



ΕΘΝΙΚΟ ΜΕΤΣΟΒΕΙΟ ΠΟΛΥΤΕΧΝΕΙΟ

ΣΧΟΛΗ ΗΛΕΚΤΡΟΛΟΓΩΝ ΜΗΧΑΝΙΚΩΝ ΚΑΙ ΜΗΧΑΝΙΚΩΝ ΥΠΟΛΟΓΙΣΤΩΝ

ΤΟΜΕΑΣ ΣΥΣΤΗΜΑΤΩΝ ΜΕΤΑΔΟΣΗΣ ΠΛΗΡΟΦΟΡΙΑΣ ΚΑΙ ΤΕΧΝΟΛΟΓΙΑΣ ΥΛΙΚΩΝ

Τεχνικές Ανάπτυξης Αλγορίθμων Μηχανικής
Μάθησης: Επισκόπηση Μεθόδων και Εφαρμογή σε
Ναυτιλιακά Δεδομένα

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ

Βασίλειος Καραμπλιάς

Επιβλέπων: Χρήστος Καψάλης

Καθηγητής Ε.Μ.Π.

Αθήνα, Σεπτέμβριος 2023



ΕΘΝΙΚΟ ΜΕΤΣΟΒΕΙΟ ΠΟΛΥΤΕΧΝΕΙΟ

ΣΧΟΛΗ ΗΛΕΚΤΡΟΛΟΓΩΝ ΜΗΧΑΝΙΚΩΝ ΚΑΙ ΜΗΧΑΝΙΚΩΝ ΥΠΟΛΟΓΙΣΤΩΝ

ΤΟΜΕΑΣ ΣΥΣΤΗΜΑΤΩΝ ΜΕΤΑΔΟΣΗΣ ΠΛΗΡΟΦΟΡΙΑΣ ΚΑΙ ΤΕΧΝΟΛΟΓΙΑΣ ΥΛΙΚΩΝ

**Τεχνικές Ανάπτυξης Αλγορίθμων Μηχανικής Μάθησης:
Επισκόπηση Μεθόδων και Εφαρμογή σε Ναυτιλιακά Δεδομένα**

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ

Βασίλειος Καραμπλιάς

Επιβλέπων: Χρήστος Καψάλης

Καθηγητής Ε.Μ.Π.

Εγκρίθηκε από την τριμελή εξεταστική επιτροπή τη 16^η Νοεμβρίου 2023

Χρήστος Καψάλης
Καθηγητής Ε.Μ.Π.

Παναγιώτης Κωττής
Καθηγητής Ε.Μ.Π.

Γεώργιος Φικιώρης
Καθηγητής Ε.Μ.Π.

Αθήνα, Σεπτέμβριος 2023

Βασίλειος Ε. Καραμπλιάς

Διπλωματούχος Ηλεκτρολόγος Μηχανικός και Μηχανικός Υπολογιστών Ε.Μ.Π.

Copyright © 2023 Βασίλειος Καραμπλιάς

Με επιφύλαξη παντός δικαιώματος. All rights reserved.

Απαγορεύεται η αντιγραφή, αποθήκευση και διανομή της παρούσας εργασίας, εξ' ολοκλήρου ή τμήματος αυτής, για εμπορικό σκοπό. Επιτρέπεται η ανατύπωση, αποθήκευση και διανομή για σκοπό μη κερδοσκοπικό, εκπαιδευτικής ή ερευνητικής φύσης, υπό την προϋπόθεση να αναφέρεται η πηγή προέλευσης και να διατηρείται το παρόν μήνυμα. Ερωτήματα που αφορούν τη χρήση της εργασίας για κερδοσκοπικό σκοπό πρέπει να απευθύνονται προς τον συγγραφέα. Οι απόψεις και τα συμπεράσματα που περιέχονται σε αυτό το έγγραφο εκφράζουν τον συγγραφέα και δεν πρέπει να ερμηνευθεί ότι αντιπροσωπεύουν τις επίσημες θέσεις του Εθνικού Μετσόβιου Πολυτεχνείου.

