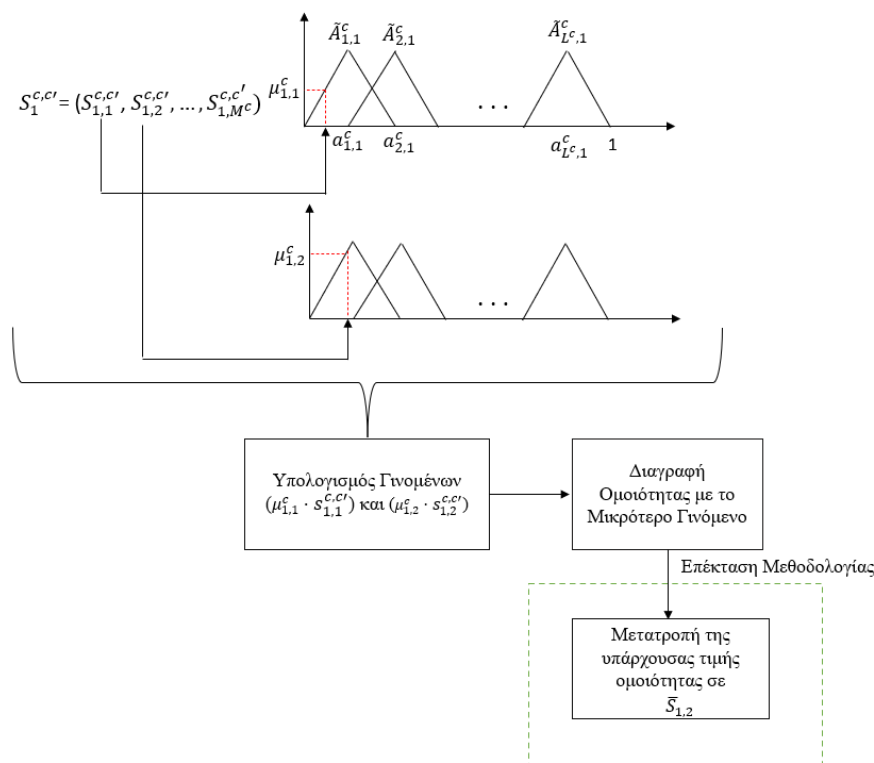
Σύμφωνα με τον ορισμό 3, ένα FSP μπορεί να περιλαμβάνει περισσότερα από ένα ασαφή σύνολα που αντιστοιχούν στον ίδιο βαθμό ομοιότητας για διαφορετικές λέξεις, για παράδειγμα, ο βαθμός ομοιότητας $s_1^{u,c}$ του x_u με τη λέξη w_1^c και ο βαθμός ομοιότητας $s_2^{u,c}$ του x_u με τη λέξη w_2^c , εκφράζονται από τα ασαφή σύνολα $A_{1,1}^c$ και $A_{1,2}^c$, αντίστοιχα, αλλά και τα δύο ασαφή σύνολα μπορεί να αντιπροσωπεύουν το λεκτικό "Χαμηλότερο". Η λογική πίσω από την προσέγγιση μείωσης φράσεων που εξετάζεται είναι ότι μόνο ένας από τους δύο βαθμούς ομοιότητας, $s_1^{u,c}$ και $s_2^{u,c}$, και τα αντίστοιχα ασαφή σύνολα, $A_{1,1}^c$ και $A_{1,2}^c$, επιλέγονται ως αρκετά για να αναπαραστήσουν τη "Χαμηλότερη" ομοιότητα του x_u με τις λέξεις της κλάσης c . Η διαδικασία μείωσης φράσεων δεν τροποποιεί τις τιμές του διανύσματος χαρακτηριστικών x_u , αλλά επιλέγει τους κατάλληλους βαθμούς ομοιότητας για να περιγράψει την ομοιότητα του x_u με μια κλάση c . Για να απλοποιηθεί η διαδικασία ταξινόμησης, προτείνεται μια μείωση του FSP, έτσι ώστε να επιλέγεται μόνο ένα από τα πολλαπλά ασαφή σύνολα για να αναπαραστήσει ένα συγκεκριμένο λεκτικό, με αντίστοιχη μείωση στον αριθμό των χρησιμοποιούμενων λέξεων. Ειδικά, για κάθε $l = 1'$ που μεγιστοποιεί την τιμή συμμετοχής $\mu_{l,m}^c$ ($s_m^{u,c}$) στην εξίσωση (5.5), το επιλεγμένο ασαφές σύνολο $A_{l',m}^c$ θα είναι αυτό που μεγιστοποιεί το γινόμενο $s_m^{u,c} \cdot \mu_{l',m}^c$ ($s_m^{u,c}$), για $m = 1, \dots, M^c$. Έτσι, προκύπτει το ακόλουθο μειωμένο FSP:

$$R^{u,c} = \{A_{l'}^c = A_{l',m}^c | m'' = \underset{\text{argmax}(s_m^{u,c} \cdot \mu_{l',m}^c(s_m^{u,c}))}{\forall l': A_{l',m}^c \in \theta^{u,c}}\} \quad (5.7)$$

Το FS $R^{u,c} \subseteq \theta^{u,c}$ που περιγράφεται από το (5.7) έχει μήκος $\Lambda(\theta^{u,c}) \leq L^C$. Αν και μπορεί να περιέχει λιγότερη πληροφορία από το $\theta^{u,c}$, μπορεί να είναι αρκετά εκφραστικό για την ταξινόμηση δεδομένων. Στα πλαίσια της παρούσας πτυχιακής εργασίας υλοποιήθηκε μια παραλλαγή στη μείωση των φράσεων. Στο μειωμένο FSP οι τιμές ομοιότητας και οι τιμές των συναρτήσεων συμμετοχής για κάθε $l = l'$ που μεγιστοποιεί την τιμή συμμετοχής $\mu_{l,m}^c(s_m^{u,c})$ στην εξίσωση (5.5) δίνονται από το μέσο όρο σύμφωνα με τη σχέση (**FSP-FEAT-SEL**, **FSP-GD-FEAT-SEL**):

$$s_{l,m}^{u,c} = \frac{s_{l,m}^{u,c} + s_{l',m}^{u,c}}{2}, \mu_{l,m}^{u,c} = \frac{\mu_{l,m}^{u,c} + \mu_{l',m}^{u,c}}{2} \quad (5.8)$$

Με αυτόν τον τρόπο, ο μηχανισμός μείωσης φράσεων που περιγράφεται καταφέρνει να απλοποιήσει τη μορφή των κανόνων τόσο σε όρους σύνθεσης όσο και σε σημασιολογία, καθώς οι κανόνες καταλήγουν με λιγότερα στοιχεία στο τμήμα της υπόθεσης, οι οποίοι είναι πιο εύκολα κατανοητοί από τους ανθρώπους. Αυτό επιβεβαιώνεται πειραματικά από τα αποτελέσματα που έχουν ληφθεί στο Κεφάλαιο 6. Αυτή η μείωση χαρακτηριστικών θα μπορούσε και να μην είχε πραγματοποιηθεί. Αποτελεί ένα κομμάτι το οποίο διερευνάται ώστε να λαμβάνονται υπόψιν όλα τα χαρακτηριστικά με κάποιο βαθμό το καθένα.



Εικόνα 15 Διαγραμματική απεικόνιση μείωσης φράσεων

5.1.4 Δημιουργία Κανόνων

Ένα FSP εξηγεί το γεγονός αν ένα εισερχόμενο διάνυσμα ανήκει σε μια κλάση, όσον αφορά την ομοιότητά του με ένα σύνολο λέξεων που χαρακτηρίζουν αυτήν την κλάση (Ορισμός 3). Ζεύγη FSP που περιγράφουν το αν ένα εισερχόμενο διάνυσμα ανήκει σε δύο διαφορετικές κλάσεις συνδυάζονται για να δημιουργήσουν κανόνες δυαδικής ταξινόμησης. Βάσει αυτών των κανόνων, ένα νέο διάνυσμα θα μπορεί να ταξινομηθεί σε μία από αυτές τις κλάσεις, ανάλογα με την ομοιότητά του με τις λέξεις της κλάσης που παρουσιάζει μεγαλύτερη ομοιότητα. Έτσι, η συλλογιστική για την ταξινόμηση μπορεί να πραγματοποιηθεί με το συνδυασμό τέτοιων κανόνων. Η συζευκτική ταξινόμηση εμπεριέχει την τετραγωνική αύξηση των κανόνων με τον αριθμό των κλάσεων. Ο ακόλουθος ορισμός περιγράφει τη διαδικασία δημιουργίας κανόνων.

Ορισμός 5: Δεδομένου ενός εισερχόμενου διανύσματος x^u και ενός FSP μοντέλου F^c , η ταξινόμηση του x^u σε μια κλάση c έναντι μιας άλλης κλάσης $c' \neq c$, όπου $c, c' = 1, 2, \dots, C$ περιγράφεται από έναν κανόνα FSP Rule $^{c,c'}(x_u)$, που παράγεται από ένα ζεύγος FSPs $(\theta^{u,c}, \theta^{u,c'})$ βασιζόμενων στο $F(c)$ ως εξής:

Αν η ομοιότητα του x^u [με την λέξη w_m^c είναι $A_{l',m}^c$], με $l = \operatorname{argmax}(\mu_{l',m}^c(s_m^{u,c}))$, για όλα τα m : $w_m^c \in M^C$ και η ομοιότητα του x_u [με την λέξη $w_m^{c'}$ είναι $A_{l',m}^{c'}$], με $l = \operatorname{argmax}(\mu_{l',m}^{c'}(s_m^{u,c'}))$, για όλα τα m : $w_m^c \in M^C$, τότε το x_u ταξινομείται στην κλάση c με το $y^{c,c'}(x_u) > 0$, όπου:

$$y^{c,c'}(x_u) = \sigma(\theta^{u,c}, q^c) - \sigma(\theta^{u,c'}, q^{c'}) \quad (5.9)$$

αναπαριστά μία μέτρηση σχετικής βεβαιότητας αυτού του κανόνα με τιμές που κυμαίνονται στο διάστημα $[-1,1]$. Οι φράσεις που βρίσκονται μέσα σε αγκύλες επαναλαμβάνονται για όλα τα m : $w_m^c \in M^C$ και για όλα τα m : $w_m^{c'} \in M^{C'}$. Αν $y^{c,c'}(x_u) \leq 0$, τότε το x_u δεν ταξινομείται στην κλάση c . Ο αριθμός των αξιολογήσεων ομοιότητας που πραγματοποιούνται ανάμεσα στο x^u και τις λέξεις $w_i = c, c'$ στο μέρος της υπόθεσης, καθορίζει το μήκος του κανόνα FSP. Στην περίπτωση μείωσης φράσης, χρησιμοποιούνται τα $(R^{u,c}, R^{u,c'})$ αντί για τα $(\theta^{u,c}, \theta^{u,c'})$ σύμφωνα με την εξίσωση (5.7). Τελικά, το διάνυσμα χαρακτηριστικών της εισερχόμενης εικόνας ταξινομείται στην κλάση που έχει το μέγιστο μέσο σταθμισμένο άθροισμα ομοιότητας, με μια σχετική βεβαιότητα.

$$y^{c,c'}(x_u) = \sigma(\theta^{u,c}, q^c) - \sigma(\theta^{u,c'}, q^{c'})$$

Ορισμός 6: Δύο κανόνες $R^{u,c}(x^u)$ και $R^{u,c'}(x^u)$ είναι ίσοι αν και μόνο αν αποτελούνται από δύο ίσα ζεύγη FSPs.

Μόλις ένας κανόνας δημιουργηθεί από ένα διάνυσμα εισόδου x^u , μπορεί να εφαρμοστεί σε οποιοδήποτε άλλο διάνυσμα εισόδου. Ένας κανόνας που περιγράφεται από τον ορισμό 5 περιλαμβάνει προσαρμόσιμες παραμέτρους στα διανύσματα βαρών q^c και $q^{c'}$. Αυτό σημαίνει πως αν η συμμετοχή ενός ή περισσότερων διανυσμάτων εισόδου σε μια κλάση είναι εκ των προτέρων γνωστή, τα βάρη του κανόνα μπορούν να προσαρμοστούν έτσι ώστε ο κανόνας να παράγει πιο ακριβή αποτελέσματα ταξινόμησης.

5.1.5 Επαλήθευση Μοντέλου και Δημιουργία Βάσης Κανόνων

Έστω το σύνολο $V^c = \{v_1^c, v_2^c, \dots, v_{K_c}^c\}$ ως το σύνολο των K^c διανυσμάτων επαλήθευσης που ανήκουν στην κλάση $c = 1, 2, \dots, C$. Το πλήρες σύνολο επαλήθευσης θα είναι $V = \cup_{c=1}^C V^c$. Δεδομένου ενός μοντέλου FSP $F(C)$ που κατασκευάστηκε κατά την φάση εκπαίδευσης, δημιουργείται ένα σύνολο FSP από όλα τα διανύσματα του συνόλου επαλήθευσης και συνδυάζονται διαδοχικά ώστε να δημιουργηθεί ένα σύνολο μοναδικών κανόνων ταξινόμησης FSP. Λαμβάνοντας υπόψη ότι η κλάση που ανήκουν τα διανύσματα επαλήθευσης θεωρείται εκ των προτέρων γνωστή, δημιουργείται και προσαρμόζεται μια βάση κανόνων κατά τη φάση

επαλήθευσης του προτεινόμενου μοντέλου FSP. Συγκεκριμένα, για κάθε διάνυσμα επαλήθευσης v_k^c που ανήκει στην κλάση c , δημιουργούνται συνολικά $C-1$ δυνατοί κανόνες ταξινόμησης $\text{Rule}^{c,c'}(v_k^c)$. Η παραγόμενη βάση κανόνων (Rules Base-RB) για τα διανύσματα της κλάσης c αναπαρίσταται ως ένα σύνολο.

$$RB^c = \{ \text{Rule}^{c,c'}(v_k^c) \mid \forall v_k^c \in V^c, K=1, \dots, K^c, c \neq c', c'=1, \dots, C \} \quad (5.10)$$

Κάθε κανόνας στη βάση κανόνων RB προσαρμόζεται έτσι ώστε κάθε διάνυσμα επαλήθευσης να ταξινομείται στην κλάση στην οποία ανήκει. Για να επιτευχθεί αυτό, τα διανύσματα βαρών q^c και $q^{c'}$ κάθε κανόνα στη βάση κανόνων (για $c = 1, 2, \dots, C$) προσαρμόζονται με τη μέθοδο βαθμωτής Κατάδυσης (Gradient Descent - GD) [102], με στόχο να μεγιστοποιηθεί η σχετική βεβαιότητα για τη σωστή ταξινόμηση των διανυσμάτων που παράγουν τον ίδιο κανόνα FSP $\text{Rule}^{c,c'}$.

Το σφάλμα ως προς τη γνωστή εκ των προτέρων κλάση (στόχος) που ανήκει το διάνυσμα v_k ορίζεται ως μια συνάρτηση της σχετικής βεβαιότητας ως εξής:

$$\varepsilon(v_k) = u(y^{c,c'}(v_k)) * y^{c,c'}(v_k) - y^t(v_k) \quad (5.11)$$

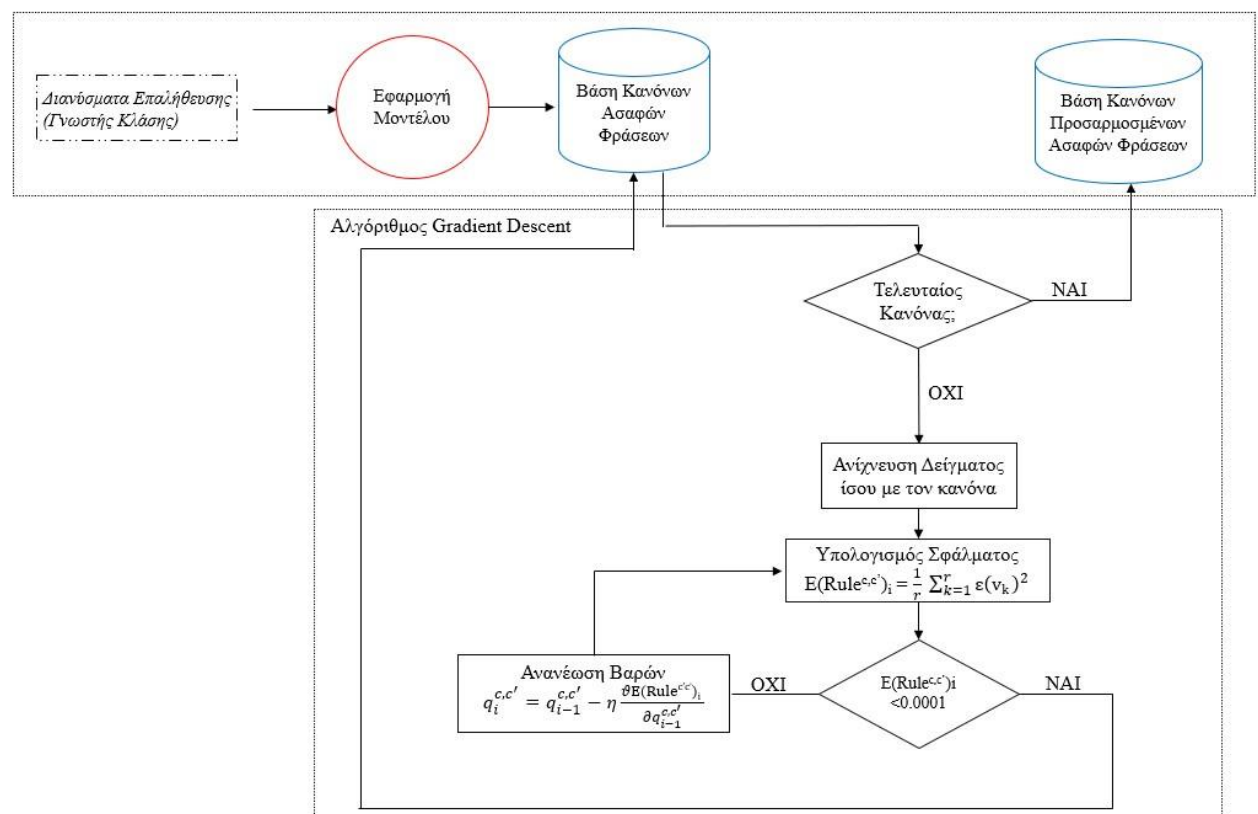
Όπου $y^t(v_k) = 1$, αν το v_k ανήκει στην κλάση c αλλιώς $y^t(v_k) = 0$ και $u(.)$ η βηματική συνάρτηση. Έστω r ο αριθμός των διανυσμάτων που δημιουργούν τους ίδιους κανόνες FSP $\text{Rule}^{c,c'}$. Τότε, το συνολικό σφάλμα του $\text{Rule}^{c,c'}$ για αυτό το σύνολο διανυσμάτων μετά από i επαναλήψεις, θα είναι:

$$E(\text{Rule}^{c,c'})_i = \frac{1}{r} \sum_{k=1}^r \varepsilon(v_k)^2 \quad (5.12)$$

Πριν από την εφαρμογή του αλγορίθμου Gradient Descent (GD), κάθε διάνυσμα βαρών $q = (q_1, \dots, q_\lambda)$ του κανόνα αρχικοποιείται έτσι ώστε τα βάρη που αντιστοιχούν σε υψηλότερο βαθμό ομοιότητας να έχουν υψηλότερες τιμές βάρους. Η αντίστοιχη εξίσωση του Gradient Descent που εφαρμόζεται επαναληπτικά ανά κανόνα για την ελαχιστοποίηση του σφάλματος που υπολογίζεται από την εξίσωση (5.12) είναι:

$$q_i^{c,c'} = q_{i-1}^{c,c'} - \eta \frac{\partial E(\text{Rule}^{c,c'})_i}{\partial q_{i-1}^{c,c'}} \quad (5.13)$$

Όπου $q_i^{c,c'}$ είναι το διάνυσμα βαρών που προκύπτει σε μια επανάληψη i από τον συνδυασμό των δύο διανυσμάτων βαρών του κανόνα $E(\text{Rule}^{c,c'})$, δηλαδή $q^{c,c'} = (q^c, q^{c'})$, και το $\eta \in [0,1]$ αντιπροσωπεύει το ρυθμό μάθησης (learning rate). Η προσαρμογή κάθε κανόνα και συνεπώς το διάνυσμα του βάρους του τερματίζεται όταν συγκλίνει ή όταν επιτυγχάνεται επαρκής ελαχιστοποίηση του σφάλματος.



Εικόνα 16 Διαγραμματική απεικόνιση για την κατασκευή της ανανεωμένης βάσης κανόνων.

5.1.6 Ταξινόμηση ενός διανύσματος άγνωστης κλάσης

Δοθέντος, λοιπόν, ενός διανύσματος εισόδου x^u άγνωστης κλάσης, του μοντέλου FSP $F(C)$ και της προσαρμοσμένης με τα βάρη Βάσης Κανόνων (RB), η ταξινόμηση του x^u στη φάση της δοκιμής (test) βασίζεται στο FSP που θα δημιουργηθεί χρησιμοποιώντας το μοντέλο FSP. Το μοντέλο FSP παρέχει μια εξήγηση γιατί το διάνυσμα εισόδου x^u ανήκει σε μια συγκεκριμένη κλάση c ($c = 1, \dots$,

C) και με ποια βεβαιότητα. Η ταξινόμηση του διανύσματος εισόδου (x^u) βασίζεται στην εξίσωση (5.9), η οποία είναι ο βασικός κανόνας που καθορίζει εάν το x^u ανήκει σε μια κλάση c ($c = 1, \dots, C$). Όπως αναφέρθηκε προηγουμένως, ο βασικός αυτός κανόνας δημιουργείται από όλα τα ζεύγη των FSPs ($\theta^{u,c}, \theta^{u,c'}$) για όλες τις κλάσεις C βασιζόμενο στο FSP $F(c)$. Στη δυαδική ταξινόμηση, το x^u μπορεί να ταξινομηθεί με βάση τη μέγιστη σχετική βεβαιότητα $y^{c,c'}(x^u)$. Στην ταξινόμηση πολλαπλών κλάσεων, το x^u μπορεί να ταξινομηθεί στην κλάση με τον μέγιστο μέσο όρο σχετικής βεβαιότητας από όλα τα ζεύγη των FSPs ($\theta^{u,c}, \theta^{u,c'}$)

$$y^c(x^u) = \frac{1}{C-1} \sum_{\substack{c'=1, \\ c \neq c'}}^C y^{c,c'}(x^u) \quad (5.14)$$

Με στόχο να βελτιωθεί η διαδικασία λήψης αποφάσεων για την ταξινόμηση του διανύσματος εισόδου (x^u), ακολουθείται μια προσέγγιση συγχώνευσης, λαμβάνοντας υπόψη την γνώση από τον βασικό κανόνα και τη γνώση που αποθηκεύεται στη βάση κανόνων RB. Τα FSPs που παράγονται από το διάνυσμα x^u χρησιμοποιούνται ως ερωτήματα για την ανάκτηση ίσων κανόνων από την RB (η ισότητα ορίζεται στον Ορισμό 6).