Περίληψη

Στόχος της παρούσας εργασίας είναι η εμβάθυνση στο πεδίο της Μηχανικής Μάθησης και των τεχνικών που περιλαμβάνει καθώς και η εφαρμογή αυτών στον τομέα της ναυτιλίας. Στα πρώτα έξι κεφάλαια, οι θεμελιώδεις έννοιες της μηχανικής μάθησης διερευνώνται διεξοδικά και μελετώνται κατηγορίες όπως η επιβλεπόμενη, η μη επιβλεπόμενη και η ενισχυτική μάθηση. Γίνεται σαφής η πολυπλοκότητα της διαδικασίας της εκπαίδευσης των μοντέλων μηχανικής μάθησης με έμφαση στην παλινδρόμηση και την κατηγοριοποίηση.

Ένα σημαντικό μέρος της εργασίας αφορά επισκόπηση αλγορίθμων επιβλεπόμενης μάθησης, όπως η γραμμική και η πολυωνυμική παλινδρόμηση, η παλινδρόμηση με δένδρα αποφάσεων και τυχαία δάση και οι αλγόριθμοι XGBoost και ADABOost. Η εργασία περιλαμβάνει επίσης τεχνικές μη επιβλεπόμενης μάθησης, συμπεριλαμβανομένων αλγορίθμων ομαδοποίησης όπως η K-means ομαδοποίηση και η ιεραρχική ομαδοποίηση καθώς και αλγορίθμων μείωσης διαστάσεων όπως η ανάλυση κύριων συστατικών, η ανάλυση ανεξάρτητων συνιστωσών και η ανάλυση παραγόντων.

Στο έβδομο κεφάλαιο παρουσιάζεται η βασική ιδέα και η δομή των νευρωνικών δικτύων, θεμελιώδεις έννοιες αυτών όπως η συνάρτηση ενεργοποίησης, η συνάρτηση κόστους και η μέθοδος οπισθοδρόμησης ενώ συγκρίνονται και τα απλά με τα βαθιά νευρωνικά δίκτυα πριν την εξέταση του τρόπου εκπαίδευσης και αξιολόγησης της απόδοσης των τελευταίων.

Το τελευταίο κομμάτι της εργασίας περιλαμβάνει μία πρακτική εφαρμογή τεχνικών μηχανικής μάθησης σε ναυτιλιακά δεδομένα. Αλγόριθμοι μηχανικής μάθησης εφαρμόζονται για την πρόβλεψη της κατανάλωσης καυσίμου, μία παράμετρος υψίστης σημασίας στις θαλάσσιες μεταφορές. Μία νέα προσέγγιση, η Ομοσπονδιακή Μάθηση, εισάγεται και αντιπαρατίθεται με άλλες μεθοδολογίες για την αξιολόγηση της ακρίβειας και της πολυπλοκότητάς της στο περιβάλλον της ναυτιλίας. Μέσω αυτής της σύγκρισης αποκτώνται πρακτικές γνώσεις, και επιλέγονται βέλτιστες στρατηγικές μηχανικής μάθησης για την πρόβλεψη της κατανάλωσης καυσίμου στις θαλάσσιες επιχειρήσεις. Επιπλέον, εντοπίζονται ανοικτά ζητήματα και παρατίθενται τα σχετικά συμπεράσματα και μελλοντικές προεκτάσεις.

Λέξεις κλειδιά: Μηχανική Μάθηση, Επιβλεπόμενη Μάθηση, Μη Επιβλεπόμενη Μάθηση, Ενισχυτική Μάθηση, Νευρωνικά Δίκτυα, Βαθιά Μάθηση, Εκπαίδευση Μοντέλου, Πρόβλεψη Κατανάλωσης Καυσίμου

Abstract

The aim of the present study is to delve into the dynamic realm of Machine Learning techniques and their diverse applications in maritime. Across the initial six chapters, the fundamental concepts of Machine Learning are explored comprehensively, covering supervised and unsupervised learning methodologies. The study navigates through the complexities of Machine Learning model training, emphasizing regression and classification.

A significant portion of the study is dedicated to supervised learning algorithms, such as linear and polynomial regression, decision tree and random forest regression and ensemble methods like XGBoost and ADABOOST as well as unsupervised learning techniques, including clustering algorithms like k-means and hierarchical clustering and dimensionality reduction algorithms like Principal Component Analysis, Independent Component Analysis and Factor Analysis.

In Chapter 7, the fundamental principles and architectural foundations of neural networks are meticulously elucidated. The chapter dissects essential concepts including activation functions, cost functions and backpropagation. A comparative analysis between simple and deep neural networks is conducted, setting the stage for an exploration into the training methodologies and performance evaluation focusing on the latter.

In Chapter 8, the focus shifts to the practical implementation of Machine Learning techniques within the specific context of maritime environments. Here, the study delves into the integration of Machine Learning algorithms for fuel consumption prediction—a critical concern for maritime sustainability. Federated Learning, a novel approach, is introduced and juxtaposed with other methodologies to assess its accuracy and complexity in the unique maritime setting. Through rigorous comparison, practical insights are gained, shedding light on the optimal Machine Learning strategies for fuel consumption prediction in maritime operations. Additionally, the chapter identifies open issues, paving the way for future research and innovation in the domain.

Keywords: Machine Learning, Supervised Learning, Unsupervised Learning, Reinforcement Learning, Neural Networks, Deep Learning, Model Training, Fuel Consumption Prediction

Ευχαριστίες

Η διπλωματική αυτή εργασία εκπονήθηκε στα πλαίσια των ερευνητικών δραστηριοτήτων του Τομέα Συστημάτων Μετάδοσης Πληροφορίας και Τεχνολογίας Υλικών της Σχολής Ηλεκτρολόγων Μηχανικών και Μηχανικών Υπολογιστών του Εθνικού Μετσόβιου Πολυτεχνείου κατά το ακαδημαϊκό έτος 2022-2023.

Αρχικά, θα ήθελα να ευχαριστήσω τον Καθηγητή κ. Χρήστο Καψάλη για την ευκαιρία που μου έδωσε να ασχοληθώ με το αντικείμενο της Μηχανικής Μάθησης και τη συγκεκριμένη εργασία, καθώς και τους Καθηγητές κ. Παναγιώτη Κωττή και κ. Γεώργιο Φικιώρη για τη συμμετοχή τους στην τριμελή εξεταστική επιτροπή.

Επίσης, θα ήθελα να ευχαριστήσω τον υποψήφιο διδάκτορα της Σχολής Ηλεκτρολόγων Μηχανικών και Μηχανικών Υπολογιστών του Εθνικού Μετσόβιου Πολυτεχνείου Αναστάσιο Γιαννόπουλο για την παρακολούθηση και τις χρήσιμες συμβουλές που προσέφερε καθ' όλη τη διάρκεια εκπόνησης της εργασίας.

Τέλος, θα ήθελα να ευχαριστήσω την οικογένειά μου και τους φίλους μου, που είναι κοντά μου και με στηρίζουν όλα αυτά τα χρόνια.