Ορισμός 7: Δεδομένου ενός ζεύγους FSPs ($\theta^{u,c}, \theta^{u,c'}$), το οποίο παράγεται από ένα διάνυσμα εισόδου x^u , η καταλληλότητα ενός κανόνα στη Βάση Κανόνων για την ταξινόμηση του x^u σε μια κλάση c ($c = 1, \dots, C$) ορίζεται ως:

$$\mathcal{F}^{c,c'} = \prod_{\forall l': A_{l',m}^c} \mu_{l',m'}^c(s_{l',m'}^c) \cdot \prod_{\forall l': A_{l',m}^{c'}} \mu_{l',m'}^{c'}(s_{l',m'}^{c'}) \quad (5.15)$$

Όπου $\mathcal{F}^{c,c'} \geq 0$. Αν $\mathcal{F}^{c,c'} = 0$ τότε ο κανόνας δεν είναι κατάλληλος για την ταξινόμηση του x^u . Εξετάζοντας τον Ορισμό 7, ένας κανόνας ανακτάται από τη βάση κανόνων RB για την ταξινόμηση του διανύσματος εισόδου (x^u) σε μια κλάση c , εάν η τιμή της μεταβλητής $\mathcal{F}^{c,c'}$ είναι μεγαλύτερη του μηδενός (0). Στη συνέχεια, η σχετική βεβαιότητα του ανακτηθέντος κανόνα θα είναι $y^{R,c,c'}(x^u)$, και η σχετική βεβαιότητα του βασικού κανόνα θα είναι $y^{c,c'}(x^u)$. Σε αυτό το σημείο πρέπει να σημειωθεί ότι υπάρχει πιθανότητα κανένας κανόνας να μην είναι κατάλληλος για την ταξινόμηση του x^u στην κλάση c . Σε αυτήν την περίπτωση μόνο ο βασικός κανόνας λαμβάνεται υπόψιν. Η συνολική σχετική βεβαιότητα λαμβάνεται από την ακόλουθη ισότητα.

$$y^{c,c'}(x^u) = \begin{cases} \frac{y^{c,c'}(x^u) + y^{R^{c,c'}}(x^u)}{2}, & \mathcal{F}^{c,c'} > 0 \\ y^{c,c'}(x^u), & \mathcal{F}^{c,c'} = 0 \end{cases} \quad (5.16)$$

Στην περίπτωση που ισχύει $\mathcal{F}^{c,c'} > 0$, η σχετική βεβαιότητα υπολογίζεται ως ο μέσος όρος της σχετικής βεβαιότητας του βασικού κανόνα και του κανόνα που ανακτήθηκε από τη βάση κανόνων. Η λογική πίσω από αυτόν τον μέσο όρο είναι να συγκεντρωθούν οι δύο σχετικές βεβαιότητες, θεωρώντας τις εξίσου σημαντικές. Έτσι, μια πιθανώς χαμηλή σχετική βεβαιότητα λόγω του βασικού κανόνα μπορεί να βελτιωθεί από μια υψηλότερη σχετική βεβαιότητα από έναν κανόνα της βάσης κανόνων και αντίστροφα.

Εάν η βάση κανόνων είναι ανεπαρκής, δηλαδή υπάρχουν δείγματα για τα οποία δεν μπορούν να βρεθούν ίσοι κανόνες, η εξίσωση (5.16) διασφαλίζει την ύπαρξη τουλάχιστον ενός κανόνα (τον βασικό) που θα χρησιμοποιηθεί για τον συλλογισμό. Οι κανόνες συμβάλλουν στην ερμηνευσιμότητα ενός αποτελέσματος ταξινόμησης από έναν ταξινομητή, στη μορφή φράσεων που είναι κατανοητές από τους ανθρώπους.

5.1.7 Ερμηνεία της Ταξινόμησης

Το αποτέλεσμα ταξινόμησης του προτεινόμενου μοντέλου FSPs μπορεί να εξηγηθεί μέσω της ανάλυσης του υποθετικού μέρους των δημιουργημένων κανόνων. Για κάθε διάνυσμα εισόδου x^u που ταξινομείται στην κλάση c έναντι μιας άλλης κλάσης c' , δημιουργείται μια εξήγηση από το ζεύγος των FSPs $(\theta^{u,c}, \theta^{u,c'})$ από τα οποία προέρχεται κάθε κανόνας, όπως:

"Το x^u ταξινομείται στην κλάση c και όχι στην κλάση c' , διότι έχει περισσότερες παρόμοιες τιμές χαρακτηριστικών με τις λέξεις της κλάσης c από ό,τι με τις λέξεις της κλάσης c' ."

Αυτή η εξήγηση προκύπτει από το γεγονός ότι η σχετική βεβαιότητα $y^{c,c'}(x^u) > 0$, που σημαίνει ότι το x^u εμφανίζει υψηλότερο συνολικό βαθμό ομοιότητας με τις λέξεις της κλάσης c . Μια πιο λεπτομερής εξήγηση μπορεί να βασίζεται στις ομοιότητες του x^u με κάθε λέξη, δηλαδή:

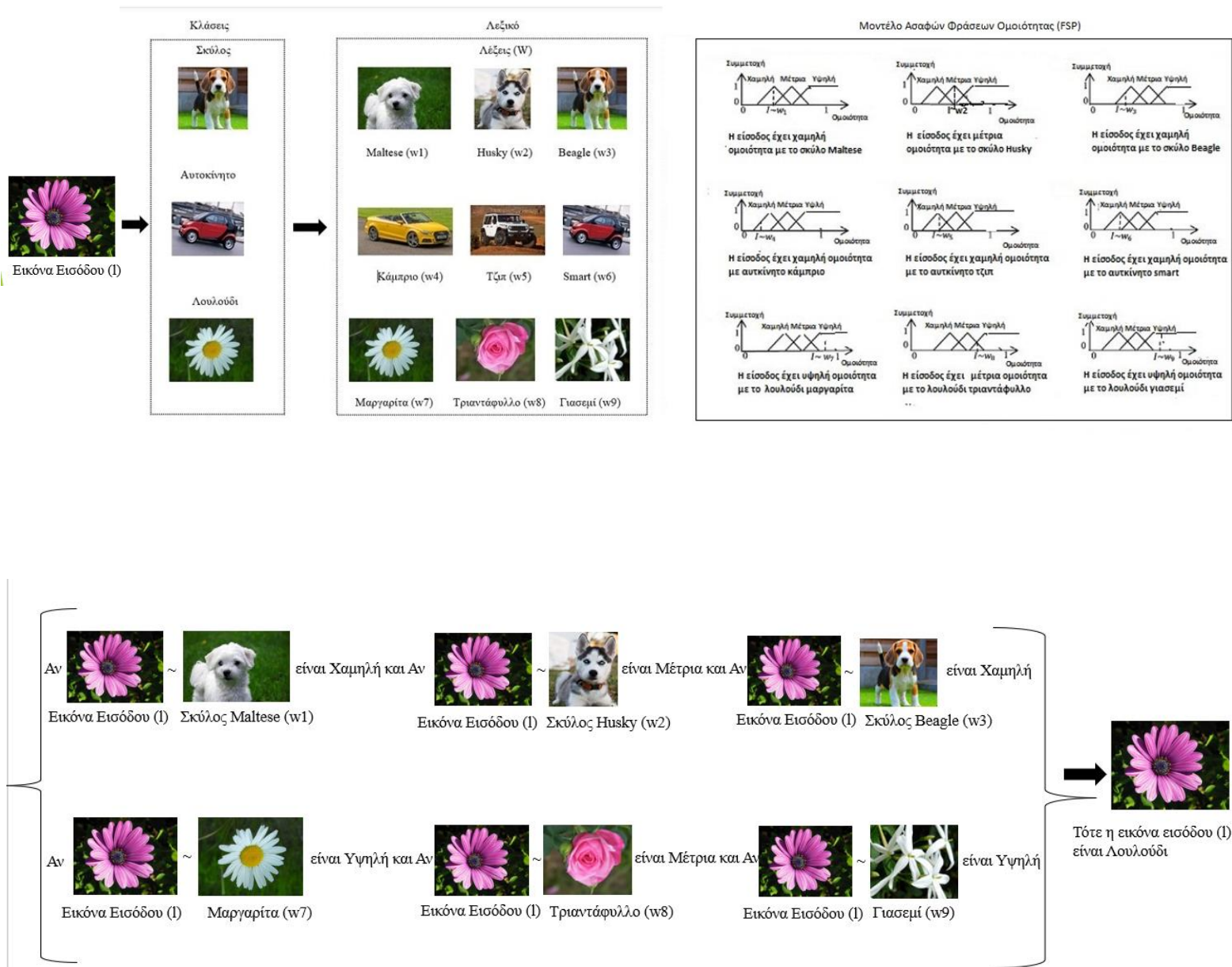
"Το x^u ταξινομείται στην κλάση c επειδή έχει $A_{l',1}^c \in \theta^{u,c}$ και $A_{l',2}^c \in \theta^{u,c}, \dots$ για $l' = \operatorname{argmax}_{l \in [1, Lc]} (\mu_{l,m}^c(s_m^{u,c}))$, $\forall m: w_m^c \in M^c$ ομοιότητα με τις λέξεις $w_m^c, m = 1, \dots, M^c$ της κλάσης c , και το x^u έχει $A_{l',1}^{c'} \in \theta^{u,c'}$ και $A_{l',2}^{c'} \in \theta^{u,c'}, \dots$ για $l' = \operatorname{argmax}_{l \in [1, Lc']} (\mu_{l,m}^{c'}(s_m^{u,c'}))$, $\forall m: w_m^{c'} \in M^{c'}$ ομοιότητα με τις λέξεις $w_m^{c'}, m = 1, \dots, M^{c'}$ της κλάσης c' ."

Στην περίπτωση ταξινόμησης πολλαπλών κλάσεων, η ερμηνεία της ταξινόμησης του x^u στην κλάση c έναντι άλλων κλάσεων, θα βασίζεται στο (5.14), το οποίο λειτουργεί ως τελεστής συγκέντρωσης πάνω σε όλους τους κανόνες δυαδικής ταξινόμησης που λαμβάνονται υπόψη, και μπορεί να εκφραστεί ως:

"Το x^u ταξινομείται στην κλάση c και όχι σε καμία άλλη κλάση επειδή έχει περισσότερες παρόμοιες τιμές χαρακτηριστικών με τις λέξεις της κλάσης c από ό,τι με τις λέξεις κάθε άλλης κλάσης."

Ένα παράδειγμα οπτικοποίησης της μεθοδολογίας παρουσιάζεται παρακάτω ώστε να γίνει πιο κατανοητό. Έστω ότι υπάρχουν τρεις κλάσεις (σκύλος, αυτοκίνητο, λουλούδι). Δημιουργούνται οι λέξεις για κάθε κλάση (w_1 - w_9) οι οποίες απαρτίζουν το λεξικό. Στη συνέχεια γίνεται ο χαρακτηρισμός των λέξεων (Χαμηλή/Μέτρια/Υψηλή Ομοιότητα) με την εικόνα εισόδου χρησιμοποιώντας FSPs. Στη δεύτερη εικόνα φαίνονται οι ασαφείς κανόνες με το σύμβολο ($\sim$) να απεικονίζει την ομοιότητα. Η ταξινόμηση μεταξύ δύο κλάσεων φαίνεται σε αυτό το σχήμα όπου σύμφωνα και με την παραπάνω φράση μπορεί να εκφραστεί ως εξής:

Η εικόνα εισόδου ταξινομείται στην κλάση λουλούδι και όχι στην κλάση σκύλος καθώς έχει περισσότερες όμοιες τιμές με τις λέξεις μαργαρίτα, τριαντάφυλλο και γιασεμί από ότι με τις λέξεις Beagle, Maltese, Husky



Εικόνα 17 Παράδειγμα οπτικοποίησης της μεθοδολογίας.

5.2 Παραδείγματα Βημάτων Υλοποίησης της Μεθόδου

Σε αυτήν την ενότητα παρουσιάζονται πρακτικά παραδείγματα όπως ακριβώς υλοποιήθηκαν σύμφωνα με τα βήματα που αναλύθηκαν παραπάνω. Στόχος είναι να γίνει κατανοητό το κάθε βήμα της μεθόδου που αναλύεται παραπάνω καθώς και τα εξαγόμενα διανύσματα του κάθε βήματος. Για την πρώτη πειραματική υλοποίηση χρησιμοποιήθηκε ένα σετ δεδομένων με ασθενείς καρδιακής ανεπάρκειας που δημοσιεύθηκε από Rousseau et al [105]. Σε επόμενο κεφάλαιο αναλύονται περαιτέρω τα πειραματικά αποτελέσματα της μεθόδου και με διαφορετικά σύνολα δεδομένων.

Φορτώνεται λοιπόν αρχικά το σύνολο δεδομένων με τους ασθενείς καρδιακής ανεπάρκειας το οποίο έχει την παρακάτω μορφή όπως φαίνεται από την Εικόνα 18. Σε αυτό το σημείο να σημειωθεί πως για την υλοποίηση και τη διαχείριση των τύπων δεδομένων χρησιμοποιήθηκαν κατά κύριο λόγο τα δημόσια διαθέσιμα πακέτα της Python τα *NumPy* [123] & *Pandas* [124].

	Sbp	Tobacco	Ldl	Adiposity	Typea	Obesity	Alcohol	Age	DEATH_EVENT
0	160	12.00	5.73	23.11	49	25.30	97.20	52	1
1	144	0.01	4.41	28.61	55	28.87	2.06	63	1
2	118	0.08	3.48	32.28	52	29.14	3.81	46	0
3	170	7.50	6.41	38.03	51	31.99	24.26	58	1
4	134	13.60	3.50	27.78	60	25.99	57.34	49	1
...	...	...	...	...	...	...	...	...	...
457	214	0.40	5.98	31.72	64	28.45	0.00	58	0
458	182	4.20	4.41	32.10	52	28.61	18.72	52	1
459	108	3.00	1.59	15.23	40	20.09	26.64	55	0
460	118	5.40	11.61	30.79	64	27.35	23.97	40	0
461	132	0.00	4.82	33.41	62	14.70	0.00	46	1

462 rows × 9 columns

Εικόνα 18 Αρχική μορφή συνόλου δεδομένων καρδιακής ανεπάρκειας

Σύμφωνα με το βήμα 1 πραγματοποιείται αυτός ο χωρισμός στα δεδομένα της κλάσης 1 και στα δεδομένα της κλάσης 0. Εφόσον διαχωριστούν οι κλάσεις, τα δεδομένα χωρίζονται σε σύνολα εκπαίδευσης, δοκιμής και επαλήθευσης με τη μέθοδο 10-fold cross validation και στη συνέχεια κανονικοποιούνται σύμφωνα με την κανονικοποίηση z-score. Έπειτα εφαρμόζεται ο kmeans και στα δύο σύνολα για να εξαχθούν τα διανύσματα των κέντρων σύμφωνα και με την υποενότητα 5.1.1. Μετά την εφαρμογή του αλγορίθμου αυτού τα κέντρα που παράγονται έχουν την παρακάτω μορφή έπειτα από ομαδοποίηση με 1 συστάδα. Ο αριθμός των συστάδων μπορεί να ποικίλει και να

είναι διαφορετικός για κάθε κλάση. Παράδειγμα φαίνεται παρακάτω στην Εικόνα 19 για την ομαδοποίηση της κλάσης 0. Τα βήματα που προαναφέρθηκαν υλοποιήθηκαν με τη βοήθεια του δημόσια διαθέσιμου πακέτου της Python το *scikit-learn* [120].

Δημιουργία Λεξικού με kmeans

```
from sklearn.cluster import KMeans
if estimator == 'kmeans':

    kmeans = KMeans(init="k-means++", n_clusters=n_cluster[i], n_init=1,
max_iter=300, random_state=None)
    kmeans.fit(df)
    list_centers=kmeans.cluster_centers_
```

Δημιουργία Λεξικού με GMM

```
from sklearn.mixture import GaussianMixture

gmm = GaussianMixture(n_components=n_cluster[i], n_init=1,max_iter=300,
random_state=None)
gmm.fit(df)
list_centers = gmm.means_
```

```
[array([[ -0.0705341 , -0.17804597, -0.1600679 , -0.14149738, -0.21699992,
        -0.06157867, -0.05365518, -0.00790376]])]
```

Εικόνα 19 1 8-διαστατό κέντρο συστάδας.

Συνεχίζοντας με το βήμα 2 αποκτώνται τα διανύσματα των ομοιοτήτων σύμφωνα με τον τύπο (5.1). Για όλα τα χαρακτηριστικά του αρχικού συνόλου παράγονται οι ομοιότητες με τα αντίστοιχα κέντρα της κάθε κλάσης. Παρακάτω φαίνονται και τα σημεία κώδικα όπου χρησιμοποιείται η συνάρτηση 5.1 καθώς και οι υπόλοιπες ομοιότητες που υλοποιήθηκαν. Σε μία τελική Εικόνα 20 έπειτα από τα κομμάτια υλοποίησης της κάθε εξίσωσης από την ενότητα 5.1.1 φαίνονται ενδεικτικά τα εξαγόμενα διανύσματα των ομοιοτήτων.

Υλοποίηση Εξίσωσης 5.1 με Ευκλείδεια Απόσταση

```
import numpy as np

for a, data in enumerate(all_centers):
    for center in range(len(all_centers[a])):
        temp=row - all_centers[a][center]
        sum_sq = np.dot(temp.T, temp)
        arr=np.sqrt(sum_sq)
        total_dist+=arr

for v_center, data in enumerate(all_centers):
    for cen in range(len(all_centers[v_center])):
        temp = row - all_centers[v_center][cen]
        sum_sq = np.dot(temp.T, temp)
        s1 = np.sqrt(sum_sq)
        s = 1 - (s1 / total_dist)
        distances.append(s)
```

Υλοποίηση Εξίσωσης 5.2 Ομοιότητας Συνημίτονου

```
import numpy as np

def cosine_similarity(x, y):
    dot_product = np.dot(x, y)
    norm_x = np.linalg.norm(x)
    norm_y = np.linalg.norm(y)

    if norm_x == 0 or norm_y == 0:
        similarity = np.nan
    else:
        similarity = dot_product / (norm_x * norm_y)

    return similarity
```

Υλοποίηση απόστασης Manhattan για εφαρμογή στην Εξίσωση 5.1

```
import numpy as np

def manhattan_distance(x, y):
    return np.sum(np.abs(x - y))
```

Υλοποίηση Απόστασης Mahalanobis για εφαρμογή στην Εξίσωση 5.1

```
import scipy as sp
import numpy as np
def mahalanobis_distance(x, y, covariance_matrix):
    diff = x - y
    inv_cov_matrix = sp.linalg.inv(covariance_matrix)
    left_term = np.dot(diff, inv_cov_matrix)
    mahal = np.dot(left_term, diff.T)
    return mahal
```

```
[array([[0.60631972, 0.39368028],
        [0.60571353, 0.39428647],
        [0.5052141 , 0.4947859 ],
        [0.51593278, 0.48406722],
        [0.4416818 , 0.5583182 ],
        [0.42634673, 0.57365327],
        [0.52221937, 0.47778063],
        [0.59711083, 0.40288917],
        [0.44979598, 0.55020402],
        [0.45344414, 0.54655586],
        [0.41374572, 0.58625428],
        [0.57781525, 0.42218475],
        [0.60179655, 0.39820345],
        [0.43977861, 0.56022139],
        [0.49580349, 0.50419651],
        [0.55731026, 0.44268974],
        [0.42007925, 0.57992075],
        [0.45177539, 0.54822461]])]
```

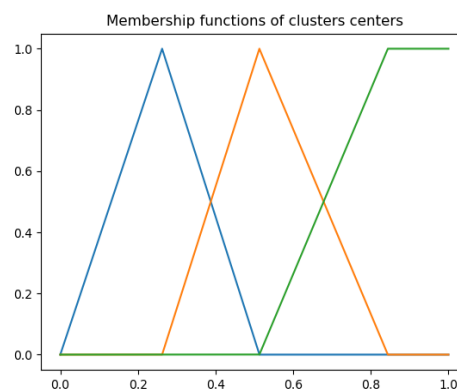
Εικόνα 20 Παραγόμενα διανύσματα ομοιοτήτων.

Όπου η πρώτη στήλη αφορά την ομοιότητα από τα κέντρα της κλάσης 1 και η άλλη τις ομοιότητες από τα κέντρα της κλάσης 0. Παρατηρώντας τις ομοιότητες, είναι λογικό πως η πρώτη στήλη έχει ελαφρώς μεγαλύτερη τιμή από την άλλη στήλη. Αυτό γιατί είναι λογικό τα δεδομένα της κλάσης 0 να έχουν μεγαλύτερη ομοιότητα από τα κέντρα της κλάσης 0 συγκριτικά με τα κέντρα της κλάσης 1 όπως αποκτήθηκαν στο προηγούμενο βήμα.

Ακολουθεί το στάδιο της ασαφοποίησης για τη δημιουργία των ασαφών συνόλων. Πριν από αυτό βέβαια τα δεδομένα κανονικοποιούνται σύμφωνα με τη μέθοδο MinMax ώστε να υπάρχει μια κλίμακα στις τιμές. Με αυτή τη μέθοδο η μέγιστη τιμή θα είναι 1 και η ελάχιστη 0. Χρησιμοποιήθηκε αριθμός συστάδων εδώ 1 για την εξαγωγή των κέντρων και 3 για τη δημιουργία των ασαφών συνόλων. Ενδεικτικά παρουσιάζονται οι τριγωνικές συναρτήσεις συμμετοχής της κλάσης 0 από τα κέντρα της κλάσης 0 όπως παράχθηκαν από αυτό το στάδιο. Για τη γραφική απεικόνιση των συναρτήσεων συμμετοχής όπως φαίνεται στην Εικόνα 21 χρησιμοποιήθηκε το δημόσια διαθέσιμο πακέτο της Python το *matplotlib* [122].

Υλοποίηση Εξίσωσης 3.3

```
def gaussian(x, m, sigma):  
    return np.exp(-((x - m) ** 2) / (sigma ** 2))
```



Εικόνα 21 Συναρτήσεις συμμετοχής για την ανάπτυξη των ασαφών συνόλων.

Όπου η κορυφή της κάθε τριγωνικής συνάρτησης αντιστοιχεί σε κέντρο συστάδας που προέρχεται από την ομαδοποίηση. Η μπλε συνάρτηση συμμετοχής αντικατοπτρίζει το λεκτικό «χαμηλή» ή τον αριθμό «0», η πορτοκαλί το λεκτικό «μεσαία» ή τον αριθμό «1» και η πράσινη το λεκτικό «υψηλή» ή τον αριθμό «2» κ.ο.κ. αν υπάρχουν περισσότερες συναρτήσεις. Για τη δημιουργία των συναρτήσεων συμμετοχής χρησιμοποιήθηκε το δημόσια διαθέσιμο πακέτο της Python το *scikit-fuzzy* [121].

Σύμφωνα με το επόμενο βήμα 5 για όλα τα δεδομένα παράγονται φράσεις της μορφής [ομοιότητα, τιμή συνάρτησης συμμετοχής, βαθμός] σύμφωνα με την εξίσωση 5.5 της υποενότητας 5.1.2. Ενδεικτικά παρακάτω φαίνονται τρία διανύσματα της κλάσης 0 όπως ακριβώς εξήχθησαν.

```
[[[0.5733578995616223, 0.9252434696573346, 1.0],  
[0.4266421004383775, 0.8186195961614644, 1.0]],  
[[0.7960183761525698, 0.9173255518372028, 2.0],  
[0.20398162384743015, 0.9335969806423562, 0.0]],  
[[0.47029021164603213, 0.6453357849493756, 1.0],  
[0.5297097883539673, 0.7524170970717389, 1.0]],  
[[0.5625853436907955, 0.9740363348792298, 1.0],  
[0.43741465630920434, 0.86098585720356, 1.0]]]
```

Εικόνα 22 Τριπλέτες εξαγόμενες από το βήμα 5.

Με αφορμή την Εικόνα 22 παρατηρείται πως στο πρώτο διάνυσμα υπάρχουν δύο τριπλέτες με τα ίδια λεκτικά («1» ή «μεσαία»). Ένα από τα δύο θα πρέπει να διαγραφεί σύμφωνα με την πράξη της υποενότητας 5.1.4 και τελικά να δημιουργηθεί ένα μειωμένο διάνυσμα σύμφωνα με την εξίσωση 5.6. Παρακάτω φαίνεται η μορφή των ίδιων φράσεων με την Εικόνα 23 έπειτα από τη μείωση τους. Στο συγκεκριμένο βήμα υλοποίησης προστέθηκε επίσης και μία ακόμη δυαδική τιμή στις φράσεις ("c1" ή "c0"). Η μεταβλητή αυτή αποδεικνύει από ποια κλάση έχει προέλθει η συγκεκριμένη φράση μετά τη διαγραφή.