Καραμπλιάς Βασίλειος, Σεπτέμβριος 2023

Πίνακας Περιεχομένων

Κεφάλαιο 1: Αντικείμενο και Δομή της εργασίας.....	16
Κεφάλαιο 2: Μηχανική Μάθηση: Βασικές έννοιες και ορισμοί.....	18
2.1 Τι είναι Μηχανική Μάθηση.....	18
2.2 Στόχοι Μηχανικής Μάθησης.....	18
2.3 Είδη Μηχανικής Μάθησης.....	18
2.4 Πλεονεκτήματα και Μειονεκτήματα Μηχανικής Μάθησης.....	19
2.4.1 Πλεονεκτήματα.....	19
2.4.2 Μειονεκτήματα.....	20
2.4.3 Γενικό Συμπέρασμα.....	20
2.5 Μοντέλο Μηχανικής Μάθησης.....	20
2.6 Εκπαίδευση μοντέλων Μηχανικής Μάθησης.....	21
2.7 Εφαρμογές Μηχανικής Μάθησης.....	22
Κεφάλαιο 3: Ιδέα Μηχανικής Μάθησης.....	24
3.1 Συλλογή δεδομένων.....	24
3.2 Προεπεξεργασία δεδομένων.....	24
3.3 Εκπαίδευση μοντέλου.....	25
3.4 Κλήση του μοντέλου σε άγνωστα δεδομένα.....	26
3.5 Υπερπροσαρμογή και Υποπροσαρμογή.....	26
Κεφάλαιο 4: Κατηγορίες αλγορίθμων Μηχανικής Μάθησης.....	28
4.1 Παλινδρόμηση και Κατηγοριοποίηση.....	28
4.2 Επιβλεπόμενη Μάθηση, Μη Επιβλεπόμενη Μάθηση και Ενισχυτική Μάθηση.....	29
Κεφάλαιο 5: Επιβλεπόμενη Μάθηση.....	30
5.1 Ορισμός και Παραδείγματα.....	30
5.2 Κατηγορίες αλγορίθμων Επιβλεπόμενης Μάθησης.....	31
5.2.1 Γραμμική παλινδρόμηση.....	31
5.2.2 Πολλαπλή γραμμική παλινδρόμηση.....	32
5.2.3 Πολυωνυμική παλινδρόμηση.....	32
5.2.4 Μέθοδος των k-πλησιέστερων γειτόνων.....	33
5.2.5 Παλινδρόμηση με διανύσματα υποστήριξης.....	34
5.2.6 Παλινδρόμηση με δένδρα αποφάσεων.....	36
5.2.7 Παλινδρόμηση με τον αλγόριθμο τυχαίου δάσους.....	37

5.2.8 Οι Αλγόριθμοι XGBoost και AdaBoost	38
Κεφάλαιο 6: Μη Επιβλεπόμενη Μάθηση	40
6.1 Ορισμός και Παραδείγματα	40
6.2 Ομαδοποίηση	40
6.2.1 Η μέθοδος ομαδοποίησης K-means.....	40
6.2.2 Ο αλγόριθμος ομαδοποίησης k-πλησιέστερων γειτόνων	41
6.2.3 Ιεραρχική ομαδοποίηση.....	41
6.2.4 Ο αλγόριθμος ομαδοποίησης DBSCAN.....	42
6.2.5 Μοντέλα Gaussian Mixture.....	43
6.3 Μείωση Διαστάσεων	44
6.3.1 Ανάλυση κύριων συνιστωσών	44
6.3.2 Ανάλυση ανεξάρτητων συνιστωσών.....	45
6.3.3 Γραμμική διακριτική ανάλυση	45
6.3.4 Ανάλυση παραγόντων	46
Κεφάλαιο 7: Νευρωνικά Δίκτυα.....	48
7.1 Εισαγωγή.....	48
7.2 Δομή νευρωνικού δικτύου	49
7.3 Απλά και βαθιά νευρωνικά δίκτυα.....	49
7.4 Δομή νευρώνα.....	50
7.5 Συνάρτηση κόστους	51
7.6 Συνάρτηση ενεργοποίησης.....	52
7.7 Μέθοδος οπισθοδρόμησης	53
7.8 Εκπαίδευση βαθιών νευρωνικών δικτύων	54
7.9 Αξιολόγηση απόδοσης βαθιών νευρωνικών δικτύων	54
Κεφάλαιο 8: Πειραματικό μέρος	56
8.1 Εισαγωγή.....	56
8.2 Πειραματικά αποτελέσματα	57
8.2.1 Περιγραφή συνόλου δεδομένων	57
8.2.2 Προεπεξεργασία και Μείωση Διαστάσεων	58
8.2.3 Ρύθμιση υπερπαραμέτρων μάθησης	59
8.2.4 Αναδρομική εξάλειψη χαρακτηριστικών	61
8.2.5 Σύγκριση απόδοσης-πολυπλοκότητας	63
8.2.6 Ζητήματα σχετικά με την εκπαίδευση και την ανάπτυξη των συστημάτων	66

8.3 Ανοικτά ζητήματα	67
8.3.1 Υλοποίηση με δεδομένα μεγάλης κλίμακας με 6G δίκτυα.....	67
8.3.2 Μεγάλες αποστάσεις διάδοσης και ποικίλα χαρακτηριστικά καναλιών	68
8.3.3 Μη πανομοιότυπα δεδομένα εκπαίδευσης προερχόμενα από διαφορετικές πηγές ...	68
8.3.4 Αξιολόγηση απόδοσης πλήρους κλίμακας.....	68
8.3.5 Κεντρικά έναντι κατανεμημένων συστημάτων Ομοσπονδιακής Μάθησης.....	69
8.3.6 Ερμηνευσιμότητα.....	69

Κατάλογος Σχημάτων

Σχήμα 1: Γραμμική παλινδρόμηση.....	31
Σχήμα 2: Γραμμική παλινδρόμηση (αριστερά), Πολυωνυμική παλινδρόμηση (δεξιά)	33
Σχήμα 3: K-Πλησιέστεροι Γείτονες	34
Σχήμα 4: Σημεία δεδομένων στον χώρο χαμηλής διάστασης	35
Σχήμα 5: Σημεία δεδομένων σε χώρο υψηλότερης διάστασης.....	35
Σχήμα 6: Μη γραμμικό μοτίβο στον χώρο χαμηλής διάστασης	35
Σχήμα 7: Παλινδρόμηση με διανύσματα υποστήριξης	36
Σχήμα 8: Παλινδρόμηση με δένδρα αποφάσεων	37
Σχήμα 9: Παλινδρόμηση με τον αλγόριθμο τυχαίου δάσους.....	38
Σχήμα 10: K-means ομαδοποίηση.....	41
Σχήμα 11: Ιεραρχική ομαδοποίηση	42
Σχήμα 12: K-means ομαδοποίηση, Ιεραρχική ομαδοποίηση, DBSCAN (από τα αριστερά προς τα δεξιά)	43
Σχήμα 13: Μοντέλο Gaussian Mixture	43
Σχήμα 14: Ανάλυση κύριων συνιστωσών	45
Σχήμα 15: Γραμμική διακριτική ανάλυση	46
Σχήμα 16: Ανάλυση παραγόντων.....	47
Σχήμα 17: Δομή νευρωνικού δικτύου.....	49
Σχήμα 18: Δομή νευρώνα.....	51
Σχήμα 19: Σιγμοειδής συνάρτηση	52
Σχήμα 20: Συνάρτηση ReLU	53
Σχήμα 21: Υπερβολική εφαπτομένη	53
Σχήμα 22: Απόδοση του μοντέλου συναρτήσει του βάθους του νευρωνικού δικτύου για τις διάφορες τιμές του ρυθμού μάθησης α	61
Σχήμα 23: Ταξινομημένη F-Value των χαρακτηριστικών όπως προέκυψε από τον αλγόριθμο XGBoost	62
Σχήμα 24: Απόδοση μοντέλου για διαφορετικό πλήθος μεταβλητών εισόδου	63
Σχήμα 25: Κανονικοποιημένη ακρίβεια (αριστερός κάθετος άξονας) και πολυπλοκότητα (δεξιός κάθετος άξονας) για τις διάφορες μεθόδους	66
Σχήμα 26: Πραγματική (actual) τιμή και προβλεφθείσες σύμφωνα με την CLI και FL μέθοδο της μεταβλητής MEP	66

Κεφάλαιο 1: Αντικείμενο και Δομή της εργασίας

Η σημασία των αλγορίθμων Μηχανικής Μάθησης έχει καταστεί πλέον μη αμφισβητήσιμη και η παρουσία τους σε ένα μεγάλο εύρος πεδίων είναι πλέον έντονη. Η επίλυση προβλημάτων με ευφυείς προσεγγίσεις βασίζεται στη δυνατότητα των συστημάτων να μαθαίνουν και η αποτελεσματικότητα των τεχνικών μηχανικής μάθησης στην επίλυση προβλημάτων καθιστά την αξιοποίηση τους απαραίτητη.

Ένα τέτοιο πεδίο αποτελεί και η βιομηχανία της ναυτιλίας, σημαντικό κομμάτι της παγκόσμιας οικονομίας το οποίο αντιμετωπίζει προκλήσεις σχετικά με τον περιορισμό των λειτουργικών εξόδων αλλά και τη μείωση των περιβαλλοντικών του επιπτώσεων. Για την αντιμετώπιση αυτών των ζητημάτων, προηγμένες τεχνολογίες όπως οι αλγόριθμοι μηχανικής μάθησης και, ιδίως, τα τεχνητά νευρωνικά δίκτυα (Artificial Neural Networks, ANNs), κατάλληλα για την επεξεργασία μεγάλου όγκου δεδομένων, έχουν αναδειχθεί ως ισχυρά εργαλεία.