Στη συνέχεια προστίθενται τα βάρη. Στη συγκεκριμένη περίπτωση των τριών λεκτικών το 0 θα πάρει τιμή 1/3, το 1 2/3 και το 2 1 όπως φαίνεται παρακάτω.

```
[[[0.5733578995616223, 0.9252434696573346, 1.0, 'c0', 0.6666666666666666],  
[0.4266421004383775, 0.8186195961614644, 1.0, 'c1', 0.6666666666666666]],  
[[0.7960183761525698, 0.9173255518372028, 2.0, 'c0', 1.0],  
[0.20398162384743015, 0.9335969806423562, 0.0, 'c1', 0.3333333333333333]],  
[[0.47029021164603213, 0.6453357849493756, 1.0, 'c0', 0.6666666666666666],  
[0.5297097883539673, 0.7524170970717389, 1.0, 'c1', 0.6666666666666666]],  
[[0.5625853436907955, 0.9740363348792298, 1.0, 'c0', 0.6666666666666666],  
[0.43741465630920434, 0.86098585720356, 1.0, 'c1', 0.6666666666666666]]]
```

Εικόνα 23 Μορφή διανυσμάτων μετά τη φάση της πρόσθεσης βαρών και τη μείωση των φράσεων.

Τέλος για τη διαδικασία της ταξινόμησης σε μετρική αποτελέσματος χρησιμοποιείται η ακρίβεια. Χρησιμοποιείται ο τύπος 5.9 της υποενότητας 5.1.4. Αφού συλλεχθούν τα αληθώς θετικά (TP), τα αληθώς αρνητικά (TN), τα ψευδώς θετικά (FP) και τα ψευδώς αρνητικά (FN) δείγματα υπολογίζεται το ποσοστό της ακρίβειας ταξινόμησης σύμφωνα με τον τύπο:

$$\text{Ακρίβεια} = (TP+TN)/(TP+TN+FN+FP) \quad (5.17)$$

Για το συγκεκριμένο πείραμα που ακολουθείται σαν παράδειγμα ο πίνακας συνάφειας ([[TP, FP], [FN, TN]]) καθώς και η ακρίβεια που εξάγεται για το σύνολο δοκιμής δίνονται παρακάτω:

```
[0.723404255319149] [0.8407258064516129] [[13, 10], [3, 21]]
```

Με σκοπό να ταξινομηθεί ένα τυχαίο δείγμα, όπως προαναφέρθηκε, δημιουργείται η βάση κανόνων σύμφωνα με την κλάση που ανήκει πραγματικά κάθε κανόνας σύμφωνα με τον τύπο 5.10 της υποενότητας 5.1.4. Η βάση αυτή έπειτα από την αφαίρεση των διπλότυπων έχει την μορφή που φαίνεται στην Εικόνα 24.

```
([['if', [0.0, 'c0', 0.333], 'and', [1.0, 'c1', 0.666], 'then Class0'],  
['if', [0.0, 'c0', 0.333], 'and', [0.0, 'c1', 0.333], 'then Class0'],  
['if', [1.0, 'c0', 0.666], 'and', [0.0, 'c1', 0.333], 'then Class0'],  
['if', [2.0, 'c0', 1.0], 'and', [0.0, 'c1', 0.333], 'then Class0']],  
[['if', [2.0, 'c0', 1.0], 'and', [0.0, 'c1', 0.333], 'then Class1'],  
['if', [0.0, 'c0', 0.333], 'and', [1.0, 'c1', 0.666], 'then Class1']])
```

Εικόνα 24 Μορφή βάσης μοναδικών κανόνων (RB).

Στη συνέχεια για κάθε κανόνα της βάσης κανόνων υπολογίζονται τα δείγματα που είναι ίσα με αυτόν τον κανόνα. Εφόσον βρεθούν τα ίσα διανύσματα υπολογίζεται το συνολικό σφάλμα σύμφωνα με την εξίσωση 5.11 της υποενότητας 5.1.5. Για τον υπολογισμό των νέων βαρών της βάσης κανόνων υλοποιείται η εξίσωση 5.13 της υποενότητας 5.1.5. Στην Εικόνα 25 λοιπόν φαίνεται η τελική βάση κανόνων έπειτα από όλες τις επαναλήψεις του αλγορίθμου Gradient Descent (GD).

Υλοποίηση Εξίσωσης 5.13

```
for index,weight in enumerate(rules_weights[a]):  
    if weight != 0:  
        new_weight=weight - (error * mul[index] * learning_rate)  
        rules_weights[a][index]=new_weight  
        tmp_weight.append(new_weight)
```

```
[['if', [0.0, 'c0', 5.703], 'and', [1.0, 'c1', 1.124], 'then Class0'],  
 ['if', [0.0, 'c0', 5.833], 'and', [0.0, 'c1', 5.915], 'then Class0'],  
 ['if', [1.0, 'c0', 3.00], 'and', [0.0, 'c1', 1.683], 'then Class0'],  
 ['if', [2.0, 'c0', 1.50], 'and', [0.0, 'c1', 0.458], 'then Class0'],  
 ['if', [2.0, 'c0', 2.512], 'and', [0.0, 'c1', 6.22], 'then Class1'],  
 ['if', [0.0, 'c0', 1.153], 'and', [1.0, 'c1', 2.28], 'then Class1']]
```

Εικόνα 25 Βάση Κανόνων έπειτα από την εφαρμογή του αλγορίθμου GD

Τέλος, σε αυτό το σημείο ακολουθείται η μέθοδος της υποενότητας 5.1.6 για την ταξινόμηση ενός άγνωστης κλάσης δείγματος. Αρχικά εξάγεται για κάθε δείγμα ποιους κανόνες ενεργοποιεί από την ανανεωμένη βάση κανόνων σύμφωνα με την εξίσωση 5.14 της υποενότητας 5.1.6. Έπειτα υπολογίζεται η σχετική βεβαιότητα όπως και πριν σύμφωνα με την εξίσωση 5.9 και επιπλέον η σχετική βεβαιότητα μαζί με τα βάρη του ή των κανόνων που ενεργοποιήθηκαν. Η συνολική σχετική βεβαιότητα για την ταξινόμηση του δείγματος τελικά πραγματοποιείται με την εξίσωση 5.14 της υποενότητας 5.1.6. Το τελικό αποτέλεσμα του μοντέλου έπειτα από όλα αυτά τα βήματα και έπειτα από το 10- fold cross validation στην οθόνη του χρήστη εμφανίζεται η μέση τιμή και η τυπική απόκλιση της ακρίβειας και της AUC για την απλή ταξινόμηση αλλά και για την ταξινόμηση με την εκπαίδευση των βαρών. Σε αυτό το σημείο προστέθηκε ο συνολικός χρόνος που απαιτείται σε sec για να εξαχθούν τα αποτελέσματα του μοντέλου με στόχο τη χρονική αξιολόγησή του.

```
-----  
Mean: 0.686308973172988 and std: 0.06370256395681305 of accuracies of dummy classification  
Mean: 0.7389516129032259 and std: 0.060696952916628334 of auc score of dummy classification  
Mean: 0.6949583718778909 and std: 0.05347039055991788 of accuracies with gradient descent method  
Mean: 0.7329905913978494 and std: 0.05333250893952363 of auc scores with gradient descent method  
Total time taken: 59.66 seconds
```

Εικόνα 26 Τελικό αποτέλεσμα μοντέλου.

Κεφάλαιο 6

Αποτελέσματα και Πειράματα

Στο κεφάλαιο αυτό αναλύονται οι μετρικές αξιολόγησης που υπολογίστηκαν με σκοπό την ερμηνεία της απόδοσης του ταξινομητή καθώς και τα πειράματα που πραγματοποιήθηκαν. Σε πρώτη φάση για λόγους εγκυρότητας της μεθοδολογίας χρησιμοποιήθηκαν 21 διαφορετικά σετ δεδομένων με ποικιλία στον αριθμό των κλάσεων. Ερμηνεύονται τα αποτελέσματα που εξήχθησαν και συγκριτικά με την ήδη υπάρχουσα υλοποιημένη μεθοδολογία. Παρουσιάζονται διαγραμματικά τα αποτελέσματα και κάποιες επιπλέον πειραματικές δοκιμές. Στη συνέχεια, αφού αποδειχθεί πως η μεθοδολογία μπορεί να θεωρηθεί έγκυρη εφαρμόζεται και σε δεδομένα σημάτων για την αναγνώριση συναισθημάτων.

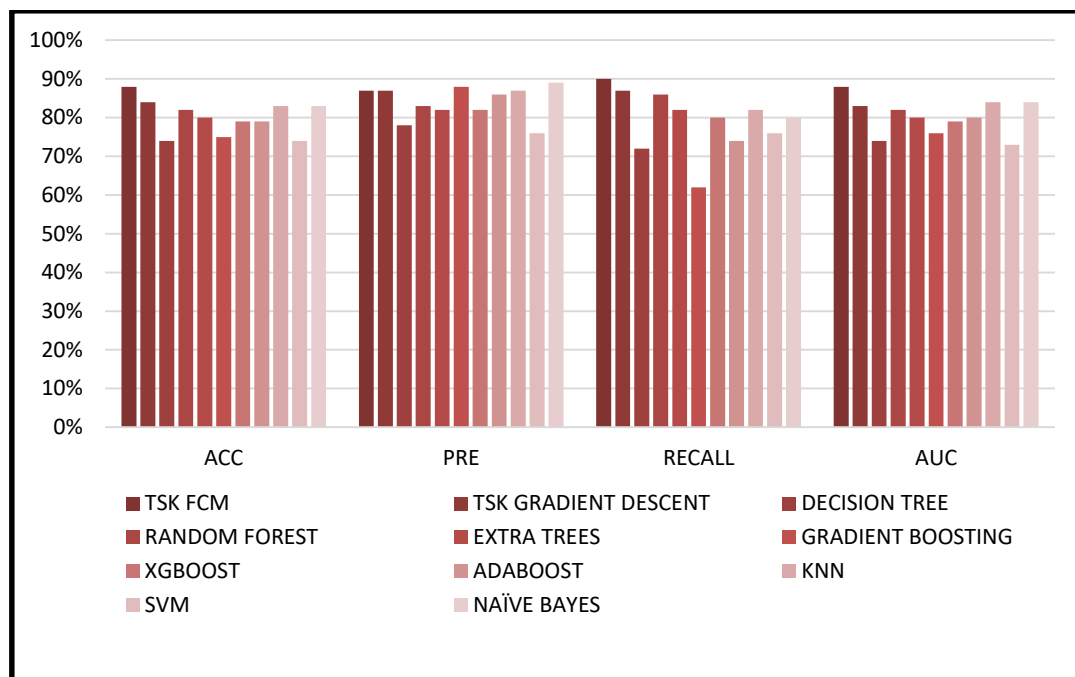
6.1 Μετρικές Αξιολόγησης

- 1) Ορθότητα (Accuracy):** Η ορθότητα αποτελεί μια από τις πιο συχνές μετρικές για την αξιολόγηση ενός ταξινομητή. Εκφράζει το συνολικό ποσοστό των ορθών ταξινομήσεων [107] και υπολογίζεται σύμφωνα με τον τύπο (5.17) όπως αναλύθηκε και προηγουμένως. Η ακρίβεια αποτελεί μια απλή μετρική για να αξιολογήσουμε την απόδοση ενός μοντέλου, αλλά μπορεί να παρουσιάσει προβλήματα σε περιπτώσεις ανισορροπίας κατηγοριών (imbalanced data), όπου οι προβλέψεις ενός μοντέλου μπορεί να επηρεαστούν από τον αριθμό των δειγμάτων σε κάθε κατηγορία. Σε τέτοιες περιπτώσεις, άλλες μετρικές, όπως η ακρίβεια εξισορρόπησης (balanced accuracy) ή η καμπύλη ROC, μπορεί να είναι πιο ενδεδειγμένες. Γι' αυτό το λόγο στην παρούσα εργασία παρουσιάζεται και η μετρική της καμπύλης ROC που αναλύεται παρακάτω.
- 2) Καμπύλη Λειτουργικού Χαρακτηριστικού Δέκτη (Receiver Operating Characteristic Curve – ROC):** Η καμπύλη αυτή αποτελεί μια τεχνική οπτικοποίησης ταξινομητών βασισμένη στην απόδοσή τους. Ουσιαστικά είναι μια γραφική παράσταση δύο διαστάσεων στην οποία τους άξονες τοποθετούνται το ποσοστό των αληθώς θετικών δειγμάτων (True Positive Rate/TPR) στον y άξονα αλλά και το ποσοστό των ψευδώς θετικών δειγμάτων (False Positive Rate/FPR) στον x άξονα. Κατασκευάζεται με τον υπολογισμό των παραπάνω ποσοστών ανά ζεύγη σε διάφορα κατώφλια (thresholds) και έτσι παράγεται και το τελικό σχήμα. Στην παρούσα εργασία χρησιμοποιείται σα μετρική αξιολόγησης η περιοχή κάτω από την καμπύλη η οποία ονομάζεται Area Under Curve (AUC) και αποτελεί

ένα μέτρο διαχωρισμού μεταξύ των κλάσεων που συμμετέχουν στην ταξινόμηση. Το εύρος της τιμής αυτής κυμαίνεται από 0-1. Όσο μεγαλύτερη είναι η τιμή της AUC σημαίνει πως το μοντέλο αποδίδει πολύ καλά στο να προβλέπει την αρνητική κλάση σαν αρνητική και τη θετική κλάση σα θετική. Όσο η τιμή της AUC πλησιάζει το 1 τόσο πιο αξιόπιστο είναι το μοντέλο, ενώ όταν πλησιάζει στο 0.5 σημαίνει ότι το μοντέλο παρουσιάζει δυσκολία ως προς τη διάκριση των θετικών και των αρνητικών δειγμάτων και ότι πιθανότατα ο διαχωρισμός ήταν τυχαίος [108].

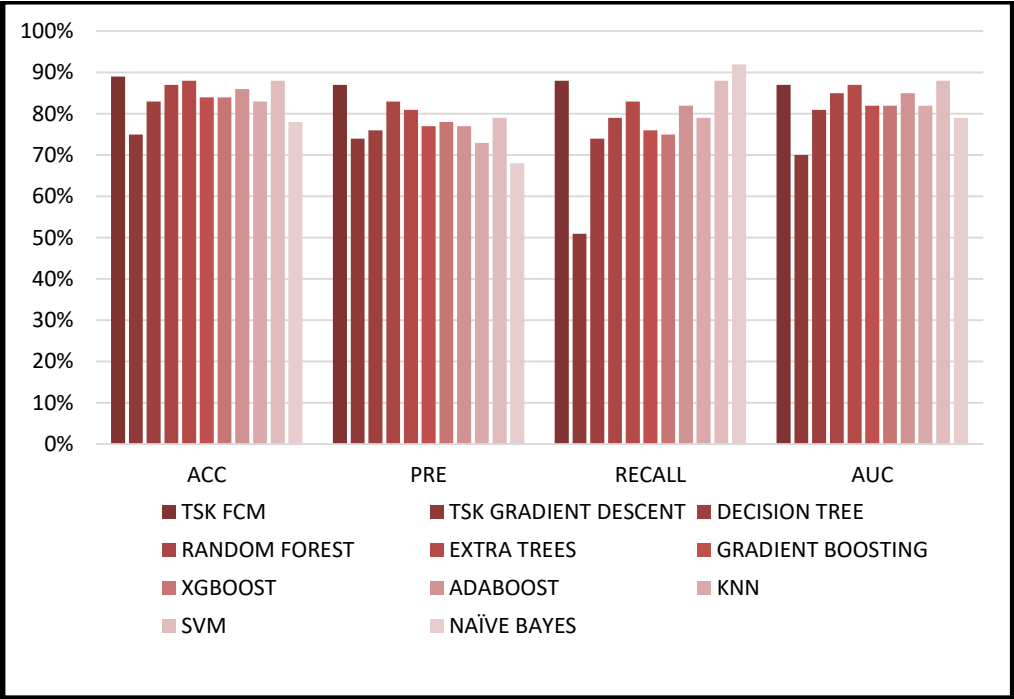
6.2 Προκαταρκτική Μελέτη

Σε αυτό το στάδιο της παρούσας εργασίας πραγματοποιήθηκαν κάποια πειράματα με διαφορετικές παραμέτρους σε ήδη υπάρχουσες μεθοδολογίες με σκοπό την εξαγωγή κάποιων συμπερασμάτων. Οι μεθοδολογίες αυτές είναι τα Δέντρα Αποφάσεων (Decision Trees), τα Τυχαία Δάση (Random Forest), τα Εξαιρετικά Τυχαία Δέντρα (Extra Trees), οι Μηχανές Διανυσμάτων Στήριξης (SVM), ο αλγόριθμος K-Κοντινότερων Γειτόνων (KNN), ο Bayes (Naïve Bayes), ο αλγόριθμος Προσαρμοστικής Ενίσχυσης (Adaboost), η μέθοδος Ενίσχυσης Κλίσης (Gradient Boosting) και η μέθοδος Εξαιρετικής Ενίσχυσης Κλίσης (XGBoost) όπως περιγράφονται στο Κεφάλαιο 2. Παρακάτω, λοιπόν, φαίνονται ενδεικτικά κάποια αποτελέσματα από διάφορα σετ δεδομένων. Φαίνεται το ποσοστό της μετρικής της ακρίβειας ταξινόμησης σε όλα τα διαγράμματα. Δοκιμάστηκαν οι ασαφείς ταξινομητές Takagi-Sugeno σύμφωνα με την υλοποιημένη μεθοδολογία [100] με τον αλγόριθμο ομαδοποίησης Fuzzy-C-means και με τη μέθοδο βελτιστοποίησης Βαθμωτής Κατάδυσης (Gradient Descent). Στόχος από τα παρακάτω διαγράμματα αποτελεί η σύγκριση των δύο αυτών μεθόδων με όλες τις υπόλοιπες που αποτελούν ταξινομητές μηχανικής μάθησης. Παρουσιάζονται παρακάτω κάποια διαγράμματα με τα ποσοστά της ακρίβειας, για κάποια ενδεικτικά σετ δεδομένων που δοκιμάστηκαν.

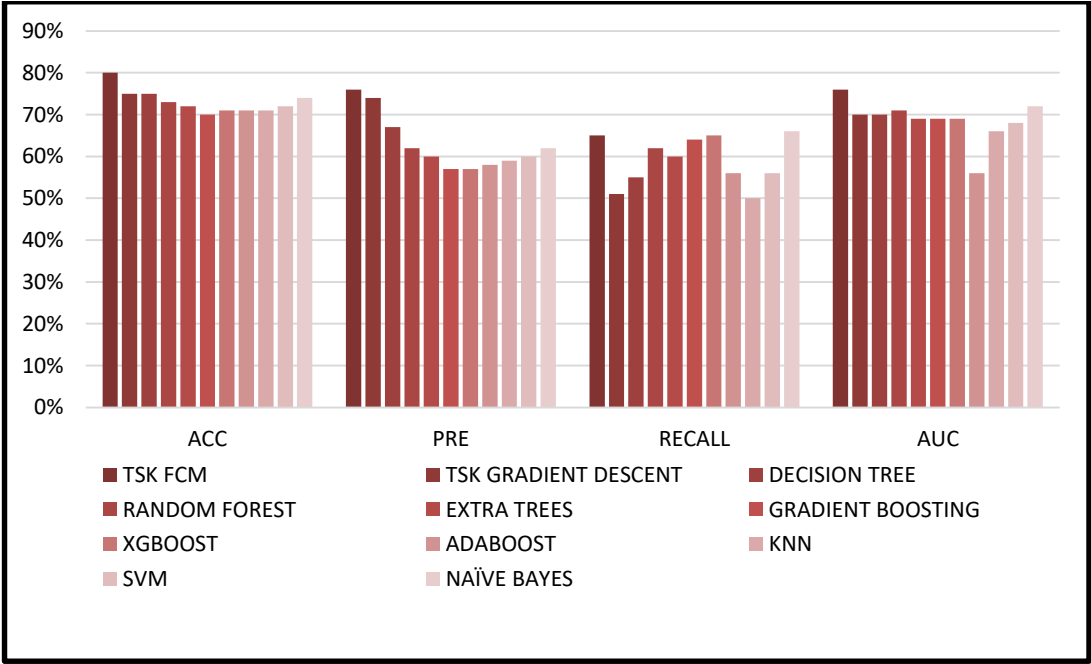


Εικόνα 27 Ποσοστά ακρίβειας σε σετ δεδομένων σχετικά με παθήσεις καρδιάς.

Από την Εικόνα 27 το οποίο αναπαριστά αποτελέσματα ακρίβειας αναφορικά με σετ δεδομένων παθήσεων της καρδιάς φαίνεται πως ο ταξινομητής Takagi-Sugeno και με τη μέθοδο βελτιστοποίησης βαθμωτής κατάδυσης υπερέχει συγκριτικά με όλους τους υπολοίπους. Σε υψηλά ποσοστά ακρίβειας είναι επίσης και ο αφελής Bayes κοντά με τα Takagi-Sugeno.



Εικόνα 28 Ποσοστά ακρίβειας σε βιολογικό σετ δεδομένων.



Εικόνα 29 Ποσοστά ακρίβειας σε σετ δεδομένων σχετικά με την ασθένεια του διαβήτη

Από την Εικόνα 28 για βιολογικό σετ δεδομένων και πάλι τα Takagi-Sugeno έχουν το μεγαλύτερο ποσοστό ακρίβειας μαζί με τις μηχανές διανυσμάτων στήριξης (SVM). Στην Εικόνα 29 και σε σετ δεδομένων που αφορά το διαβήτη με διαφορά από τους υπόλοιπους ταξινομητές τα Takagi – Sugeno κατέχουν το μεγαλύτερο ποσοστό ακρίβειας 80%. Από τα παραπάνω πειράματα σε ήδη υπάρχουσες μεθοδολογίες μπορεί να εξαχθεί ένα πολύ σημαντικό συμπέρασμα. Φαίνεται από όλα τα διαγράμματα πως οι ασαφείς ταξινομητές Takagi-Sugeno εξάγουν τα καλύτερα αποτελέσματα σε σύγκριση με όλους τους υπόλοιπους ταξινομητές. Για αυτό και η σύγκριση της παρούσας μεθοδολογίας πραγματοποιείται με το συγκεκριμένο είδος ασαφών ταξινομητών.

6.3 Πειραματικά Αποτελέσματα Μεθοδολογίας FSP

Όπως προαναφέρθηκε η μεθοδολογία εφαρμόστηκε σε πρώτη φάση σε 21 διαφορετικά σετ δεδομένων ώστε να συγκριθεί με τη δημοσιευμένη εργασία. Σε κάθε σετ δεδομένων, πραγματοποιείται μία τελική σύγκριση μεταξύ των δημοσιευμένων αποτελεσμάτων της μεθοδολογίας υλοποιημένα σε γλώσσα MATLAB και των αποτελεσμάτων της παρούσας εργασίας υλοποιημένα σε γλώσσα Python. Τα αποτελέσματα εξήχθησαν έπειτα από 10 fold cross validation και παρακάτω φαίνεται ο μέσος όρος της μετρικής της ακρίβειας και της AUC καθώς και η τυπική τους απόκλιση. Στον πίνακα 5 εκτός από τα αποτελέσματα της ακρίβειας και της AUC φαίνεται και το πλήθος των κανόνων που εξάγονται ώστε να συνεισφέρουν στην ταξινόμηση. Υπολογίστηκαν αυτές οι δύο μετρικές αρχικά με την εικονική ταξινόμηση και έπειτα με τη συνδρομή της βάσης κανόνων και τη μέθοδο gradient descent για την εκπαίδευση των βαρών του κάθε κανόνα. Σε κάθε πίνακα που ακολουθεί το καλύτερο αποτέλεσμα που αποκτήθηκε αναπαρίσταται είτε με έντονο χρωματισμό είτε με κόκκινο χρώμα.