Στην παρούσα εργασία, στο 2^ο Κεφάλαιο ορίζεται η Μηχανική Μάθηση, παρουσιάζονται τα διάφορα είδη και οι στόχοι της καθώς και τα πλεονεκτήματα και τα μειονεκτήματά της ενώ καταδεικνύεται και η σημασία της μέσα από την πληθώρα των εφαρμογών της.

Στο 3^ο Κεφάλαιο αναλύεται η διαδικασία συλλογής, επεξεργασίας και εκμετάλλευσης των δεδομένων κατά την εκπαίδευση και τον έλεγχο των μοντέλων μηχανικής μάθησης και επεξηγούνται βασικοί όροι όπως η Υπερπροσαρμογή και η Υποπροσαρμογή.

Στη συνέχεια παρατίθενται στο 4^ο Κεφάλαιο οι διάφορες κατηγορίες μηχανικής μάθησης όπως η επιβλεπόμενη, η μη επιβλεπόμενη και η ενισχυτική μάθηση και γίνεται διαχωρισμός μεταξύ των αλγορίθμων που χρησιμοποιούνται για παλινδρόμηση και αυτών που χρησιμοποιούνται για κατηγοριοποίηση.

Το 5^ο Κεφάλαιο εμβαθύνει στην επιβλεπόμενη μάθηση και γίνεται επισκόπηση των σημαντικότερων αλγορίθμων που περιλαμβάνει όπως η γραμμική και η πολυωνυμική παλινδρόμηση, η μέθοδος των K-πλησιέστερων γειτόνων (k-NN), η παλινδρόμηση με διανύσματα υποστήριξης, δένδρα αποφάσεων και τυχαία δάση καθώς και οι αλγόριθμοι XGBoost και AdaBoost.

Το 6^ο Κεφάλαιο επικεντρώνεται στους αλγορίθμους μη επιβλεπόμενης μάθησης. Γίνεται επισκόπηση τόσο αλγορίθμων ομαδοποίησης (μέθοδος k-means, αλγόριθμος k-NN στη μη επιβλεπόμενη μάθηση, ιεραρχική ομαδοποίηση, αλγόριθμος DBSCAN, μοντέλα Gaussian Mixture) όσο και αλγορίθμων μείωσης διαστάσεων (ανάλυση κύριων και ανάλυση ανεξάρτητων συνιστωσών, γραμμική διακριτική ανάλυση, ανάλυση παραγόντων).

Στη συνέχεια, στο 7^ο Κεφάλαιο γίνεται μία εισαγωγή στο πεδίο των Νευρωνικών Δικτύων και παρουσιάζεται η δομή τους, καθώς και η δομή των μεμονωμένων νευρώνων που το αποτελούν ενώ γίνεται και σύγκριση μεταξύ των απλών και των βαθιών νευρωνικών δικτύων.

Μελετώνται βασικά εργαλεία των νευρωνικών δικτύων όπως η συνάρτηση κόστους και η συνάρτηση ενεργοποίησης και εξηγείται η μέθοδος οπισθοδρόμησης αλλά και η διαδικασία εκπαίδευσης και αξιολόγησης της απόδοσης των βαθιών νευρωνικών δικτύων.

Μετά την κατανόηση όλων των παραπάνω εννοιών, ακολουθεί στο 8^ο Κεφάλαιο το πειραματικό μέρος της εργασίας, κατά το οποίο παρουσιάζεται μία πρακτική εφαρμογή της Μηχανικής Μάθησης με πραγματικά δεδομένα μίας ναυτιλιακής επιχείρησης με στόχο την εύρεση του αλγορίθμου ο οποίος διατηρεί μια ισορροπία μεταξύ πολυπλοκότητας και απόδοσης του μοντέλου και με στόχο τον όσο το δυνατόν μεγαλύτερο περιορισμό της κατανάλωσης καυσίμου. Παρουσιάζεται η έννοια της Ομοσπονδιακής Μάθησης, συγκρίνεται με άλλες τεχνικές με βάση την πολυπλοκότητα και την απόδοσή της στη συγκεκριμένη εφαρμογή και, τέλος, αναφέρονται ορισμένα ανοιχτά ζητήματα σε ό,τι αφορά τη χρήση των τεχνικών αυτών στο ναυτιλιακό περιβάλλον.