Δεδομένα	L2-TSK-FS	FCM-TSK	B-TSK	HID-TSK	zero-order-TSK	LFA	NBA	DC*	EasleR	FSP-NRB	FSP	Αποτελέσματα παρούσας εργασίας	
												Ακρίβεια χωρίς βάση κανόνων	Ακρίβεια με GD
AUS	0.777	0.846	0.869	0.834	0.73	–	–	–	0.855	0.85	0.861	0.876	0.877
DIA/PIMA	0.634	0.76	0.772	0.718	0.669	0.783	0.773	0.678	0.729	0.725	0.734	0.74	0.739
LIV	0.634	0.655	0.667	0.656	0.552	0.588	–	0.592	0.594	0.604	0.621	0.646	0.647
HAB	0.722	0.729	0.759	–	–	–	0.761	–	0.74	0.66	0.771	0.761	0.755
RIN	0.97	0.971	0.974	–	–	–	–	–	–	0.888	0.977	0.844	0.845
SAH	0.707	0.701	0.702	–	–	–	–	0.654	–	0.684	0.741	0.707	0.71
HEA	0.831	0.845	0.856	–	–	0.856	0.828	–	0.837	0.833	0.868	0.848	0.87
SEI	0.921	–	–	0.932	0.932	–	–	–	–	0.603	0.922	0.907	0.91
SON	0.662	–	–	0.659	0.493	0.855	–	0.629	0.759	0.822	0.832	0.764	0.769
SPE	0.745	–	–	0.752	0.752	–	–	–	–	0.58	0.806	0.756	0.756
BAL	0.711	–	–	0.842	0.791	–	–	0.726	–	0.723	0.88	0.876	0.887
WIN	–	–	–	–	–	0.996	–	0.767	0.95	0.976	0.976	0.955	0.955
ION	–	–	–	–	–	0.932	–	0.75	–	0.779	0.911	0.91	0.906
APP	–	–	–	–	–	–	–	0.846	0.822	0.825	0.891	0.884	0.893

Πίνακας 4 Σύγκριση αποτελεσμάτων παρούσας μεθοδολογίας με άλλους ασαφείς ταξινομητές [104].

Από όλα τα εξαγόμενα αποτελέσματα φαίνεται πως η προσπάθεια αναπαραγωγής της ήδη υπάρχουσας μεθοδολογίας ανταποκρίνεται αρκετά καλά και επιπλέον για κάποια σετ δεδομένων εξάγονται καλύτερα αποτελέσματα. Κάποιες πληροφορίες για κάθε σετ δεδομένων παρουσιάζεται στον παρακάτω πίνακα αλλά και πιο αναλυτικά υπάρχουν στην ήδη δημοσιευμένη μεθοδολογία [104].

Στον πίνακα 4, αρχικά, συγκρίνεται το ποσοστό της ακρίβειας μόνο μεταξύ ασαφών ταξινομητών σε όσα σετ δεδομένων αφορούν δύο κλάσεις. Παρατηρείται από τον πίνακα αυτό πως γενικά η ακρίβεια της προτεινόμενης μεθοδολογίας [104] και της παρούσας εργασίας έχει λάβει το μεγαλύτερο ποσοστό συγκριτικά με όλους τους υπόλοιπους ταξινομητές για όλα τα σετ δεδομένων. Μάλιστα σε 6 σετ δεδομένων η ακρίβεια της παρούσας εργασίας ξεπέρασε σε ποσοστό τη μεθοδολογία σε MATLAB. Η μεθοδολογία που υλοποιήθηκε σε Python φαίνεται να αναπαράγει πολύ καλά τα αποτελέσματα της μεθοδολογίας που υλοποιήθηκε σε MATLAB. Σε 6/18 σετ δεδομένων και πιο συγκεκριμένα στα appendicitis, ionosphere, liver, diabetes/pima, balance, seismic η παρούσα εργασία πέτυχε καλύτερα αποτελέσματα. Σε έναν τελευταίο πίνακα αποτελεσμάτων παρουσιάζονται συνολικά τα αποτελέσματα της παρούσας εργασίας με τη μεθοδολογία που ακολουθήθηκε αλλά και με διάφορους άλλους ταξινομητές. Εκτός από τις μετρικές σε αυτόν τον πίνακα φαίνεται και ο αριθμός των κανόνων που χρησιμοποιήθηκε στο κάθε πείραμα. Αυτό το πλήθος των κανόνων ποσοτικοποιεί την ερμηνευσιμότητα της μεθοδολογίας. Φαίνεται από τον παρακάτω πίνακα πως για μικρό πλήθος κανόνων τα αποτελέσματα ταξινόμησης είναι αρκετά υψηλά. Παρατηρείται γενικά πως η βάση κανόνων συνεισφέρει σε ένα καλύτερο αποτέλεσμα ταξινόμησης καθώς το μοντέλο γίνεται αυτόματα πιο ερμηνεύσιμο. Στην περίπτωση

για παράδειγμα του liver ο ταξινομητής SVM-G ο οποίος χρησιμοποιεί τη γκαουσσιανή συνάρτηση πυρήνα, πέτυχε τη μεγαλύτερη ακρίβεια (0.687) σε σύγκριση το αποτέλεσμα της παρούσας εργασίας που φαίνεται να είναι (0.647). Ωστόσο, το αποτέλεσμα αυτό του SVM-G δε θεωρείται πλήρες ερμηνεύσιμο σε σχέση με τον ασαφή ταξινομητή που αναπτύχθηκε στην παρούσα εργασία. Μία μέση τιμή της ακρίβειας συνολικά όπως φαίνεται και από τον πίνακα είναι (0.80) χωρίς τη βοήθεια της βάσης κανόνων και (0.81) με την εκπαίδευση των βαρών των κανόνων και αντίστοιχα για τη μετρική AUC είναι (0.789) & (0.777).

																		Αποτελέσματα μεθοδολογίας που ακολουθήθηκε				Αποτελέσματα παρούσας εργασίας			
																		FSP-NRB		FSP		Χωρίς βάση κανόνων		Με εκπαίδευση βαρών με GD	
SVM L		SVM G		KNN		SGERD		GFS-AB		FH-GBML		FURIA		IVTURS		ACC	AUC	STD	ACC	AUC	STD	ACC	AUC	STD	
DAT	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	ACC	STD	
APP	0.88	0.093	0.877	0.06	0.86	0.094	0.887	0.067	0.89	0.087	0.846	0.076	0.831	0.05	0.83	0.077	0.825	0.08	0.891	0.07	0.889	0.090	0.893	0.09	
	0.75	0.116	0.785	0.2	0.78	0.195	0.769	0.182	0.79	0.213	0.715	0.206	0.71	0.19	0.7	0.189	0.791	0.186	0.803	0.195	0.854	0.130	0.86	0.133	
							2.6				6		3.3		15.4		1		3.7				3.65		
AUS	0.85	0.042	0.841	0.04	0.83	0.024	0.856	0.054	0.81	0.037	0.849	0.047	0.864	0.04	0.85	0.047	0.85	0.033	0.861	0.03	0.876	0.030	0.877	0.03	
	0.92	0.049	0.901	0.03	0.79	0.045	0.864	0.048	0.8	0.04	0.851	0.042	0.862	0.04	0.35	0.044	0.91	0.04	0.91	0.038	0.923	0.040	0.905	0.04	
							2.000				6.9		8.3		31.9		1		5.1				3.000		
BANA	0.87	0.014	0.868	0.02	0.86	0.014	0.635	0.072	0.75	0.023	0.816	0.025	0.866	0.01	0.85	0.016	0.763	0.019	0.877	0.013	0.757	0.02	0.761	0.023	
	0.82	0.017	0.883	0.02	0.83	0.023	0.617	0.084	0.74	0.024	0.808	0.029	0.876	0.02	0.84	0.018	0.845	0.019	0.912	0.012	0.819	0.01	0.745	0.056	
							3.000				10.000		17.2		15.1		1.000		104.000				21.75		
BAN	0.7	0.053	0.696	0.04	0.7	0.077	0.63	0.044	0.55	0.076	0.665	0.04	0.691	0.07	0.68	0.848	0.669	0.086	0.688	0.05	0.685	0.05	0.677	0.055	
	0.66	0.087	0.768	0.07	0.67	0.099	0.511	0.046	0.56	0.065	0.559	0.054	0.631	0.07	0.63	0.057	0.704	0.095	0.697	0.098	0.630	0.09	0.635	0.097	
							2.5				6.9		15		28.5		1		4.000				4.000		
DIA/PIMA	0.77	0.05	0.71	0.04	0.73	0.048	0.71	0.039	0.76	0.039	0.752	0.041	0.749	0.04	0.76	0.033	0.725	0.048	0.734	0.057	0.740	0.05	0.739	0.05	
	0.84	0.045	0.771	0.04	0.76	0.06	0.612	0.052	0.72	0.051	0.7	0.043	0.707	0.05	0.71	0.045	0.792	0.047	0.795	0.046	0.747	0.07	0.743	0.06	
							4.1				6.9		7.8		29.8		1		5.7				10.2		
HAB	0.72	0.026	0.725	0.03	0.69	0.048	0.676	0.146	0.73	0.05	0.725	0.036	0.729	0.12	0.75	0.036	0.66	0.095	0.771	0.037	0.761	0.03	0.755	0.03	
	0.69	0.135	0.682	0.12	0.53	0.097	0.536	0.062	0.55	0.065	0.51	0.035	0.592	0.13	0.56	0.062	0.692	0.095	0.694	0.094	0.622	0.13	0.585	0.123	
							3.1				5		3.2		12.9		1		6950.000				9.500		
ION	0.88	0.027	0.951	0.03	0.85	0.098	0.848	0.041	0.65	0.017	0.912	0.056	0.892	0.04	0.93	0.046	0.779	0.098	0.911	0.045	0.910	0.05	0.906	0.06	
	0.88	0.055	0.965	0.02	0.95	0.016	0.796	0.067	0.51	0.018	0.734	0.152	0.868	0.05	0.93	0.054	0.853	0.09	0.959	0.043	0.889	0.06	0.886	0.06	
							4.3				5.7		12.9		26.3		1		4.8				4.45		
LIV	0.67	0.086	0.687	0.1	0.65	0.075	0.566	0.05	0.67	0.065	0.667	0.071	0.675	0.05	0.66	0.073	0.604	0.098	0.621	0.089	0.646	0.08	0.647	0.08	
	0.7	0.117	0.701	0.14	0.61	0.122	0.5	0.046	0.64	0.077	0.648	0.081	0.653	0.06	0.63	0.075	0.619	0.136	0.621	0.115	0.633	0.09	0.630	0.1	
							2.8				8.1		10.6		9.3		1		5.4				2.350		
RIN	0.76	0.017	0.756	0.02	0.72	0.005	0.727	0.025	0.85	0.015	0.876	0.026	0.941	0.01	0.9	0.009	0.888	0.011	0.977	0.026	0.844	0.02	0.845	0.02	
	0.81	0.016	0.829	0	0.5	0.019	0.727	0.025	0.84	0.015	0.876	0.026	0.94	0.01	0.9	0.009	0.961	0.013	0.965	0.015	0.925	0.009	0.915	0.01	
							4.5				7.6		90.9		21.1		1		4.7				61.45		
SAH	0.72	0.047	0.697	0.04	0.67	0.076	0.681	0.051	0.68	0.055	0.689	0.067	0.729	0.04	0.72	0.05	0.684	0.054	0.741	0.065	0.707	0.048	0.71	0.04	
	0.77	0.058	0.701	0.08	0.71	0.089	0.597	0.046	0.57	0.072	0.615	0.074	0.655	0.05	0.67	0.058	0.735	0.073	0.735	0.076	0.740	0.105	0.735	0.09	
							2.5				7.3		6		35.9		1		8.6				7.95		
SEI	0.93	0	0.935	0	0.92	0.009	0.934	0	0.91	0.01	0.934	0.003	0.926	0.01	0.93	0.003	0.603	0.164	0.922	0.027	0.907	0.008	0.91	0.1	
	0.62	0.104	0.694	0.09	0.68	0.066	0.5	0	0.55	0.04	0.5	0	0.52	0.04	0.51	0.013	0.628	0.107	0.696	0.07	0.625	0.07	0.592	0.08	
							2				2.5		9.6		6.1		1		131				9.05		
SON	0.73	0.108	0.808	0.07	0.81	0.086	0.716	0.097	0.47	0.016	0.685	0.097	0.822	0.08	0.8	0.083	0.822	0.074	0.832	0.084	0.764	0.102	0.769	0.098	
	0.84	0.101	0.892	0.05	0.72	0.062	0.709	0.099	0.5	0	0.694	0.074	0.817	0.08	0.8	0.08	0.913	0.043	0.929	0.081	0.739	0.102	0.734	0.1	
							3.5				6.6		11.3		20.3		1		3.8				4.65		
SPE	0.76	0.072	0.794	0.01	0.72	0.055	0.782	0.027	0.21	0.017	0.749	0.068	0.787	0.07	0.78	0.041	0.58	0.126	0.806	0.047	0.756	0.06	0.756	0.069	
	0.78	0.11	0.802	0.11	0.41	0.164	0.493	0.011	0.5	0	0.521	0.084	0.589	0.08	0.63	0.086	0.601	0.157	0.741	0.152	0.820	0.08	0.818	0.079	
							2.1				5.2		12.1		15.2		1		8.4				2.75		
HEA	0.83	0.068	0.822	0.09	0.82	0.066	0.819	0.049	0.73	0.05	0.782	0.076	0.796	0.07	0.85	0.07	0.833	0.064	0.868	0.059	0.848	0.058	0.86	0.046	
	0.89	0.058	0.897	0.07	0.79	0.127	0.807	0.057	0.71	0.052	0.775	0.078	0.788	0.08	0.85	0.076	0.896	0.07	0.899	0.069	0.868	0.06	0.854	0.065	
							3				5.9		8.1		34.3		1		4.8				3.000		
TWO	0.98	0.005	0.969	0.01	0.96	0.005	0.722	0.024	0.87	0.017	0.885	0.006	0.942	0.01	0.93	0.005	0.973	0.005	0.985	0.004	0.979	0.004	0.98	0.004	
	0.99	0.001	0.992	0	0.9	0.012	0.722	0.024	0.87	0.018	0.888	0.009	0.942	0.01	0.92	0.006	0.989	0.002	0.993	0.002	0.994	0.002	0.989	0.005	
							4.2				8.8		84.3		112		1		10				2.8		
BAL	0.88	0.033	0.884	0.02	0.82	0.032	0.757	0.057	0.69	0.035	0.843	0.04	0.843	0.03	0.85	0.039	0.723	0.05	0.88	0.034	0.876	0.019	0.887	0.022	
							4.2				9.8		24.3		63.2		1		9.1						
CLE	0.57	0.05	0.58	0.05	0.58	0.055	0.495	0.1	0.54	0.038	0.535	0.034	0.559	0.05	0.59	0.059	0.533	0.103	0.62	0.072	0.552	0.079	0.548	0.06	
							7.1				5.3		7.4		68.2		1		17.2				11.6		
WIN	0.96	0.027	0.973	0.03	0.96	0.047	0.921	0.055	0.88	0.076	0.91	0.065	0.972	0.03	0.94	0.038	0.976	0.042	0.976	0.042	0.955	0.041	0.955	0.04	
							4				7.7		6.6		13.7		1		10.3						
AVE	0.81	0.043	0.816	0.04	0.79	0.048	0.74	0.053	0.68	0.038	0.788	0.046	0.817	0.04	0.82	0.04	0.757	0.066	0.836	0.047	0.803	0.047	0.804	0.0509	
	0.8	0.071	0.817	0.07	0.71	0.079	0.65	0.056	0.66	0.05	0.692	0.065	0.743	0.06	0.74	0.058	0.795	0.078	0.823	0.073	0.789	0.068	0.777	0.0718	

Πίνακας 5 Σύγκριση αποτελεσμάτων με άλλους ταξινομητές [105].

Γενικά, από όλα τα παραπάνω διαγράμματα παρατηρείται πως η υλοποίηση της παρούσας εργασίας παράγει συγκρίσιμα αποτελέσματα με την ήδη δημοσιευμένη. Όλα τα παραπάνω πειραματικά αποτελέσματα επιβεβαιώνουν, λοιπόν, την ορθότητα της υλοποιημένης μεθοδολογίας σε γλώσσα Python. Σε επόμενο κεφάλαιο η συγκεκριμένη μεθοδολογία εφαρμόζεται σε καρδιολογικά δεδομένα με διάφορες παραλλαγές με σκοπό την εξαγωγή καλύτερων αποτελεσμάτων.

Κεφάλαιο 7

Καρδιολογικά Δεδομένα

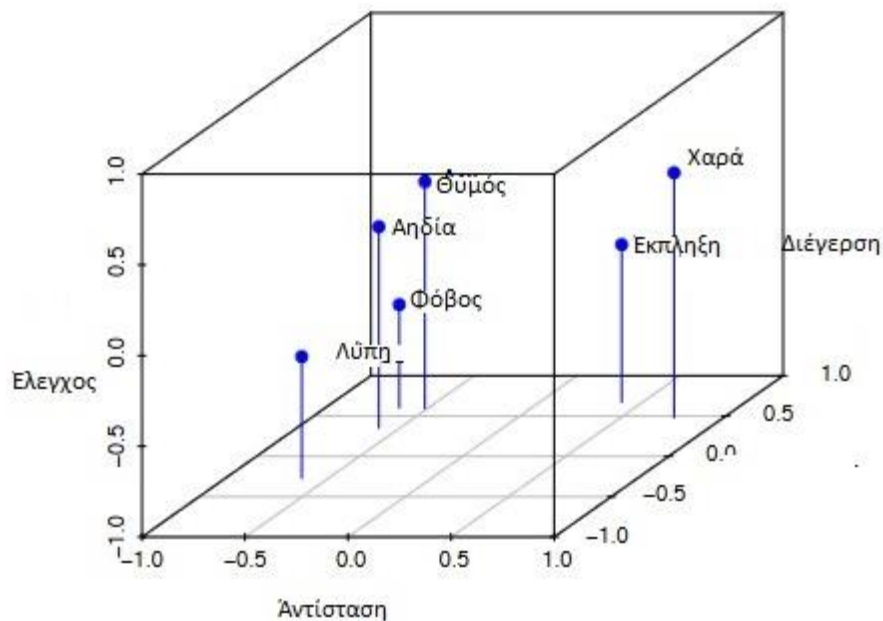
7.1 Δεδομένα καρδιολογικών σημάτων

7.1.1 Περιγραφή συνόλου Δεδομένων (DREAMER)

Η υλοποίηση της παρούσας μεθοδολογίας δοκιμάστηκε σε δεδομένα καρδιολογικών σημάτων. Συγκεκριμένα για το παρόν πείραμα χρησιμοποιήθηκε το σετ δεδομένων DREAMER [109] το οποίο περιέχει καρδιολογικά σήματα. Το σετ δεδομένων αυτό δημιουργήθηκε από 23 συμμετέχοντες (9 γυναίκες και 14 άνδρες) οι οποίοι καταγράφηκαν σε 18 βίντεο που κλήθηκαν να παρακολουθήσουν. Καταγράφηκαν εγκεφαλικά και καρδιολογικά σήματα καθώς και τα συναισθήματα που ένιωθε ο κάθε συμμετέχων σε κάθε βίντεο. Τα συναισθήματα «στόχοι» φαίνονται στον παρακάτω πίνακα και μετρήθηκαν σύμφωνα με ένα από τα πιο διάσημα μοντέλα διαστάσεων συναισθημάτων, το VAD των Bradley and Lang [110], το οποίο φαίνεται σχηματικά στο χώρο στην Εικόνα 30 . Αυτό το μοντέλο διαστάσεων συναισθημάτων αναφέρεται σε τρεις ορθογώνιες διαστάσεις, αυτή της αντίστασης (valence) δηλαδή του βαθμού της πολικότητας του ανθρώπου, αυτή του βαθμού διέγερσης (arousal) δηλαδή της ηρεμίας του ανθρώπου και αυτή του ελέγχου (dominance) δηλαδή του αντιλαμβανόμενου από τον άνθρωπο βαθμό ελέγχου σε μία - κοινωνική- κατάσταση [111]. Οι τρεις αυτές διαστάσεις μετρήθηκαν σε μια κλίμακα από 1-5.

ID	Βίντεο	Συναίσθημα	Αντίσταση (Valence)	Βαθμός Διέγερσης (Arousal)	Έλεγχος (Dominance)
1	Searching for Bobby Fischer	ηρεμία	3.17 ± 0.72	2.26 ± 0.75	2.09 ± 0.73
2	D.O.A.	έκπληξη	3.04 ± 0.88	3.00 ± 1.00	2.70 ± 0.88
3	The Hangover	διασκέδαση	4.57 ± 0.73	3.83 ± 0.83	3.83 ± 0.72
4	The Ring	φόβος	2.04 ± 1.02	4.26 ± 0.69	4.13 ± 0.87
5	300	ενθουσιασμός	3.22 ± 1.17	3.70 ± 0.70	3.52 ± 0.95
6	National Lampoon's VanWilder	αηδία	2.704 ± 1.55	3.83 ± 0.83	4.04 ± 0.98
7	Wall-E	χαρά	4.52 ± 0.59	3.17 ± 0.98	3.57 ± 0.99
8	Crash	θυμός	1.35 ± 0.65	3.96 ± 0.77	4.35 ± 0.65
9	My Girl	λύπη	1.39 ± 0.66	3.004 ± 1.09	3.48 ± 0.95
10	The Fly	αηδία	2.179 ± 1.15	3.309 ± 1.02	3.619 ± 0.89
11	Pride and Prejudice	ηρεμία	3.96 ± 0.64	1.96 ± 0.82	2.61 ± 0.89
12	Modern Times	διασκέδαση	3.96 ± 0.56	2.61 ± 0.89	2.70 ± 0.82
13	Remember the 'Titans	χαρά	4.39 ± 0.66	3.709 ± 0.97	3.749 ± 0.96
14	Gentlemen Agreement	θυμός	2.35 ± 0.65	2.22 ± 0.85	2.39 ± 0.72
15	Psycho	φόβος	2.48 ± 0.85	3.09 ± 1.00	3.22 ± 0.9
16	The Bourne Identity	ενθουσιασμός	3.65 ± 0.65	3.35 ± 1.07	3.269 ± 1.14
17	The Shawshank Redemption	λύπη	1.52 ± 0.59	3.00 ± 0.74	3.96 ± 0.77
18	The Departed	έκπληξη	2.65 ± 0.78	3.91 ± 0.85	3.57 ± 1.04

Πίνακας 6 Μέση τιμή και τυπική απόκλιση από όλους του συμμετέχοντες για κάθε συναίσθημα και για κάθε κλίμακα βαθμολόγησης.



Εικόνα 30 Το μοντέλο VAD στο χώρο ως προς τα 6 πιο βασικά συναισθήματα.

Δεδομένα	DREAMER
Συμμετέχοντες	23
Μέγεθος Δεδομένων	414
Μέθοδος Εξαγωγής Συναισθημάτων	βίντεο
Διαθέσιμα Δεδομένα	Χωρίς προεπεξεργασία
Φυσιολογικά δεδομένα, 'D'	Καρδιολογικά σήματα με ρυθμό δειγματοληψίας 128/256 Hz
Εξαγωγή Χαρακτηριστικών, 'DI'	27 δεδομένα καρδιακού ρυθμού
Αριθμός κλάσεων & κατηγοριών	Αντίσταση (Valence) (0,1)
	Βαθμός Διέγερσης (Arousal) (0, 1)
	Έλεγχος (Dominance) (0, 1)

Πίνακας 7 Πληροφορίες του συνόλου δεδομένων.