Κεφάλαιο 2: Μηχανική Μάθηση: Βασικές έννοιες και ορισμοί

2.1 Τι είναι Μηχανική Μάθηση

Η Μηχανική Μάθηση (Machine Learning) είναι μια συλλογή εργαλείων και τεχνικών που μετατρέπει τα δεδομένα σε αποφάσεις κάνοντας α) κατηγοριοποιήσεις και β) ποσοτικές προβλέψεις. Στην περίπτωση των κατηγοριοποιήσεων η διαδικασία λέγεται Κατηγοριοποίηση (Classification) ενώ στις ποσοτικές προβλέψεις (quantitative predictions) λέγεται Παλινδρόμηση (Regression).

2.2 Στόχοι Μηχανικής Μάθησης

Ο κύριος σκοπός της μηχανικής μάθησης είναι να επιτρέπει στους υπολογιστές να μαθαίνουν από δεδομένα και να κάνουν προβλέψεις ή να λαμβάνουν αποφάσεις χωρίς να προγραμματίζονται ρητά για την εκτέλεση της εκάστοτε εργασίας.[1,3] Η μηχανική μάθηση χρησιμοποιείται σε ένα ευρύ φάσμα εφαρμογών, συμπεριλαμβανομένης της αναγνώρισης εικόνας και ομιλίας, της επεξεργασίας φυσικής γλώσσας, της ανίχνευσης ανωμαλιών, των συστημάτων συστάσεων και της λήψης αποφάσεων. Γενικά, η μηχανική μάθηση χρησιμοποιείται για την αυτοματοποίηση εργασιών των οποίων η χειροκίνητη εκτέλεση θα ήταν χρονοβόρα, δαπανηρή ή και ανέφικτη για τους ανθρώπους με σκοπό τη βελτίωση της αποτελεσματικότητας, τη μείωση των σφαλμάτων και την απόκτηση γνώσεων από τα δεδομένα.

2.3 Είδη Μηχανικής Μάθησης

Υπάρχουν τρία κύρια είδη μηχανικής μάθησης: επιβλεπόμενη μάθηση (supervised learning), μη επιβλεπόμενη μάθηση (unsupervised learning) και ενισχυτική μάθηση (reinforcement learning).

Επιβλεπόμενη μάθηση: στην επιβλεπόμενη μάθηση, ένα μοντέλο εκπαιδεύεται σε επισημασμένα δεδομένα (labeled data), πράγμα που σημαίνει ότι εκτός από κάθε είσοδο παρέχεται και η αντίστοιχη σωστή έξοδος. Στη συνέχεια, το μοντέλο κάνει προβλέψεις για νέα, άγνωστα δεδομένα με βάση αυτά που έχει μάθει από τα δεδομένα εκπαίδευσης (training data). Παραδείγματα επιβλεπόμενης μάθησης περιλαμβάνουν τη γραμμική παλινδρόμηση (linear regression), τη λογιστική παλινδρόμηση (logistic regression) και τα δέντρα αποφάσεων (decision trees).

Μη επιβλεπόμενη μάθηση: στη μη επιβλεπόμενη μάθηση, στο μοντέλο δεν παρέχονται επισημασμένα δεδομένα. Αντ' αυτού, το ίδιο το μοντέλο αναζητεί μοτίβα και δομή στα δεδομένα εισόδου από μόνο του. Παραδείγματα μη επιβλεπόμενης μάθησης περιλαμβάνουν την k-means ομαδοποίηση (k-means clustering) και την ανάλυση κύριων συνιστωσών (Principal Component Analysis, PCA).