Οι μεθοδολογίες που έχουν εφαρμοστεί στο συγκεκριμένο σετ δεδομένων αφορούν κυρίως ανάπτυξη νευρωνικών δικτύων βαθιάς μάθησης [112-114]. Βέβαια έχουν γίνει και εφαρμογές με μεθόδους μηχανικής μάθησης. Η παρούσα πτυχιακή εργασία ακολουθεί την προεπεξεργασία των χαρακτηριστικών του συγκεκριμένου σετ δεδομένων με την εργασία [115] ώστε να συγκριθούν τα τελικά αποτελέσματα μεταξύ τους. Έτσι, λοιπόν, πραγματοποιήθηκαν τρεις ταξινομήσεις σύμφωνα με κάθε διάσταση συναισθημάτων. Τέθηκε σύμφωνα και με τη μεθοδολογία [115] ένα κατώφλι 3 ώστε να μετατραπεί ο βαθμός της κάθε διάστασης σε υψηλός ή χαμηλός. Οι τιμές του βαθμού της κάθε διάστασης μετατράπηκαν τελικά σε 0 (χαμηλός) και 1 (υψηλός). Τα χαρακτηριστικά για κάθε ταξινόμηση είναι 21 καθώς αφαιρέθηκαν κάποια που στην πλειονότητα των δειγμάτων είχαν ελλειπείς (nan) ή την τιμή του απείρου (inf).

7.1.2 Αποτελέσματα Συνόλου Δεδομένων DREAMER

Σύμφωνα με τα παραπάνω εξήχθησαν αποτελέσματα του συγκεκριμένου συνόλου δεδομένων για κάθε διάσταση συναισθημάτων (valence-arousal-dominance). Εκτελέστηκαν τρεις διαφορετικές ταξινομήσεις για το συγκεκριμένο σύνολο δεδομένων, μία για κάθε διάσταση συναισθημάτων (DREAMER-arousal, DREAMER-valence, DREAMER-dominance). Πραγματοποιήθηκαν δοκιμές με όλες τις παραλλαγές της μεθοδολογίας που περιγράφονται αναλυτικά στο κεφάλαιο 5. Έπειτα από μία εξαντλητική μελέτη παρουσιάζονται παρακάτω τα καλύτερα αποτελέσματα από κάθε παραλλαγή. Τα αποτελέσματα της κάθε έκδοσης παρουσιάζονται σύμφωνα με τις συντομογραφίες που αναφέρονται στις υποενότητες του κεφαλαίου 5.1. Ειδικότερα όλοι οι παρακάτω πίνακες χωρίζονται στις εκδόσεις σε γλώσσα MATLAB που παράχθηκαν από τη

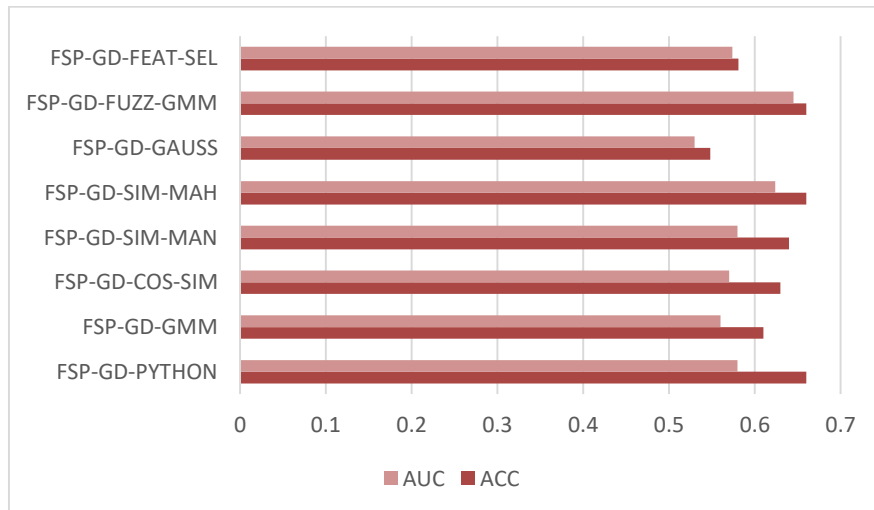
δημοσιευμένη μεθοδολογία [105], στις εκδόσεις σε γλώσσα Python που παράχθηκαν από την παρούσα πτυχιακή εργασία και στις εκδόσεις που παράχθηκαν από τη δοκιμή ενός δημόσια διαθέσιμου πακέτου της Python το PYTSK που αφορά ταξινομητές Takagi-Sugeno όπως περιγράφεται στο κεφάλαιο 4.1.2. Η μορφή του πίνακα των αποτελεσμάτων είναι ίδια για όλα τα σύνολα δεδομένων.

Η πρώτη ταξινόμηση αφορά την πρώτη διάσταση συναισθημάτων η οποία είναι αυτή του βαθμού της αντίστασης (valence) του ανθρώπου. Στη συγκεκριμένη περίπτωση η κατανομή των κλάσεων είναι 246 με βαθμό αντίστασης χαμηλό και 161 με βαθμό αντίστασης υψηλό. Παρακάτω φαίνονται τα αποτελέσματα που εξήχθησαν έπειτα από διάφορες δοκιμές στη μορφή του πίνακα όπως αναλύθηκε παραπάνω. Συγκριτικά φαίνεται πως τα αποτελέσματα της μεθοδολογίας σε γλώσσα MATLAB και της παρούσας μεθοδολογίας είναι αρκετά κοντά, αλλά αισθητά αυξημένα από το πακέτο pytsk.

	DREAMER-valence			
	10-fold cross validation	ACC MEAN (STD)	AUC MEAN (STD)	Total time (sec)
MATLAB	FSP-MATLAB	0.624	0.56	2.12
	FSP-GD-MATLAB	0.644	0.57	18.77
PYTHON	FSP-PYTHON	0.637	0.57	2.55
	FSP-GD-PYTHON	0.66	0.58	14.67
	FSP-GMM	0.58	0.56	2.73
	FSP-GD-GMM	0.61	0.56	14.87
	FSP-COS-SIM	0.61	0.554	3.32
	FSP-GD-COS-SIM	0.63	0.57	31.01
	FSP-SIM-MAN	0.6	0.55	3.28
	FSP-GD-SIM-MAN	0.64	0.58	32.88
	FSP-SIM-MAH	0.64	0.61	3.36
	FSP-GD-SIM-MAH	0.66	0.624	32.86
	FSP-GAUSS	0.53	0.532	4.47
	FSP-GD-GAUSS	0.548	0.53	41.67
	FSP-FUZZ-GMM	0.625	0.615	4.42
	FSP-GD-FUZZ-GMM	0.65	0.645	66.39
	FSP-FEAT-SEL	0.571	0.563	2.63
	FSP-GD-FEAT-SEL	0.581	0.574	25.36
PAPER- [125]	MLP-ALL FEATURES	0.59		
	EXTRA-TREES-EFS	0.74	0.63	
	SVM-PEARSON CORRELATION	0.602	0.53	
PUBLICLY AVAILABLE	PYTSK	0.541	0.514	-
	PYTSK-GD	0.5	0.469	-

Πίνακας 8 Αποτελέσματα ακρίβειας και AUC.

Συγκρίνοντας τις εκδόσεις σε γλώσσα Python μεταξύ τους αυτές που αλλαγές δε βελτίωσαν το αποτέλεσμα είναι οι **FSP-GMM**, **FSP-GD-GMM**, **FSP-GAUSS**, **FSP-GD-GAUSS**, **FSP-FEAT-SEL**, **FSP-GD-FEAT-SEL**. Ειδικότερα οι αλλαγές στον αλγόριθμο συσταδοποίησης, στη συνάρτηση συμμετοχής, και στην επιλογή χαρακτηριστικών δε βελτίωσε επιθυμητά το αποτέλεσμα. Αντιθέτως, η αλλαγή που συνέβαλε στη σημαντική βελτίωση του αποτελέσματος είναι αυτή στον αλγόριθμο συσταδοποίησης για τη δημιουργία των ασαφών συνόλων (**FSP-GD-GMM**). Από τον παραπάνω πίνακα παρατηρείται πως συγκριτικά με τη μεθοδολογία του [115] η μετρική της AUC είναι αυξημένη γεγονός που αποδεικνύει πως η παρούσα μεθοδολογία έχει μεγαλύτερη διακριτική ικανότητα μεταξύ των κλάσεων. Ωστόσο, η μετρική της ακρίβειας είναι ελάχιστα μειωμένη συγκριτικά με τη δημοσιευμένη μεθοδολογία [115] στην περίπτωση του ταξινομητή Extra Trees με μέθοδο επιλογής χαρακτηριστικών (EXTRA-TREES-EFS). Βέβαια στην τελευταία περίπτωση ο ταξινομητής που εξήγαγε μεγαλύτερη ακρίβεια δε θεωρείται ερμηνεύσιμος. Η παρούσα μεθοδολογία, λοιπόν, υπερτερεί καθώς είναι εξ ολοκλήρου ερμηνεύσιμη και ταυτόχρονα δε χάνει πολύ από την ακρίβεια σε σύγκριση με τη δημοσιευμένη μεθοδολογία [115].



Εικόνα 31 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.

Από το παραπάνω διάγραμμα παρατηρείται πως η μετρική της AUC καταλαμβάνει ποσοστά περίπου στο 65%. Αυτό είναι πιθανόν να συμβαίνει καθώς τα χαρακτηριστικά του συγκεκριμένου συνόλου δεδομένων έχουν υψηλή συσχέτιση μεταξύ τους, γεγονός που αποδεικνύεται από τον πίνακα συσχέτισης στη μεθοδολογία των Khan et al στο [115]. Η συγκεκριμένη μεθοδολογία με

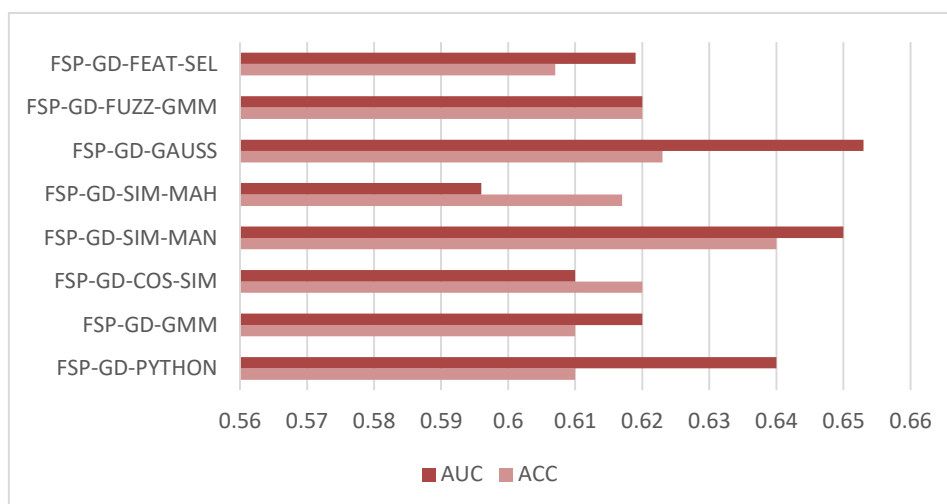
μέθοδο επιλογής χαρακτηριστικών πέτυχε καλύτερη απόδοση και στη μετρική της ακρίβειας αλλά και στη μετρική της AUC.

Επόμενη διάσταση συναισθημάτων είναι αυτή του βαθμού διέγερσης του ανθρώπου (arousal). Στη συγκεκριμένη περίπτωση η κατανομή των κλάσεων είναι 227 με βαθμό αντίστασης χαμηλό και 180 με βαθμό αντίστασης υψηλό. Στον πίνακα των αποτελεσμάτων φαίνονται και τα καλύτερα αποτελέσματα της μεθοδολογίας [115] με και χωρίς μέθοδο επιλογής χαρακτηριστικών. Συγκρίνοντας τις εκδόσεις σε γλώσσα Python μεταξύ τους αυτές που αλλαγές δε βελτίωσαν το αποτέλεσμα είναι οι **FSP-GMM**, **FSP-GD-GMM**, **FSP-SIM-MAH**, **FSP-GD-SIM-MAH**. Ειδικότερα οι αλλαγές στον αλγόριθμο συσταδοποίησης και η ομοιότητα με την απόσταση Mahalanobis δε βελτίωσε επιθυμητά το αποτέλεσμα. Αντιθέτως, η αλλαγή που συνέβαλε στη σημαντική βελτίωση του αποτελέσματος είναι αυτή της ομοιότητας με απόσταση Manhattan (**FSP-GD-SIM-MAH**). Συγκριτικά, με τις τρεις περιπτώσεις της μεθοδολογίας [115] η μετρική της AUC είναι σε καλύτερο ποσοστό στην παρούσα εργασία γεγονός που σημαίνει πως ο ταξινομητής που υλοποιήθηκε στην παρούσα εργασία έχει την ικανότητα να διακρίνει τις κλάσεις σε μεγαλύτερο ποσοστό από ότι στη συγκρινόμενη μεθοδολογία. Ωστόσο, η μετρική της ακρίβειας είναι ελάχιστα μειωμένη συγκριτικά και με τις τρεις μεθοδολογίες του [115]. Βέβαια οι ταξινομητές αυτοί δε θεωρούνται ερμηνεύσιμοι. Η παρούσα μεθοδολογία, λοιπόν, υπερτερεί καθώς είναι εξ ολοκλήρου ερμηνεύσιμη και ταυτόχρονα δε χάνει πολύ από την ακρίβεια σε σύγκριση με τους ταξινομητές της δημοσιευμένης μεθοδολογίας [115].

	DREAMER-arousal			
	10-fold cross validation	ACC MEAN (STD)	AUC MEAN (STD)	Total time (sec)
MATLAB	FSP-MATLAB	0.612	0.64	2.74
	FSP-GD-MATLAB	0.627	0.64	25.88
PYTHON	FSP-PYTHON	0.615	0.641	3.73
	FSP-GD-PYTHON	0.61	0.64	19.47
	FSP-GMM	0.61	0.63	2.63
	FSP-GD-GMM	0.61	0.62	18.51
	FSP-COS-SIM	0.61	0.6	5.58
	FSP-GD-COS-SIM	0.62	0.61	35.55
	FSP-SIM-MAN	0.63	0.646	2.28
	FSP-GD-SIM-MAN	0.64	0.65	20.31
	FSP-SIM-MAH	0.602	0.575	2.63
	FSP-GD-SIM-MAH	0.617	0.596	21.65
	FSP-GAUSS	0.618	0.634	2.22
	FSP-GD-GAUSS	0.623	0.653	16.02
	FSP-FUZZ-GMM	0.6	0.62	2.56
	FSP-GD-FUZZ-GMM	0.62	0.62	18.87
	FSP-FEAT-SEL	0.625	0.634	2.45
	FSP-GD-FEAT-SEL	0.607	0.619	19.37
PAPER- [115]	MLP-ALL FEATURES	0.62		
	DECISION-TREE-EFS	0.74	0.63	
	MLP-PEARSON CORRELATION	0.69	0.63	
PUBLICLY AVAILABLE	PYTSK	0.62	0.46	
	PYTSK-GD	0.63	0.42	

Πίνακας 9 Αποτελέσματα ακρίβειας και AUC.

Από το παρακάτω διάγραμμα παρατηρείται πως η μετρική της AUC καταλαμβάνει ποσοστά περίπου στο 65%. Αυτό είναι πιθανόν να συμβαίνει καθώς τα χαρακτηριστικά του συγκεκριμένου συνόλου δεδομένων έχουν υψηλή συσχέτιση μεταξύ τους, γεγονός που αποδεικνύεται από τον πίνακα συσχέτισης στη μεθοδολογία των Khan et al στο [115] όπως προαναφέρθηκε.



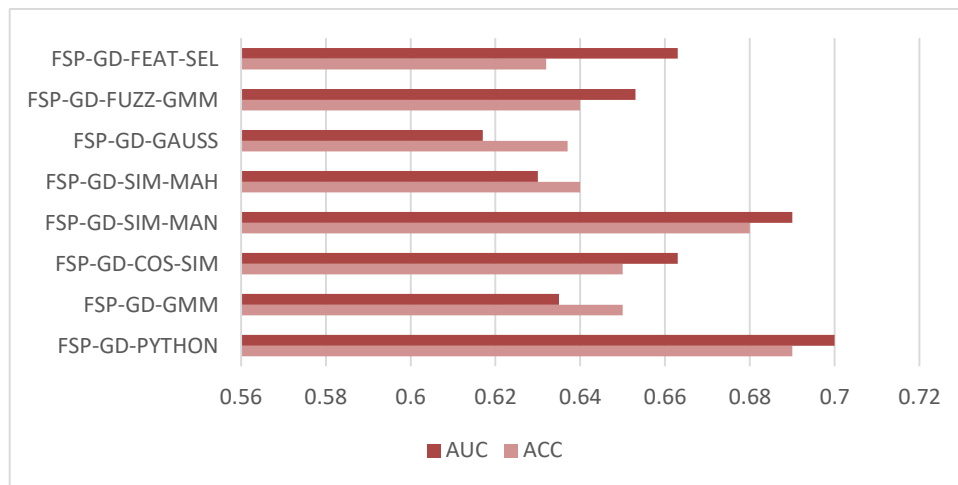
Εικόνα 32 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.

Η τελευταία διάσταση συναισθημάτων είναι αυτή του βαθμού ελέγχου του ανθρώπου (dominance). Στη συγκεκριμένη περίπτωση η κατανομή των κλάσεων είναι 212 με βαθμό αντίστασης χαμηλό και 195 με βαθμό αντίστασης υψηλό. Στον πίνακα των αποτελεσμάτων φαίνονται και τα καλύτερα αποτελέσματα της μεθοδολογίας [115] με και χωρίς μέθοδο επιλογής χαρακτηριστικών. Συγκρίνοντας τις εκδόσεις σε γλώσσα Python μεταξύ τους αυτές που αλλαγές δε βελτίωσαν το αποτέλεσμα είναι οι **FSP-GMM**, **FSP-GD-GMM**, **FSP-GAUSS**, **FSP-GD-GAUSS**, **FSP-FEAT-SEL**, **FSP-GD-FEAT-SEL**. Ειδικότερα οι αλλαγές στον αλγόριθμο συσταδοποίησης, στη συνάρτηση συμμετοχής (γκαουσιανή) και η αλλαγή στο κομμάτι της επιλογής χαρακτηριστικών δε βελτίωσε επιθυμητά το αποτέλεσμα. Στη συγκεκριμένη περίπτωση καμία αλλαγή δε βελτίωσε αρκετά το αποτέλεσμα από την αρχική έκδοση υλοποίησης σε γλώσσα Python. Συγκριτικά, με τις τρεις περιπτώσεις της μεθοδολογίας [115] η μετρική της AUC είναι σε καλύτερο ποσοστό στην παρούσα εργασία γεγονός που σημαίνει πως ο ταξινομητής που υλοποιήθηκε στην παρούσα εργασία έχει την ικανότητα να διακρίνει τις κλάσεις σε μεγαλύτερο ποσοστό από ότι στη συγκρινόμενη μεθοδολογία. Και η μετρική της ακρίβειας κυμαίνεται στο ίδιο επίπεδο με τη μεθοδολογία γεγονός που σημαίνει πως τα μοντέλα προβλέπουν το ίδιο σωστά τις κλάσεις.

	DREAMER-dominance			
	10-fold cross validation	ACC MEAN (STD)	AUC MEAN (STD)	Total time (sec)
MATLAB	FSP-MATLAB	0.66	0.692	3.24
	FSP-GD-MATLAB	0.69	0.69	49.45
PYTHON	FSP-PYTHON	0.667	0.682	4.22
	FSP-GD-PYTHON	0.69	0.7	45.28
	FSP-GMM	0.639	0.642	5.34
	FSP-GD-GMM	0.65	0.635	66.48
	FSP-COS-SIM	0.65	0.674	3.87
	FSP-GD-COS-SIM	0.65	0.663	32.85
	FSP-SIM-MAN	0.67	0.687	4.37
	FSP-GD-SIM-MAN	0.68	0.69	47.57
	FSP-SIM-MAH	0.615	0.621	4.45
	FSP-GD-SIM-MAH	0.64	0.63	47.83
	FSP-GAUSS	0.627	0.62	4.89
	FSP-GD-GAUSS	0.637	0.617	59.12
	FSP-FUZZ-GMM	0.65	0.65	3.34
	FSP-GD-FUZZ-GMM	0.64	0.653	30.08
	FSP-FEAT-SEL	0.627	0.652	3.98
	FSP-GD-FEAT-SEL	0.632	0.663	30.62
PAPER- [115]	DECISION-TREE-ALL FEATURES	0.61		
	DECISION-TREE-EFS	0.69	0.54	
	LOGISTIC REGRESSION-PEARSON CORRELATION	0.61	0.65	
PUBLICLY AVAILABLE	PYTSK	0.68	0.49	
	PYTSK-GD	0.62	0.44	

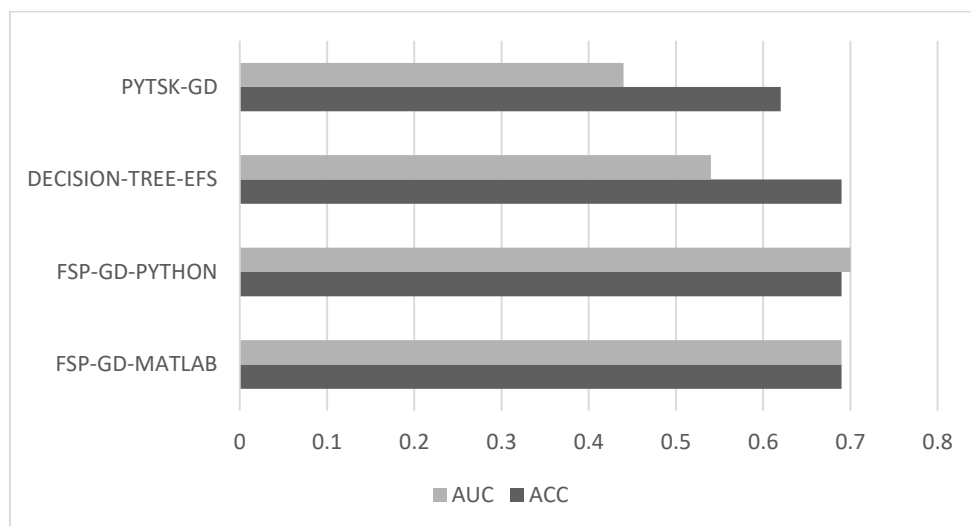
Πίνακας 10 Αποτελέσματα ακρίβειας και AUC.

Όπως προαναφέρθηκε αλλά και γίνεται εμφανές από το παρακάτω διάγραμμα καμία αλλαγή στην αρχική μεθοδολογία δε βελτίωσε το αποτέλεσμα της ταξινόμησης στη συγκεκριμένη διάσταση συναισθήματος, αυτή του ελέγχου. Η αρχική μεθοδολογία υλοποιημένη σε γλώσσα Python έφερε τα καλύτερα αποτελέσματα.



Εικόνα 33 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.

Στο παρακάτω διάγραμμα φαίνεται σχηματικά η σύγκριση της μεθοδολογίας υλοποιημένης σε γλώσσα Python, υλοποιημένης σε γλώσσα MATLAB, το καλύτερο αποτέλεσμα της μεθοδολογίας [115] που πάρθηκε με δέντρα αποφάσεων και του δημόσια διαθέσιμου πακέτου της Python που υλοποιεί ασαφείς ταξινομητές Takagi-Sugeno. Από το παρακάτω διάγραμμα φαίνεται πως η μεθοδολογία σε γλώσσα Python και MATLAB υπερέχει των άλλων δύο.



Εικόνα 34 Σύγκριση τελικών αποτελεσμάτων μεθοδολογιών με τα καλύτερα αποτελέσματα.

Έπειτα από τις τρεις ταξινομήσεις σε όλες τις διαστάσεις συναισθημάτων παρατηρήθηκε πως σε κάποιες περιπτώσεις η ακρίβεια ήταν ελάχιστα μειωμένη σε σύγκριση με άλλους μη ερμηνεύσιμους ταξινομητές. Βέβαια, ο ταξινομητής της παρούσας εργασίας υπερτερεί καθώς θεωρείται ολοκληρωτικά ερμηνεύσιμος, μπορεί δηλαδή να ερμηνεύσει το αποτέλεσμα. Σε αυτό το σημείο αξίζει να σημειωθεί πως οι παράμετροι του συστήματος παίζουν πολύ σημαντικό ρόλο στην απόδοση του. Επομένως αν στη μεθοδολογία χρησιμοποιηθεί ένας αλγόριθμος βελτιστοποίησης παραμέτρων ώστε να βρεθούν οι καλύτερες ενδεχομένως να αυξηθεί η ακρίβεια του συστήματος και να ξεπεράσει τους άλλους ταξινομητές.

7.2 Δεδομένα νόσων της καρδιάς στη Νότια Αφρική (SAHeart-SAH)

7.2.1 Περιγραφή Συνόλου Δεδομένων SAHeart

Πρόκειται για ένα υποδειγματικό δείγμα ανδρών σε μια περιοχή υψηλού κινδύνου για νόσο της καρδιάς στο Δυτικό Ακρωτήριο, Νότια Αφρική. Υπάρχουν περίπου δύο άτομα χωρίς νόσο καρδιάς ανά περίπτωση νόσου της καρδιάς (CHD). Πολλοί από τους άνδρες με θετικό ιστορικό CHD έχουν υποβληθεί σε θεραπεία μείωσης της αρτηριακής πίεσης και άλλα προγράμματα για τη μείωση των παραγόντων κινδύνου μετά το συμβάν της CHD. Σε ορισμένες περιπτώσεις, οι μετρήσεις έγιναν μετά από αυτές τις θεραπείες. Αυτά τα δεδομένα προέρχονται από ένα μεγαλύτερο σύνολο δεδομένων, περιγράφονται στο Rousseauw et al, [117]. Οι Hastie και Tibshirani (1987) επέλεξαν ένα υποσύνολο 465 υποκειμένων από τους 3,357 λευκούς άνδρες (σε αυτές τις κοινότητες, τα ποσοστά θνησιμότητας των ανδρών ήταν περίπου δύο και μισό φορές αυτό των γυναικών, βλ. Rossouw κ.ά., 1983). Τα 465 υποκείμενα αποτελούνταν από όλες τις 162 περιπτώσεις που είχαν νόσο της καρδιάς, καθώς και 303 ελέγχους που επιλέχθηκαν από το υπόλοιπο σύνολο των υποκειμένων της έρευνας.

Τα ίδια (ή παρόμοια) δεδομένα φαίνεται να χρησιμοποιούνται και πάλι για επίδειξη στο Hastie, Tibshirani και Friedman (2009) και αυτό είναι αυτό που τώρα χρησιμοποιείται στην παρούσα εργασία. Παρατηρητικά, αυτό το σύνολο δεδομένων περιέχει τιμές μόνο για 462 παρατηρήσεις, από τις οποίες τώρα μόνο 160 είναι ασθενείς και 302 είναι υγιείς. Το παρόν σετ δεδομένων, λοιπόν, περιέχει 462 παρατηρήσεις με 10 χαρακτηριστικά εκ των οποίων το 1 είναι δυαδικό τα 8 αριθμητικά και το 1 χαρακτηριστικό «στόχος» για την ταξινόμηση της κλάσης.

Δεδομένα νόσων της καρδιάς στη Νότια Αφρική (SAHeart)	
Χαρακτηριστικά	Ελάχιστη-Μέγιστη Τιμή
Συστολική Αρτηριακή Πίεση (sbp)	101-218
Καταλάνωση Καπνού (kg) (tobacco)	0-31.2
Χαμηλή πυκνότητα λιποπρωτεϊνών χοληστερόλης (ldl)	0.98-15.33
Δείκτης Λίπους	6.74-42.49
Οικογενειακό Ιστορικό Καρδιαγγειακής νόσου (famhist)	0 ή 1
Συμπεριφορά Τύπου-A (typea-A)	13-78
Δείκτης Παχυσαρκίας (obesity)	14.7-46.58
Τρέχουσα κατανάλωση αλκοόλ (alcohol)	0.00-147.19
Ηλικία (age)	15-64
Καρδιαγγειακή Νόσος (chd)	0 ή 1

Πίνακας 11 Χαρακτηριστικά συνόλου Δεδομένων SAHeart.

7.2.2 Αποτελέσματα Συνόλου Δεδομένων SAHeart

Όπως και στο προηγούμενο σύνολο δεδομένων στον παρακάτω πίνακα φαίνονται τα αποτελέσματα που εξήχθησαν έπειτα από διάφορες δοκιμές. Ο πίνακας είναι ακριβώς στην ίδια μορφή όπως περιγράφηκε παραπάνω. Πρώτο κομμάτι αποτελεί η μεθοδολογία σε γλώσσα MATLAB, επόμενο κομμάτι οι εκδόσεις σε γλώσσα Python και τελευταίο κομμάτι η δοκιμή με το δημόσια διαθέσιμο πακέτο της Python για ασαφείς ταξινομητές. Από τον παρακάτω πίνακα μπορούν να βγουν κάποια συμπεράσματα για τα αποτελέσματα των εκδόσεων. Σε αυτό το σημείο αξίζει να σημειωθεί πως για όλες τις εκδόσεις παρουσιάζεται και το αποτέλεσμα με εκπαίδευση των βαρών (GD). Για αυτό και όλες οι εκδόσεις είναι με ένα επιπλέον αποτέλεσμα με -GD-. Παρατηρώντας τον παρακάτω πίνακα, λοιπόν, φαίνεται γενικά πως τα αποτελέσματα όλων των εκδόσεων σε γλώσσα Python είναι αρκετά κοντά και ελαφρά αυξημένα με αυτά της μεθοδολογίας σε MATLAB. Παράλληλα σε σύγκριση με τα αποτελέσματα του πακέτου PyTSK τα αποτελέσματα της παρούσας εργασίας είναι σαφώς καλύτερα και αισθητά αυξημένα.

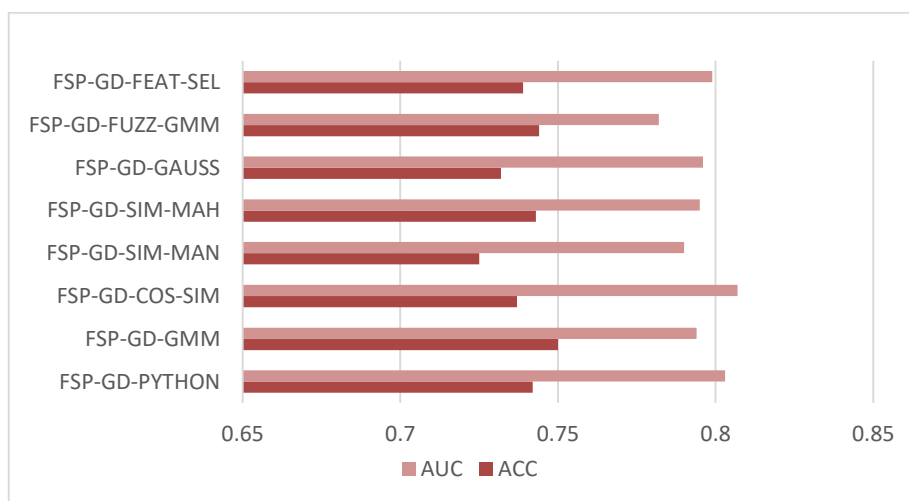
	SAHeart(SAH)			
	VERSIONS	ACC MEAN	AUC MEAN	Total time (sec)
MATLAB	FSP-MATLAB	0.684	0.735	3.23
	FSP-GD-MATLAB	0.741	0.735	45.403
PYTHON	FSP-PYTHON	0.736	0.806	2.8
	FSP-GD-PYTHON	0.742	0.803	72.33
	FSP-GMM	0.738	0.808	3.24
	FSP-GD-GMM	0.75	0.794	38.96
	FSP-COS-SIM	0.742	0.804	3.54
	FSP-GD-COS-SIM	0.737	0.807	65.74
	FSP-SIM-MAN	0.72	0.797	3.53
	FSP-GD-SIM-MAN	0.725	0.79	40.53
	FSP-SIM-MAH	0.746	0.8	3.87
	FSP-GD-SIM-MAH	0.743	0.795	40.85
	FSP-GAUSS	0.723	0.792	3.23
	FSP-GD-GAUSS	0.732	0.796	32.29
	FSP-FUZZ-GMM	0.746	0.8	5.36
	FSP-GD-FUZZ-GMM	0.744	0.782	55.66
	FSP-FEAT-SEL	0.761	0.803	3.23
	FSP-GD-FEAT-SEL	0.739	0.799	6.43
	FSP-COS-SIM-GMM-GAUSS	0.735	0.822	2.73
	FSP-GD-COS-SIM-GMM-GAUSS	0.77	0.826	18.57
PUBLICLY AVAILABLE	PYTSK	0.677	0.621	-
	PYTSK-GD	0.666	0.632	-

Πίνακας 12 Αποτελέσματα ακρίβειας και AUC

Αρχικά παρατηρώντας τον παραπάνω πίνακα φαίνεται πως τα δευτερόλεπτα που χρειάζεται για να ολοκληρωθεί η απλή ταξινόμηση είναι πολύ λιγότερα σε σύγκριση με την ταξινόμηση που γίνεται με εκπαίδευση σε όλες τις εκδόσεις. Αυτό είναι λογικό και αναμένεται να συμβεί σε όλα τα σύνολα δεδομένων καθώς η ταξινόμηση με εκπαίδευση απαιτεί και τη δημιουργία των κανόνων αλλά και την ανανέωση των βαρών τους.

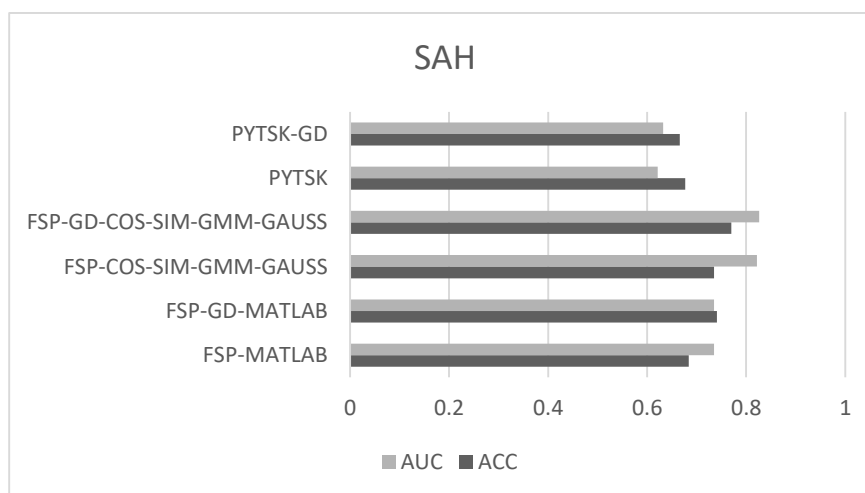
Συγκρίνοντας τις εκδόσεις σε γλώσσα Python μεταξύ τους αυτές που αλλαγές δε βελτίωσαν το αποτέλεσμα είναι οι **FSP-SIM-MAN**, **FSP-GD-SIM-MAN**, **FSP-GAUSS**, **FSP-GD-GAUSS**, **FSP-FEAT-SEL**, **FSP-GD-FEAT-SEL**. Ειδικότερα οι αλλαγές στη συνάρτηση συμμετοχής, η ομοιότητα με απόσταση Manhattan και η αλλαγή στην επιλογή χαρακτηριστικών δε βελτίωσαν επιθυμητά το αποτέλεσμα. Αυτό φαίνεται και στο παρακάτω διάγραμμα στο οποίο συγκρίνονται μεταξύ τους όλες οι εκδόσεις που δοκιμάστηκαν στην παρούσα εργασία με εκπαίδευση των βαρών στη βάση κανόνων (GD). Αντιθέτως, οι αλλαγές που συνέβαλαν στη βελτίωση του αποτελέσματος

είναι αναμφίβολα η ομοιότητα συνημίτονου καθώς από το παρακάτω διάγραμμα (**FSP-GD-COS-SIM**) κατέχει το μεγαλύτερο ποσοστό στη μετρική της AUC. Ποσοστό που ξεπερνά το 80% γεγονός που σημαίνει πως ο ταξινομητής διακρίνει σε πολύ καλό ποσοστό τις θετικές από τις αρνητικές κλάσεις. Ακόμη μία αλλαγή που συνεισφέρει σε ένα καλό αποτέλεσμα γενικότερα είναι αυτή στον αλγόριθμο της πρώτης συσταδοποίησης που χρησιμοποιείται ο γκαουσιανός αλγόριθμος (GMM). Από το παρακάτω διάγραμμα φαίνεται πως η έκδοση **FSP-GD-GMM** κατέχει ένα καλό ποσοστό ακρίβειας ($\approx 75\%$) σε σύγκριση με τις υπόλοιπες εκδόσεις αλλά και η μετρική της AUC φτάνει ένα πολύ καλό ποσοστό της τάξης του 80%. Τα δύο αυτά ποσοστά συνδυαστικά μαρτυρούν το γεγονός ότι ο ταξινομητής αποδίδει καλά τόσο σε σχέση με τη συνολική ορθότητα όσο και στη δυνατότά του στη διάκριση μεταξύ των κλάσεων.



Εικόνα 35 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.

Το καλύτερα αποτέλεσμα, λοιπόν, πάρθηκε από συνδυασμό της ομοιότητας συνημίτονου, γκαουσιανής συνάρτησης συμμετοχής και γκαουσιανό αλγόριθμο ομαδοποίησης (GMM). Παρακάτω φαίνεται και μία οπτικοποίηση του αποτελέσματος στην οποία παρουσιάζεται η καλύτερη εκδοχή της παρούσας εργασίας που δοκιμάστηκε, το αποτέλεσμα σε γλώσσα MATLAB καθώς και το αποτέλεσμα του πακέτου Pytsk.



Εικόνα 36 Σύγκριση αποτελεσμάτων

Από το παραπάνω διάγραμμα φαίνεται πως ο συνδυασμός των εκδόσεων της παρούσας μεθοδολογίας πέτυχε καλύτερα αποτελέσματα σε σύγκριση με τη μεθοδολογία σε γλώσσα MATLAB αλλά και με το πακέτο Pytsk.

7.3 Δεδομένα αξονικής τομογραφίας εκπομπής μονήρους φωτονίου της καρδιάς (SPECTF heart-SPEC)

7.3.1 Περιγραφή Συνόλου Δεδομένων SPECTF heart

Το σύνολο δεδομένων περιγράφει τη διάγνωση εικόνων αξονικής τομογραφίας εκπομπής μονήρους φωτονίου (SPECT) της καρδιάς [118]. Κάθε ασθενής κατατάσσεται σε δύο κατηγορίες: φυσιολογικός και μη φυσιολογικός. Η βάση δεδομένων περιλαμβάνει 267 σύνολα εικόνων SPECT (ασθενείς) που επεξεργάστηκαν για να εξαχθούν χαρακτηριστικά που περιλαμβάνουν τις αρχικές εικόνες SPECT. Ως αποτέλεσμα, δημιουργήθηκαν 44 συνεχή χαρακτηριστικά μοτίβου για κάθε ασθενή. Το μοτίβο επεξεργάστηκε περαιτέρω για να προκύψουν 22 δυαδικά χαρακτηριστικά μοτίβα.

Το SPECT είναι ένα καλό σύνολο δεδομένων για δοκιμή αλγορίθμων μηχανικής μάθησης. Περιλαμβάνει 267 παραδείγματα που περιγράφονται από 44 χαρακτηριστικά και 1 χαρακτηριστικό «στόχος» που αφορά την κλάση στην οποία ταξινομείται το δείγμα. Από τις 267 παρατηρήσεις οι 212 αφορούν φυσιολογικά δείγματα ενώ οι υπόλοιπες 55 αφορούν μη φυσιολογικά δείγματα. Στο

συγκεκριμένο σύνολο δεδομένων υπάρχει ανισορροπία μεταξύ των κλάσεων και αυτό είναι ένα ενδιαφέρον κομμάτι για να παρατηρηθεί η συμπεριφορά του ταξινομητή.

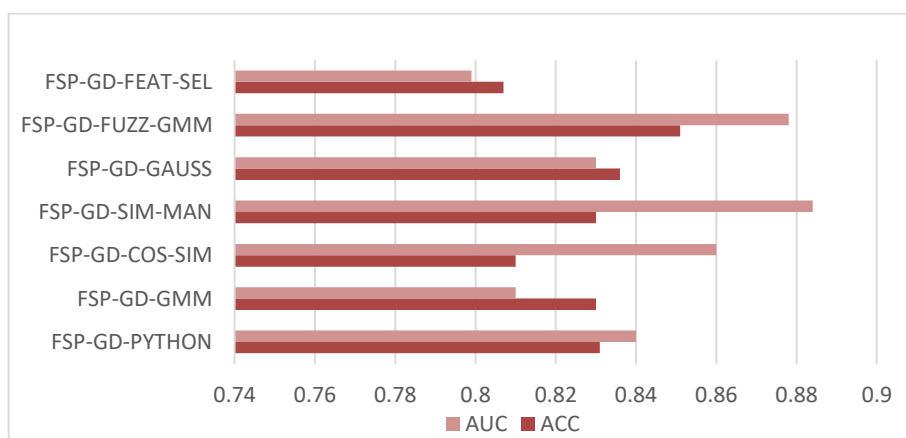
7.3.2 Αποτελέσματα Συνόλου Δεδομένων SPECTF Heart

Έπειτα από διάφορες δοκιμές σε όλες τις παραλλαγές της μεθοδολογίας παρακάτω παρουσιάζονται τα καλύτερα αποτελέσματα για το σύνολο δεδομένων SPECTF Heart. Όλα τα αποτελέσματα εξήχθησαν έπειτα από 10-fold cross validation και στον πίνακα φαίνεται η μέση τιμή της ακρίβειας και της μετρικής της AUC. Στην τελευταία στήλη παρουσιάζεται ο χρόνος σε δευτερόλεπτα που χρειάζεται για να τρέξει το μοντέλο. Με έντονη γραφή φαίνονται οι εκδόσεις της μεθοδολογίας που έφεραν τα καλύτερα αποτελέσματα έπειτα από διάφορους συνδυασμούς των παραλλαγών. Παρατηρώντας τον παρακάτω πίνακα, λοιπόν, φαίνεται γενικά πως τα αποτελέσματα όλων των εκδόσεων σε γλώσσα Python είναι αρκετά κοντά και ελαφρά αυξημένα με αυτά της μεθοδολογίας σε MATLAB. Παράλληλα σε σύγκριση με τα αποτελέσματα του πακέτου PyTSK τα αποτελέσματα της παρούσας εργασίας είναι σαφώς καλύτερα και αισθητά αυξημένα.

	SPECTFHeart (SPEC)			
	10-fold cross validation	ACC MEAN	AUC MEAN	Total time (sec)
MATLAB	FSP-MATLAB	0.58	0.601	2.33
	FSP-GD-MATLAB	0.806	0.741	18.52
PYTHON	FSP-PYTHON	0.793	0.81	2.96
	FSP-GD-PYTHON	0.831	0.84	10.89
	FSP-GMM	0.778	0.81	1.78
	FSP-GD-GMM	0.83	0.81	9.88
	FSP-COS-SIM	0.794	0.87	2.85
	FSP-GD-COS-SIM	0.81	0.86	11.97
	FSP-SIM-MAN	0.823	0.82	2.57
	FSP-GD-SIM-MAN	0.83	0.884	10.66
	FSP-GAUSS	0.747	0.765	2.35
	FSP-GD-GAUSS	0.836	0.83	3.79
	FSP-FUZZ-GMM	0.82	0.876	3.98
	FSP-GD-FUZZ-GMM	0.851	0.878	11.85
	FSP-FEAT-SEL	0.808	0.765	5.76
	FSP-GD-FEAT-SEL	0.807	0.799	32.37
	FSP-GMM-GAUSS	0.81	0.78	1.38
	FSP-GD-GMM-GAUSS	0.841	0.8	10.43
PUBLICLY AVAILABLE	PYTSK	0.722	0.486	
	PYTSK-GD	0.518	0.551	

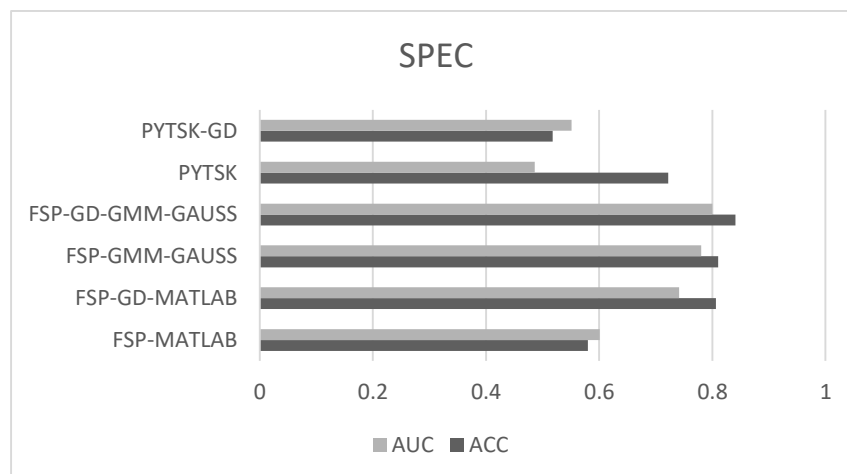
Πίνακας 13 Αναλυτικά αποτελέσματα όλων των εκδόσεων.

Συγκρίνοντας τις εκδόσεις σε γλώσσα Python μεταξύ τους αυτές οι παραλλαγές που δε βελτίωσαν το αποτέλεσμα είναι οι **FSP-COS-SIM**, **FSP-GD-COS-SIM**, **FSP-GAUSS**, **FSP-GD-GAUSS**, **FSP-FEAT-SEL**, **FSP-GD-FEAT-SEL**. Ειδικότερα οι αλλαγές στη συνάρτηση συμμετοχής, η ομοιότητα συνημίτονου και η αλλαγή στην επιλογή χαρακτηριστικών δε βελτίωσε επιθυμητά το αποτέλεσμα. Αυτό φαίνεται και στο παρακάτω διάγραμμα στο οποίο συγκρίνονται μεταξύ τους όλες οι εκδόσεις που δοκιμάστηκαν στην παρούσα εργασία με εκπαίδευση των βαρών στη βάση κανόνων (GD). Αντιθέτως, η αλλαγή που συνέβαλε στη βελτίωση του αποτελέσματος είναι αναμφίβολα η γκαουσιανή συνάρτηση συμμετοχής καθώς από το παρακάτω διάγραμμα (**FSP-GD-FUZZ-GMM**) κατέχει το μεγαλύτερο ποσοστό ακρίβειας. Αυτό το ποσοστό φτάνει το 88% γεγονός που σημαίνει πως ο ταξινομητής μπόρεσε να προβλέψει σωστά μεγάλο αριθμό δειγμάτων. Παρόλα αυτά λόγω της άνισης κατανομής των κλάσεων του συνόλου δεδομένων πρέπει να διερευνηθεί περεταίρω η τιμή της ακρίβειας και ειδικότερα η περίπτωση ο ταξινομητής να προβλέπει πάντα την κλάση με τη μεγαλύτερη αναλογία. Υπολογίζοντας, όμως, την τιμή της AUC (87.8%) αποδεικνύεται η ικανότητα του ταξινομητή να διαχωρίζει τις αρνητικές από τις θετικές κλάσεις. Επομένως, η τιμή της ακριβείας δεν είναι πλασματική και μπορεί να ληφθεί υπόψιν για την αξιολόγηση του ταξινομητή. Τα δύο αυτά ποσοστά συνδυαστικά μαρτυρούν το γεγονός ότι ο ταξινομητής αποδίδει καλά τόσο σε σχέση με τη συνολική ορθότητα όσο και στη δυνατότητά του στη διάκριση μεταξύ των κλάσεων.



Εικόνα 37 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.

Από τον παραπάνω πίνακα φαίνεται πως η παρούσα εργασία με τους συνδυασμούς των εκδόσεών της πέτυχε τα καλύτερα αποτελέσματα σε σύγκριση με τη μεθοδολογία σε γλώσσα MATLAB αλλά και με το πακέτο pytsk. Αυτή η υπεροχή φαίνεται και σχηματικά παρακάτω σε ένα διάγραμμα.



Εικόνα 38 Σύγκριση διάφορων εκδόσεων.

7.4 Δεδομένα νόσων της καρδιάς (Statlog heart)

7.4.1 Περιγραφή Συνόλου Δεδομένων Statlog Heart

Το σύνολο δεδομένων για τη νόσο της καρδιάς του Statlog χρησιμοποιήθηκε από το *UCI Machine learning Repository* [119]. Συγκεκριμένα, περιλαμβάνει 270 δείγματα από ασθενείς, από τα οποία τα 120 αφορούν ασθενείς ενώ τα υπόλοιπα 150 δείγματα υγιείς. Τα δείγματα από ασθενείς και υγιείς περιλαμβάνουν 13 χαρακτηριστικά τα οποία είναι: 1. ηλικία, 2. φύλο, 3. τύπος πόνου στο στήθος (4 τιμές), 4. αρτηριακή πίεση σε κατάσταση ηρεμίας, 5. Χοληστερόλη ορού σε mg/dl, 6. αρτηριακή πίεση νηστείας > 120 mg/dl, 7. αποτελέσματα ηλεκτροκαρδιογραφικής ηρεμίας (τιμές 0,1,2), 8. μέγιστη επιτευχθείσα καρδιακή συχνότητα, 9. πόνος στο στήθος λόγω άσκησης, 10. Ισχαιμία του μυοκαρδίου κατά τη διάρκεια άσκησης 11. η κλίση του τμήματος ST, 12. αριθμός κύριων αγγείων (0-3) που έχουν χρωματιστεί με ακτινοσκόπηση, 13. Ασθένεια θαλασσαιμίας : 3 = φυσιολογικό, 6 = σταθερή ανωμαλία, 7 = αναστρέψιμη ανωμαλία. Αυτό το σύνολο δεδομένων περιέχει 13 χαρακτηριστικά και 2 κατηγορίες. Οι πληροφορίες για την κατηγορία συμπεριλαμβάνονται στο σύνολο δεδομένων ως 1 και 2 αντίστοιχα, για την απουσία και την παρουσία της νόσου αντίστοιχα. Από τα 13 αυτά χαρακτηριστικά τα 3 είναι δυαδικά και τα υπόλοιπα 10 αριθμητικά.

Δεδομένα νόσων της καρδιάς (Statlog heart)	
Χαρακτηριστικά	Ελάχιστη-Μέγιστη Τιμή
Ηλικία (age)	29-77
Φύλο (sex)	0 ή 1
Πόνος στο στήθος (chest pain)	1.00-4.00
Αρτηριακή Πίεση σε κατάσταση Ηρεμίας (resting blood pressure)	94-200
Χοληστερόλη σε mg/dL (cholesterol)	126-564
Δείκτης Σακχάρου >120 mg/dL (blood sugar)	0.00-1.00
Ηλεκτροκαρδιογράφημα σε κατάσταση ηρεμίας (resting electrocardiographic results)	0.0-2.00
Μέγιστος καρδιακός ρυθμός (maximum heart rate achieved)	71-202
Ισχαιμία του μυοκαρδίου κατά τη διάρκεια άσκησης (oldpeak = ST depression induced by exercise relative to rest)	0.00-6.20
Κλίση του τμήματος ST (ST slop)	1.00-3.00
Αριθμός των μέγιστων αγγείων (The number of major vessels)	0.00-3.00
Θαλασσαιμία (thal)	3 (φυσιολογικό), 6(μη φυσιολογικό),
Καρδιαγγειακή Νόσος (chd)	1 ή 2

Πίνακας 14 Χαρακτηριστικά Συνόλου Δεδομένων Statlog Heart

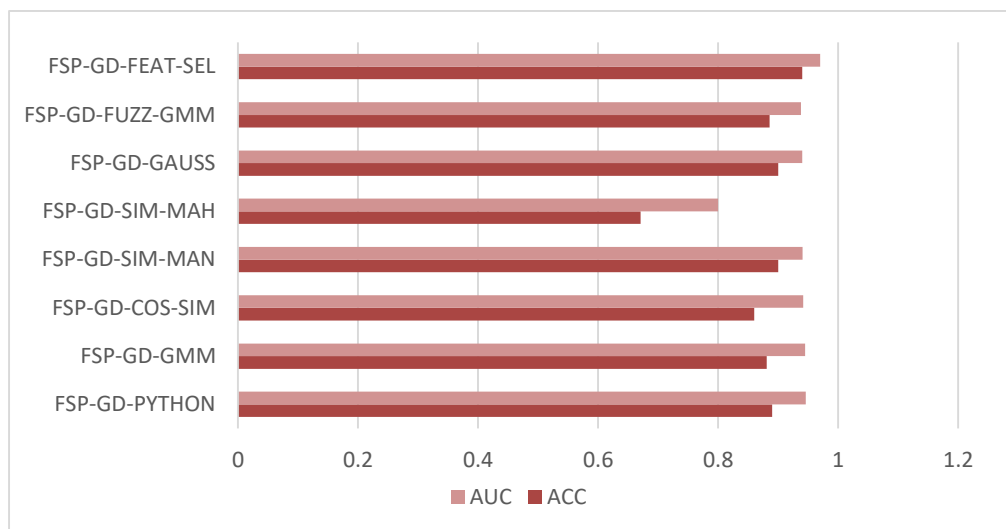
7.4.2 Αποτελέσματα Συνόλου Δεδομένων Statlog Heart

Όπως και σε όλα τα προηγούμενα σύνολα δεδομένων έτσι κι εδώ παρουσιάζεται ο παρακάτω πίνακας στην ίδια μορφή. Συγκεκριμένα για την κάθε έκδοση φαίνονται τα αποτελέσματα και της εικονικής ταξινόμησης αλλά και με την εκπαίδευση των βαρών με τον αλγόριθμο Gradient Descent (-GD-). Παρατηρώντας τον παρακάτω πίνακα, λοιπόν, φαίνεται γενικά πως τα αποτελέσματα όλων των εκδόσεων σε γλώσσα Python είναι αρκετά κοντά και ελαφρά αυξημένα με αυτά της μεθοδολογίας σε MATLAB. Παράλληλα σε σύγκριση με τα αποτελέσματα του πακέτου PyTSK τα αποτελέσματα της παρούσας εργασίας είναι σαφώς καλύτερα και αισθητά αυξημένα.

	Statlog Heart (STAT)			
	10-fold cross validation	ACC MEAN (STD)	AUC MEAN (STD)	Total time (sec)
MATLAB	FSP-MATLAB	0.833	0.896	2.831
	FSP-GD-MATLAB	0.868	0.899	25.369
PYTHON	FSP-PYTHON	0.88	0.951	2.79
	FSP-GD-PYTHON	0.89	0.946	18.81
	FSP-GMM	0.879	0.94	3.39
	FSP-GD-GMM	0.881	0.945	17.41
	FSP-COS-SIM	0.864	0.941	2.22
	FSP-GD-COS-SIM	0.86	0.942	18.19
	FSP-SIM-MAN	0.883	0.94	3.32
	FSP-GD-SIM-MAN	0.9	0.941	19.35
	FSP-SIM-MAH	0.72	0.8	2.23
	FSP-GD-SIM-MAH	0.671	0.8	17.37
	FSP-GAUSS	0.9	0.93	3.29
	FSP-GD-GAUSS	0.9	0.94	15.34
	FSP-FUZZ-GMM	0.879	0.938	3.83
	FSP-GD-FUZZ-GMM	0.886	0.938	18.78
	FSP-FEAT-SEL	0.93	0.92	3.15
	FSP-GD-FEAT-SEL	0.94	0.97	18.1
PUBLICLY AVAILABLE	PYTSK	0.759	0.756	
	PYTSK-GD	0.796	0.794	

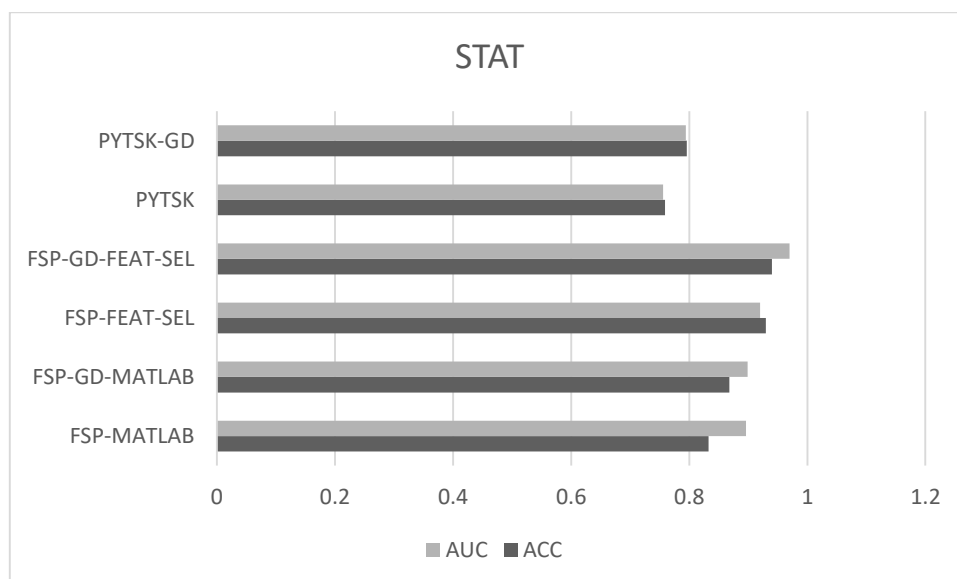
Πίνακας 15 Αναλυτικά αποτελέσματα όλων των εκδόσεων.

Βλέποντας τις εκδόσεις σε γλώσσα Python μεταξύ τους αυτές οι παραλλαγές που δε βελτίωσαν το αποτέλεσμα είναι οι **FSP-SIM-MAH**, **FSP-GD-SIM-MAH**. Ειδικότερα η ομοιότητα με απόσταση Mahalanobis δε βελτίωσε επιθυμητά το αποτέλεσμα. Αυτό φαίνεται και στο παρακάτω διάγραμμα στο οποίο συγκρίνονται μεταξύ τους όλες οι εκδόσεις που δοκιμάστηκαν στην παρούσα εργασία με εκπαίδευση των βαρών στη βάση κανόνων (GD). Αντιθέτως, η αλλαγή που συνέβαλε στη βελτίωση του αποτελέσματος είναι αυτή στο κομμάτι της επιλογής χαρακτηριστικών. Στην έκδοση αυτή (**FSP-GD-FEAT-SEL**) το ποσοστό της ακρίβειας φτάνει στο 94% γεγονός που αποδεικνύει πως ο ταξινομητής προέβλεψε σωστά μεγάλο αριθμό δειγμάτων και η AUC 97% γεγονός που αποδεικνύει πως ο ταξινομητής διαχωρίζει τις θετικές από τις αρνητικές κλάσεις σε αρκετά μεγάλο ποσοστό.



Εικόνα 39 Αποτελεσμάτων εκδόσεων με εκπαίδευση σε γλώσσα Python.

Τα καλύτερα αποτελέσματα πάρθηκαν με την αλλαγή στην επιλογή των χαρακτηριστικών (**FSP-FEAT-SEL** και **FSP-GD-FEAT-SEL**).



Εικόνα 40 Αποτελέσματα διάφορων εκδόσεων

Από τον παραπάνω πίνακα φαίνεται πως η παρούσα εργασία με τους συνδυασμούς των εκδόσεών της πέτυχε τα καλύτερα αποτελέσματα με μικρή διαφορά σε σύγκριση με τη μεθοδολογία σε γλώσσα MATLAB. Βέβαια σε σύγκριση με το πακέτο pytsk η διαφορά είναι αισθητή γεγονός που αποδεικνύει τον ταξινομητή της παρούσας εργασίας αποδοτικότερο από τα απλά Takagi-Sugeno. Αυτή η υπεροχή φαίνεται και σχηματικά παρακάτω σε ένα διάγραμμα.

Κεφάλαιο 8

Συμπεράσματα και Προοπτικές

8.1 Συμπεράσματα

Η παρούσα πτυχιακή διερεύνησε μεθόδους μηχανικής μάθησης, οι οποίες μπορούν να χρησιμοποιηθούν σε συστήματα στήριξης ιατρικών αποφάσεων και εφαρμόζονται σε σήματα, όπως καρδιολογικά σήματα συλλεγμένα μέσω τηλεπαρακολούθησης αποσκοπώντας στην πρόβλεψη ανεπιθύμητων καταστάσεων. Η πτυχιακή επικεντρώθηκε κυρίως στην μελέτη και υλοποίηση της μεθοδολογίας των Ασαφών Φράσεων Ομοιότητας [105], που αποτελεί έναν γενικό ασαφή ταξινομητή εφαρμόσιμο σε πληθώρα εφαρμογών ταξινόμησης σημάτων. Ο ασαφής ταξινομητής είναι ερμηνεύσιμος και καλύπτει την ανάγκη για ερμηνεία και επεξήγηση των εξαγόμενων αποτελεσμάτων στο πεδίο της μηχανικής μάθησης. Το μοντέλο περιλαμβάνει μεθοδολογίες όπως για παράδειγμα εξαγωγή χαρακτηριστικών, εξαγωγή κανόνων, επιλογή χαρακτηριστικών και πρόβλεψη κλάσεων. Στην παρούσα εργασία εκτός από την αναπαραγωγή της υλοποίησης εξ ολοκλήρου σε γλώσσα Python πραγματοποιήθηκε και μία εξαντλητική μελέτη με σκοπό την εξαγωγή καλύτερων αποτελεσμάτων. Δοκιμάστηκαν διάφορες παραλλαγές με αλλαγές που αναλύονται παραπάνω και συγκρίνονται τα τελικά αποτελέσματα. Από την υλοποίηση αυτού του ασαφούς μοντέλου σε γλώσσα Python εξάγονται κάποια χρήσιμα συμπεράσματα:

- Παρέχει αρκετά καλή απόδοση ταξινόμησης σε σύγκριση με άλλους ασαφείς ταξινομητές και σε σύγκριση με την υλοποίηση σε MATLAB. Η παρούσα εργασία έδειξε πως, σε όσα σύνολα δεδομένων δοκιμάστηκαν, η ακρίβεια φτάνει σε ποσοστό 94% με την εκπαίδευση των βαρών στους κανόνες σε σύγκριση με τα Takagi Sugeno που φτάνουν 78%. Πιο συγκεκριμένα η ακρίβεια στο σύνολο δεδομένων Statlog heart έφτασε σε ποσοστό 94% με την παραλλαγή στο κομμάτι της επιλογής χαρακτηριστικών (**FSP-GD-FEAT-SEL**). Για τη συγκεκριμένη έκδοση η μετρική της AUC έφτασε, επίσης, ένα πολύ καλό ποσοστό της τάξης του 97% γεγονός που δείχνει πως το μοντέλο διακρίνει πολύ καλά τις αρνητικές και τις θετικές κλάσεις. Στο ίδιο παράδειγμα η δημοσιευμένη μεθοδολογία [105] για το ίδιο σύνολο δεδομένων το Statlog heart πέτυχε ακρίβεια 86% και AUC 89%. Συνεπώς, η υλοποίηση σε Python πέτυχε αύξηση στην ακρίβεια και στη μετρική της AUC κατά 8%.
- Σε σύγκριση με άλλους ταξινομητές η απόδοση μπορεί να είναι σε μικρότερο ποσοστό κάποιες φορές (όπως για παράδειγμα στα δέντρα αποφάσεων) κυρίως στο σύνολο

δεδομένων DREAMER. Όλοι οι ταξινομητές που συγκρίθηκαν από τη μεθοδολογία [115] δεν ανήκουν στους ασαφείς ταξινομητές και επομένως δε θεωρούνται ερμηνεύσιμοι από μόνοι τους. Ωστόσο, ο ταξινομητής που υλοποιήθηκε στην παρούσα εργασία μπορεί να ερμηνεύσει το αποτέλεσμα χωρίς να χάνει πολύ από την ακρίβεια για αυτό και υπερτερεί σε σύγκριση με άλλους μη ερμηνεύσιμους.

- Σε κάθε σύνολο δεδομένων η αλλαγή που έφερε το καλύτερο αποτέλεσμα είναι διαφορετική. Για παράδειγμα στο σύνολο δεδομένων SAHeart το καλύτερο αποτέλεσμα προήλθε από την εφαρμογή τριών αλλαγών ταυτόχρονα, αυτή της ομοιότητας συνημίτονου, του αλγορίθμου συσταδοποίησης GMM και της συνάρτησης συμμετοχής GAUSS. Το ποσοστό της ακρίβειας σε αυτήν την περίπτωση είναι αυξημένο κατά 3% και της AUC κατά 9% σε σύγκριση με την υλοποίηση σε MATLAB. Παράλληλα στο σύνολο δεδομένων SPECTFheart το καλύτερο αποτέλεσμα προήλθε από την εφαρμογή δύο αλλαγών ταυτόχρονα, αυτή του αλγορίθμου συσταδοποίησης GMM για τη δημιουργία του λεξικού και αυτή της συνάρτησης συμμετοχής GAUSS. Τελικά το ποσοστό της ακρίβειας αυξήθηκε κατά 4% και της AUC κατά 6% σε σύγκριση με την αρχική έκδοση σε MATLAB. Συνεπώς, με βάση τα πειραματικά αποτελέσματα αναδείχθηκε η δυνατότητα προσαρμογής της μεθοδολογίας των Ασαφών Φράσεων Ομοιότητας στην εκάστοτε εφαρμογή ταξινόμησης δεδομένων.
- Η υλοποίηση σε Python χρειάζεται λιγότερο χρόνο να εκτελεστεί από ότι σε γλώσσα MATLAB. Αυτό φαίνεται σε όλους τους πίνακες με τα αποτελέσματα στο κεφάλαιο 7. Αυτό συμβαίνει καθώς η υλοποίηση σε Python χειριζόταν τα δεδομένα με πίνακες (NumPy arrays) με αποτέλεσμα την καλύτερη χρήση της μνήμης RAM του συστήματος κατά την ταξινόμηση των δεδομένων σε σχέση με αυτή της MATLAB που υλοποιήθηκε χρησιμοποιώντας δομές (structs). Συνεπώς, η υλοποίηση σε Python, μπορεί ευκολότερα να εφαρμοστεί σε μεγαλύτερα σύνολα δεδομένα, ακόμη και σε δεδομένα εικόνων.
- Η υλοποίηση σε Python είναι διαχειρίσιμη και εύκολα επεκτάσιμη καθώς κατασκευάστηκε σε διαφορετικές συναρτήσεις ανάλογα με το στάδιο της μεθοδολογίας. Στοιχείο το οποίο δεν υπάρχει στην υλοποίηση σε γλώσσα MATLAB. Μπορεί, λοιπόν, εύκολα να τροποποιηθεί και να επεκταθεί ανάλογα με το στάδιο στο οποίο χρειάζονται αλλαγές. Είναι, επίσης, εύκολο να παραλληλοποιηθεί ώστε να εκτελείται είτε στη CPU είτε στη GPU.

- Η ερμηνευσιμότητα της μεθοδολογίας μπορεί να ποσοτικοποιηθεί με τον αριθμό των κανόνων που εξάγονται. Παρατηρήθηκε πως για μικρό πλήθος κανόνων τα αποτελέσματα ταξινόμησης είναι αρκετά υψηλά.

8.2 Μελλοντικές Προοπτικές

Η παρούσα πτυχιακή εργασία αποτελεί μια μελέτη και μετάφραση της μεθόδου [105] από τη γλώσσα MATLAB στη γλώσσα προγραμματισμού Python. Μία πρώτη εφαρμογή της μεθοδολογίας με δυνατότητες επιπλέον μελλοντικής διερεύνησης που πραγματοποιήθηκε στο πλαίσιο της παρούσας πτυχιακής εργασίας αφορούσε δεδομένα σημάτων για την αναγνώριση συναισθημάτων. Άλλος μελλοντικός στόχος αναφορικά με την εφαρμογή της αποτελεί η δοκιμή σε διάφορα άλλα πεδία όπως για παράδειγμα στην αναγνώριση κειμένου αλλά και σε εικόνες. Πέρα από αυτό, γενικά, στόχος αποτελεί η βελτίωση της απόδοσης του ταξινομητή με διάφορες άλλες τεχνικές. Κάποιες από αυτές είναι η δοκιμή ενός άλλου αλγορίθμου για την εκπαίδευση των βαρών των κανόνων αλλά και ενός άλλου αλγορίθμου για τη δημιουργία των συστάδων. Ένα ανοιχτό πεδίο για διερεύνηση είναι ακόμη και το κομμάτι της επιλογής χαρακτηριστικών το οποίο μπορεί να βελτιώσει κι αυτό την απόδοση.

Βιβλιογραφία

- [1] “Cardiovascular diseases (cvds),” *World Health Organization*, 11-Jun-2021. [Online]. Available: [https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-\(cvds\)](https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-(cvds)). [Accessed: 06-Feb-2022].
- [2] “Heart failure,” *National Heart Lung and Blood Institute*. [Online]. Available: <https://www.nhlbi.nih.gov/health-topics/heart-failure>. [Accessed: 06-Feb-2022].
- [3] P. Ponikowski et al., “Heart failure: preventing disease and death worldwide”, *ESC Heart Failure*, vol 1, no 1, bll 4–25, 2014.
- [4] C. Zednik, “Solving the black box problem: A normative framework for explainable artificial intelligence,” *Philosophy & Technology*, vol. 34, no. 2, pp. 265–288, 2019.
- [5] K.-A. Bowles, E. H. Skinner, D. Mitchell, R. Haas, M. Ho, K. Salter, K. May, D. Markham, L. O’Brien, S. Plumb, T. P. Haines, and M. N. Sarkies, “Data Collection Methods in Health Services Research,” *Applied Clinical Informatics*, vol. 06, no. 01, pp. 96–109, 2015.
- [6] V. Chaurasia, “Data Mining Approach to Detect Heart Dieses”, *International Journal of Advanced Computer Science and Information Technology*, vol.2, pp.56-66, 2013.
- [7] K. Vembandasamy, E. Deepa, and R. Sasipriya, “Heart diseases detection using naive Bayes algorithm - IJISSET,” *International Journal of Innovative Science, Engineering & Technology*, 09-Sep-2015. [Online]. Available: http://ijiset.com/vol2/v2s9/IJISSET_V2_I9_54.pdf. [Accessed: 08-Feb-2022].
- [8] K. Vembandasamy, E. Deepa, and R. Sasipriya, “Heart diseases detection using naive Bayes algorithm - IJISSET,” *International Journal of Innovative Science, Engineering & Technology*, 09-Sep-2015. [Online]. Available: http://ijiset.com/vol2/v2s9/IJISSET_V2_I9_54.pdf. [Accessed: 08-Feb-2022].
- [9] L. Hussain, K. Lone, I. Awan, A. Abbasi, and J.-u.-R. Pirzada, “Detecting congestive heart failure by extracting multimodal features with synthetic minority oversampling technique (SMOTE) for imbalanced data using robust machine learning techniques,” *Waves in Random and Complex Media*, vol. 32, pp. 1–24, Aug. 2020.
- [10] F. S. Alotaibi, “Implementation of Machine Learning Model to Predict Heart Failure Disease,” *International Journal of Advanced Computer Science and Applications*, vol. 10, Art. no. 6, 2019.

- [11] B. Mahesh, “Machine learning algorithms-a review,” *International Journal of Science and Research (IJSR)*.*[Internet]*, vol. 9, pp. 381–386, 2020.
- [12] G. Wiederhold and J. McCarthy, “Arthur Samuel: Pioneer in Machine Learning,” *IBM Journal of Research and Development*, vol. 36, Art. no. 3, 1992.
- [13] T. G. Dietterich, “Machine learning,” *Annual review of computer science*, vol. 4, Art. no. 1, 1990.
- [14] M. Mohammed, M. B. Khan, and E. B. M. Bashier, *Machine learning: algorithms and applications*. Crc Press, 2016.
- [15] L. P. Kaelbling, M. L. Littman, and A. W. Moore, “Reinforcement learning: A survey,” *Journal of artificial intelligence research*, vol. 4, pp. 237–285, 1996.
- [16] S. F. Dessai, “Intelligent heart disease prediction system using probabilistic neural network,” *International Journal on Advanced Computer Theory and Engineering (IJACTE)*, vol. 2, no. 3, pp. 2319–2526, 2013.
- [17] T. Yiu, “Understanding Random Forest,” *Towards Data Science*, 2019.
- [18] M. Pal, “Random forest classifier for remote sensing classification,” *International Journal of Remote Sensing*, vol. 26, no. 1, pp. 217–222, 2005.
- [19] L. Breiman, “Random forests - Random Features,” *Statistics Department University of California Berkeley*, 1999.
- [20] Nouman, “Random Forest — Machine Learning Algorithms with Implementation in Python,” *Medium Daily Digest*, 2021.
- [21] P. Geurts, D. Ernst, and L. Wehenkel, “Extremely randomized trees,” *Machine Learning*, vol. 63, no. 1, pp. 3–42, 2006.
- [22] N. Bhandari, “Extra Trees Classifier- How does Extra Trees Classifier reduce the risk of overfitting?,” *Towards Data Science*, 2018.
- [23] C. J. Burges, “A Tutorial on Support Vector Machines for Pattern Recognition,” *Data mining and knowledge discovery*, 1999.
- [24] Dimca, “Hyperplane arrangements.,” *Springer International Publishing*, 2017.
- [25] S. Firmin, “An Introduction to Support Vector Machines,” *Alteryx Community*, 2019
- [26] Friedrichs Frauke and C. Igel, “Evolutionary tuning of multiple SVM parameters,” *Neurocomputing* 64, 2005.
- [27] T. Cover and P. Hart, “Nearest neighbor pattern classification,” *IEEE transactions on information theory*, vol. 13, Art. no. 1, 1967.

- [28] S. B. Kotsiantis, I. Zaharakis, P. Pintelas, and others, “Supervised machine learning: A review of classification techniques,” *Emerging artificial intelligence applications in computer engineering*, vol. 160, Art. no. 1, 2007.
- [29] B. Gnaneswar and M. R. E. Jebarani, “A review on prediction and diagnosis of heart failure,” *2017 International Conference on Innovations in Information, Embedded and Communication Systems (ICIIECS)*, 2017.
- [30] L. C. Van Der Gaag, “Bayesian belief networks: odds and ends,” *The Computer Journal*, vol. 39, Art. no. 2, 1996.
- [31] S. A. Pattekari and A. Parveen, “Prediction system for heart disease using Naïve Bayes.”
- [32] O. Sagi and L. Rokach, “Ensemble learning: A survey,” *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, vol. 8, Art. no. 4, 2018.
- [33] J. Rocca, “Ensemble methods: bagging, boosting and stacking,” *Towards Data Science*, 2019.
- [34] Y. Freund and R. E. Schapire, “A Decision-Theoretic Generalization of On-Line Learning and an Application to Boosting,” *Journal of Computer and System Sciences*, vol. 55, Art. no. 1, 1997.
- [35] Y. Freund and R. E. Schapire, “A Decision-Theoretic Generalization of On-Line Learning and an Application to Boosting,” *Proceedings of the Second European Conference on Computational Learning Theory*, vol. 55, Art. no. 1, 1995.
- [36] R. Wang, “AdaBoost for Feature Selection, Classification and Its Relation with SVM, A Review,” *Physics Procedia*, vol. 25, pp. 800–807, 2012.
- [37] T.-K. An and M.-H. Kim, “A New Diverse AdaBoost Classifier,” in *2010 International Conference on Artificial Intelligence and Computational Intelligence*, 2010, vol. 1, pp. 359–363.
- [38] J. H. Friedman, “Greedy function approximation: A gradient boosting machine.,” *The Annals of Statistics*, vol. 29, Art. no. 5, 2001.
- [39] J. H. Friedman, “Stochastic gradient boosting,” *Computational statistics & data analysis*, vol. 38, Art. no. 4, 2002.
- [40] Aler, R., Galvn, I.M., Ruiz-Arias, J.A., Gueymard, C.A.: Improving the separation of direct and diffuse solar radiation components using machine learning by gradient boosting. *Solar Energy* 150, 558–569 (2017).
- [41] T. Chen and C. Guestrin, “Xgboost: A scalable tree boosting system,” in *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, 2016, pp. 785–794.

- [42] C. Molnar, “Interpretable machine learning. A Guide for Making Black Box Models Explainable,” p. 328, Feb. 2022.
- [43] H. Alemzadeh, J. Raman, N. Leveson, Z. Kalbarczyk, and R. K. Iyer, “Adverse events in robotic surgery: A retrospective study of 14 years of FDA Data,” *PLOS ONE*, vol. 11, no. 4, 2016.
- [44] W. J. Murdoch, C. Singh, K. Kumbier, R. Abbasi-Asl, and B. Yu, “Definitions, methods, and applications in interpretable machine learning,” *Proceedings of the National Academy of Sciences*, vol. 116, no. 44, pp. 22071–22080, 2019.
- [45] D. V. Carvalho, E. M. Pereira, and J. S. Cardoso, “Machine Learning Interpretability: A Survey on Methods and Metrics,” *Electronics*, vol. 8, no. 8, Art. no. 8, Aug. 2019, doi: 10.3390/electronics8080832
- [46] J. H. Friedman, “Greedy function approximation: A gradient boosting machine.,” *The Annals of Statistics*, vol. 29, Art. no. 5, 2001.
- [47] S. M. Lundberg and S.-I. Lee, “A Unified Approach to Interpreting Model Predictions,” *31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA.*, 2017.
- [48] Z. Q. Yuan Meng Nianhua Yang and G. Zhang, “What Makes an Online Review More Helpful: An Interpretation Framework Using XGBoost and SHAP Values,” *Journal of Theoretical and Applied Electronic Commerce Research*, 2020.
- [49] M. T. Ribeiro, S. Singh, and C. Guestrin, ““ Why should i trust you?” Explaining the predictions of any classifier,” in *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining*, 2016, pp. 1135–1144.
- [50] L. A. Zadeh, “Fuzzy sets,” *Information and Control*, vol. 8, Art. no. 3, 1965.
- [51] L. A. Zadeh, “Fuzzy logic,” *Computer*, vol. 21, Art. no. 4, 1988.
- [52] Mousavi J., Ponnambalam K., Karray F., “Inferring operating rules for reservoir operations using fuzzy regression and ANFIS,” *Fuzzy Sets and Systems*, Elsevier, vol. 158, pp. 1064–1082, 2007. T. Bilgiç and I. B. Türkşen, “Measurement of membership functions: theoretical and empirical work,” in *Fundamentals of fuzzy sets*, Springer, 2000, pp. 195–227.
- [53] L. I. Kuncheva, “How Good Are Fuzzy If-Then Classifiers?,” *IEEE TRANSACTIONS ON SYSTEMS, MAN, AND CYBERNETICS—PART B: CYBERNETICS*, VOL. 30, NO. 4, 2000.
- [54] T. Takagi and M. Sugeno, “Fuzzy identification of systems and its applications to modeling and control,” *IEEE transactions on systems, man, and cybernetics*, Art. no. 1, 1985.

- [55] J. Iglesias, P. Angelov, A. Ledezma, and A. Sanchis, “Human activity recognition based on evolving fuzzy systems,” *Int. J. Neural Syst.*, vol. 20, no. 5, pp. 355–364, Oct. 2010.
- [56] W. Pedrycz and F. Gomide, *Fuzzy Systems Engineering: Toward Human-Centric Computing*. Hoboken, NJ, USA: Wiley, 2007.
- [57] A. Lemos, W. Caminhas, and F. Gomide, “Adaptive fault detection and diagnosis using an evolving fuzzy classifier,” *Inf. Sci.*, vol. 220, no. 1, pp. 64–85, Jan. 2013.
- [58] M. Sugeno and G. Kang, “Structure identification of fuzzy model,” *Fuzzy sets and systems*, vol. 28, Art. no. 1, 1988.
- [59] J. Liu, F.-I. Chung, and S. Wang, “Bayesian zero-order TSK fuzzy system modeling,” *Applied Soft Computing*, vol. 55, pp. 253–264, 2017.
- [60] N. Wulandari and A. Abdullah, “Design and Simulation of Washing Machine using Fuzzy Logic Controller (FLC),” *IOP Conference Series: Materials Science and Engineering*, vol. 384, p. 12044, Jul. 2018.
- [61] N. Sabri, S. Aljunid, M. Salim, R. Badlishah, R. Kamaruddin, and M. Malek, “Fuzzy inference system: Short review and design,” *Int. Rev. Autom. Control*, vol. 6, Art. no. 4, 2013.
- [62] S. Oh, Y. Park, K. J. Cho, and S. J. Kim, “Explainable machine learning model for glaucoma diagnosis and its interpretation,” *Diagnostics*, vol. 11, no. 3, p. 510, 2021.
- [63] Y. Niu, L. Gu, Y. Zhao, and F. Lu, “Explainable diabetic retinopathy detection and retinal image generation,” *IEEE Journal of Biomedical and Health Informatics*, vol. 26, no. 1, pp. 44–55, 2022.
- [64] S. M. Shankaranarayana and D. Runje, “Alime: Autoencoder based approach for local interpretability,” *Intelligent Data Engineering and Automated Learning – IDEAL 2019*, pp. 454–463, 2019.
- [65] J. Duell, X. Fan, B. Burnett, G. Aarts, and S.-M. Zhou, “A comparison of explanations given by explainable artificial intelligence methods on analysing electronic health records,” *2021 IEEE EMBS International Conference on Biomedical and Health Informatics (BHI)*, Aug. 2021.
- [66] C.-A. Hu, C.-M. Chen, Y.-C. Fang, S.-J. Liang, H.-C. Wang, W.-F. Fang, C.-C. Sheu, W.-C. Perng, K.-Y. Yang, K.-C. Kao, C.-L. Wu, C.-S. Tsai, M.-Y. Lin, and W.-C. Chao, “Using a machine learning approach to predict mortality in critically ill influenza patients: A cross-sectional retrospective multicentre study in Taiwan,” *BMJ Open*, vol. 10, no. 2, 2020.

- [67] P. Gemmar, “An interpretable mortality prediction model for COVID-19 patients – alternative approach,” 2020.
- [68] M. R. Karim, T. Dohmen, M. Cochez, O. Beyan, D. Rebholz-Schuhmann, and S. Decker, “Deep covid explainer: Explainable covid-19 diagnosis from chest X-ray images,” *2020 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*, 2020.
- [69] K. Duarte, J.-M. Monnez, and E. Albuissou, “Methodology for constructing a short-term event risk score in heart failure patients,” *Applied Mathematics*, vol. 09, no. 08, pp. 954–974, 2018.
- [70] A. Ghorbani, D. Ouyang, A. Abid, B. He, J. Chen, R. Harrington, D. Liang, E. Ashley, and J. Zou, “Deep learning interpretation of echocardiograms,” *NPJ Digit, Med*, 2019.
- [71] M. Athanasiou, K. Sfrintzeri, K. Zarkogianni, A. Thanopoulou, and K. S. Nikita, “An explainable XGBoost-based approach towards assessing the risk of cardiovascular disease in patients with type 2 diabetes mellitus,” 2020.
- [72] D. Zhang, S. Yang, X. Yuan, and P. Zhang, “Interpretable deep learning for automatic diagnosis of 12-lead electrocardiogram,” *iScience*, vol. 24, no. 4, p. 102373, 2021.
- [73] X. Feng, Y. Hua, J. Zou, S. Jia, J. Ji, Y. Xing, J. Zhou, and J. Liao, “Intelligible models for healthcare: Predicting the probability of 6-month unfavorable outcome in patients with ischemic stroke,” *Neuroinformatics*, 2021.
- [74] S. M. Miran, S. J. Nelson, and Q. Zeng-Treitler, “A model-agnostic approach for understanding heart failure risk factors,” *BMC Res. Notes*, 2021.
- [75] P. A. Moreno-Sanchez, “Development of an explainable prediction model of heart failure survival by using ensemble trees,” *2020 IEEE International Conference on Big Data (Big Data)*, 2020.
- [76] T. Ahmad, A. Munir, S. H. Bhatti, M. Aftab, and M. A. Raza, “Survival analysis of heart failure patients: A case study,” *PLOS ONE*, vol. 12, no.7, p. e0181001, Jul. 2017, doi: 10.1371/journal.pone.0181000
- [77] D. Dua and C. Graff, *UCI Machine Learning Repository*. University of California, Irvine, School of Information and Computer Sciences, 2017.
- [78] J. Rashid, S.M.A. Shah, A. Irtaza, Fuzzy topic modeling approach for text mining over short text, *Inf. Process. Manage.* (2019)102060, <https://dl.acm.org/doi/10.1016/j.ipm.2019.102060>.
- [79] S. Abri and R. Abri, “Providing a personalization model based on fuzzy topic modeling,” *Arabian Journal for Science and Engineering*, vol. 46, Art. no. 4, 2021.

- [80] A. Karami, A. Gangopadhyay, B. Zhou, and H. Kharrazi, "Flatm: A fuzzy logic approach topic model for medical documents," in *2015 Annual Conference of the North American Fuzzy Information Processing Society (NAFIPS) held jointly with 2015 5th World Conference on Soft Computing (WConSC)*, 2015, pp. 1–6.
- [81] Attur, S. S., & Harikumar, M. E. (2022, January). Detection of Alzheimer's Disease using Fuzzy Inference System. In *2022 4th International Conference on Smart Systems and Inventive Technology (ICSSIT)* (pp. 1235-1241). IEEE.
- [82] Ahmed, T. I., Bhola, J., Shabaz, M., Singla, J., Rakhra, M., More, S., & Samori, I. A. Fuzzy Logic-Based Systems for the Diagnosis of Chronic Kidney Disease. *BioMed Research International*, 2022.
- [83] Damodara, K., & Thakur, A. (2021, March). Adaptive Neuro Fuzzy Inference System based Prediction of Chronic Kidney Disease. In *2021 7th International Conference on Advanced Computing and Communication Systems (ICACCS)* (Vol. 1, pp. 973-976). IEEE.
- [84] Jayade, S., Ingole, D. T., Ingole, M. D., & Thakare, A. (2021, November). Cholera Disease Detection using Fuzzy Logic Technique. In *2021 3rd International Conference on Electrical, Control and Instrumentation Engineering (ICECIE)* (pp. 1-7). IEEE.
- [85] Alhammadi, F., Alkhanbashi, F., & Shatnawi, M. (2020, December). COVID-19 Fuzzy Inference System. In *2020 International Conference on Computational Science and Computational Intelligence (CSCI)* (pp. 849-852). IEEE.
- [86] Singh, D., Verma, S., & Singla, J. (2021, January). A Neuro-fuzzy based Medical Intelligent System for the Diagnosis of Hepatitis B. In *2021 2nd International Conference on Computation, Automation and Knowledge Management (ICCAKM)* (pp. 107- 111). IEEE.
- [87] Priyadarshini, L., & Shrinivasan, L. (2020, July). Design of an ANFIS based Decision Support System for Diabetes Diagnosis. In *2020 International Conference on Communication and Signal Processing (ICCSP)* (pp. 1486-1489). IEEE.
- [88] Ghosh, A., Rahman, N., Awadalla, N., Sagahyroon, A., Aloul, F., & Dhou, S. Asthma Diagnosis Using Neuro-Fuzzy Techniques. In *2020 Advances in Science and Engineering Technology International Conferences (ASET)* (pp. 1-4). IEEE.
- [89] Iancu, I. (2018). Heart disease diagnosis based on mediative fuzzy logic. *Artificial Intelligence in Medicine*, 89, 51-60
- [90] Muhammad, L. J., & Algehyne, E. A. (2021). Fuzzy based expert system for diagnosis of coronary artery disease in Nigeria. *Health and technology*, 11(2), 319-329.

- [91] Laurentinus., Juniawan, F. P., Sylfania, D. Y., Kurniawan, P., & Pradana, H. A. (2020, October). Design Fuzzy Expert System And Certainty Factor In Early Detection Of Stroke Disease. In *2020 8th International Conference on Cyber and IT Service Management (CITSM)* (pp. 1-7). IEEE.
- [92] S. Srivastava, M. Pant, and N. Agarwal, "A review on role of fuzzy logic in Psychology," *Advances in Intelligent Systems and Computing*, vol. 437, pp. 783–794, 2016.
- [93] A. Hentout, A. Maoudj, and M. Aouache, "A review of the literature on fuzzy-logic approaches for collision-free path planning of manipulator robots," *Artificial Intelligence Review*, 2022.
- [94] S. Hosseinpour and A. Martynenko, "Application of fuzzy logic in drying: A review," *Drying Technology*, vol. 40, Art. no. 5, 2022.
- [95] Z. Lu, J. Wang, R. Mao, M. Lu, and J. Shi, "Jointly Composite Feature Learning and Autism Spectrum Disorder Classification Using Deep Multi-Output Takagi-Sugeno-Kang Fuzzy Inference Systems," *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, p. 1, 2022.
- [96] Y. Gu, K. Xia, K. -W. Lai, Y. Jiang, P. Qian and X. Gu, "Transferable Takagi-Sugeno-Kang Fuzzy Classifier With Multi-Views for EEG-Based Driving Fatigue Recognition in Intelligent Transportation," in *IEEE Transactions on Intelligent Transportation Systems*, 2022, doi: 10.1109/TITS.2022.3220597.
- [97] Z. Boreiri, A. N. Azad and A. Ghodousian, "A Convolutional Neuro-Fuzzy Network Using Fuzzy Image Segmentation for Acute Leukemia Classification," *2022 27th International Computer Conference, Computer Society of Iran (CSICC)*, 2022, pp. 1-7, doi: 10.1109/CSICC55295.2022.9780525.
- [98] Xiaoyi Guo, Yizhang Jiang, Quan Zou, Structured Sparse Regularized TSK Fuzzy System for predicting therapeutic peptides, *Briefings in Bioinformatics*, Volume 23, Issue 3, May 2022, bbac135.
- [99] X. Song, F. Gu, X. Wang, S. Ma, and L. Wang, "Interpretable Recognition for Dementia Using Brain Images," *Frontiers in Neuroscience*, vol. 15, 2021.
- [100] Y. Cui, D. Wu, X. Jiang, and Y. Xu, "PyTSK: A Python Toolbox for TSK Fuzzy Systems." Jun. 2022. doi: 10.48550/arXiv.2206.03310.
- [101] S. Akbarzadeh, A. K. Arof, and T. subramaniam, "Prediction of Conductivity by Adaptive Neuro-Fuzzy Model," *PloS one*, vol. 9, p. e92241, Mar. 2014.

- [102] P. Baldi, “Gradient descent learning algorithm overview: a general dynamical systems perspective,” *IEEE Transactions on Neural Networks*, vol. 6, Art. no. 1, 1995.
- [103] N. Kovač and S. Bauk, “The Anfis based route preference estimation in sea navigation,” *Journal of maritime research: JMR, ISSN 1697-4840, Vol. 3, N°. 3, 2006, pags. 69-86*, Jan. 2023.
- [104] R. Babuška and H. Verbruggen, “Neuro-fuzzy methods for nonlinear system identification,” *Annual Reviews in Control*, vol. 27, pp. 73-85, 2003.
- [105] M. D. Vasilakakis and D. K. Iakovidis, “Fuzzy similarity phrases for interpretable data classification,” *Information Sciences*, vol. 624, pp. 881–907, 2023.
- [106] d. P. Rousseauw J., “Coronary risk factor screening in three rural communities,” *South African Medical Journal* 64, 1983.
- [107] P. Pandey, “Understanding the Mathematics behind Gradient Descent.,” *Towards Data Science*, 2019.
- [108] C. Liu, M. White, and G. Newell, “Measuring and comparing the accuracy of species distribution models with presence-absence data,” *Ecography (Cop.)*, vol. 34, no. 2, pp. 232–243, 2011.
- [109] T. Fawcett, “An introduction to ROC analysis,” *Pattern Recognit. Lett.*, vol. 27, no. 8, pp. 861–874, 2006.
- [110] N. R. S. Katsigiannis, “DREAMER: A Database for Emotion Recognition Through EEG and ECG Signals from Wireless Low-cost Off-the-Shelf Devices,” *IEEE Journal of Biomedical and Health Informatics*, 2018.
- [111] M. M. Bradley and P. J. Lang., “Measur-ing emotion: The self-assessment manikin and thesemantic differential,” *Journal of Behavior Therapyand Experimental Psychiatry*, 1994.
- [112] S. Buechel and U. Hahn, “Readers vs. Writers vs. Texts: Coping with Different Perspectives of Text Understanding in Emotion Annotation,” Jan. 2017.
- [113] Ruchilekha, M. K. Singh, and M. Singh, “A deep learning approach for subject-dependent & subject-independent emotion recognition using brain signals with dimensional emotion model,” *Biomedical Signal Processing and Control*, vol. 84, p. 104928, 2023.
- [114] Z. Du R., “Valence-arousal classification of emotion evoked by Chinese ancient-style music using 1D-CNN-BiLSTM model on EEG signals for college students,” *Multimed Tools Appl*, 2023.
- [115] T. Fan et al., “A new deep convolutional neural network incorporating attentional mechanisms for ECG emotion recognition,” *Computers in Biology and Medicine*, vol. 159, p. 106938, 2023.

- [116] J. E. R. Chy Mohammed Tawsif Khan Nor Azlina Ab Aziz, “Evaluation of Machine Learning Algorithms for Emotions Recognition using Electrocardiogram,” *Emerging Science Journal*, 2023
- [117] d. P. Rousseauw J., “Coronary risk factor screening in three rural communities,” *South African Medical Journal*, 1938.
- [118] K. Cios Krzysztof, “SPECT Heart,” *UCI Machine Learning Repository*, 2001.
- [119] Statlog (Heart). *UCI Machine Learning Repository*. <https://doi.org/10.24432/C57303>.
- [120] F. Pedregosa et al., “Scikit-learn: Machine Learning in Python,” *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011.
- [121] J. Warner et al., “JDWarner/scikit-fuzzy: Scikit-Fuzzy version 0.4.2.” Zenodo, Nov. 2019. doi: 10.5281/zenodo.3541386.
- [122] J. D. Hunter, “Matplotlib: A 2D graphics environment,” *Computing in Science & Engineering*, vol. 9, Art. no. 3, 2007.
- [123] C. R. Harris et al., “Array programming with NumPy,” *Nature*, vol. 585, Art. no. 7825, Sep. 2020.
- [124] W. McKinney, “Data Structures for Statistical Computing in Python,” in *Proceedings of the 9th Python in Science Conference*, 2010, pp. 56–61.
- [125] C. M. T. Khan, N. A. Ab Aziz, J. E. Raja, S. W. B. Nawawi, and P. Rani, “Evaluation of Machine Learning Algorithms for Emotions Recognition using Electrocardiogram,” *Emerging Science Journal*, vol. 7, Art. no. 1, 2022.
- [126] C. J. Burges, “A tutorial on support vector machines for pattern recognition,” *Data mining and knowledge discovery*, vol. 2, Art. no. 2, 1998.
- [127] R. Fletcher, *Practical methods of optimization*. John Wiley & Sons, 2000.
- [128] G. C. Chow, *Dynamic economics: optimization by the Lagrange method*. Oxford University Press, 1997.