Ενισχυτική μάθηση: στην ενισχυτική μάθηση, το μοντέλο μαθαίνει να λαμβάνει αποφάσεις αλληλεπιδρώντας με το περιβάλλον του και λαμβάνοντας ανατροφοδότηση με τη μορφή ανταμοιβών ή ποινών. Ο στόχος είναι η υιοθέτηση από το μοντέλο μιας πολιτικής η οποία μεγιστοποιεί τη συνολική ανταμοιβή με την πάροδο του χρόνου. Η ενισχυτική μάθηση χρησιμοποιείται συχνά στη ρομποτική, στη θεωρία παιγνίων και στη λήψη αποφάσεων.

Αξίζει να σημειωθεί ότι ορισμένα μοντέλα μπορεί να ανήκουν σε περισσότερες από μία κατηγορίες και ορισμένοι αλγόριθμοι μπορούν να προσαρμοστούν σε διαφορετικά είδη μάθησης με την κατάλληλη προσαρμογή και των δεδομένων εισόδου και εξόδου.

2.4 Πλεονεκτήματα και Μειονεκτήματα Μηχανικής Μάθησης

2.4.1 Πλεονεκτήματα

Τα σημαντικότερα πλεονεκτήματα της μηχανικής μάθησης αφορούν τους παρακάτω τομείς:

Αυτοματοποίηση: Η μηχανική μάθηση μπορεί να αυτοματοποιήσει εργασίες που διαφορετικά θα απαιτούσαν ανθρώπινη παρέμβαση, όπως αναγνώριση εικόνας ή επεξεργασία φυσικής γλώσσας.

Αποδοτικότητα: Τα μοντέλα μηχανικής μάθησης μπορούν να επεξεργάζονται και να αναλύουν μεγάλες ποσότητες δεδομένων γρήγορα και με ακρίβεια, καθιστώντας τα πιο αποτελεσματικά από τους ανθρώπους σε ορισμένες εργασίες.

Εξατομίκευση: Η μηχανική μάθηση μπορεί να χρησιμοποιηθεί για την παροχή εξατομικευμένων συστάσεων, γεγονός που την καθιστά πολύ σημαντική στον τομέα των συστημάτων συστάσεων.

Προβλέψεις: Τα μοντέλα μηχανικής μάθησης μπορούν να χρησιμοποιηθούν για να κάνουν προβλέψεις σε διάφορους τομείς όπως στις τεχνικές προβλέψεις, τον εντοπισμό και την αντιμετώπιση πιθανής απόπειρας για απάτη, την πρόβλεψη του καιρού κ.ά.

Συνεχής μάθηση: Τα μοντέλα μηχανικής μάθησης με την πάροδο του χρόνου και την εκπαίδευση σε νέα δεδομένα μαθαίνουν, εξελίσσονται και βελτιώνονται συνεχώς.

2.4.2 Μειονεκτήματα

Από την άλλη πλευρά, η μηχανική μάθηση περιλαμβάνει και μειονεκτήματα ή δυσκολίες σε τομείς όπως οι παρακάτω:

Πολυπλοκότητα: Η μηχανική μάθηση αποτελεί ένα πολύπλοκο πεδίο που απαιτεί εξειδικευμένες γνώσεις και δεξιότητες για την υλοποίηση, την εφαρμογή και τη συντήρησή του.

Προκατάληψη δεδομένων (Data Bias): Τα μοντέλα μηχανικής μάθησης είναι τόσο καλά όσο τα δεδομένα στα οποία εκπαιδεύονται και εάν τα δεδομένα δεν είναι αντιπροσωπευτικά, το μοντέλο θα είναι προκατειλημμένο (biased). Γενικά η ποιότητα των δεδομένων που χρησιμοποιούνται για την εκπαίδευση (training) και τις δοκιμές (testing) επηρεάζει σε μεγάλο βαθμό και την απόδοση του μοντέλου.[7]

Περιορισμένη δυνατότητα ερμηνείας και κατανόησης: Ορισμένα μοντέλα μηχανικής μάθησης, όπως τα Νευρωνικά Δίκτυα Βαθιάς Μάθησης (Deep Learning Neural Networks), μπορεί να είναι δύσκολο να ερμηνευθούν και να κατανοηθεί ο τρόπος με τον οποίο κατέληξαν στις προβλέψεις τους.

Υψηλές υπολογιστικές απαιτήσεις: Η υψηλή πολυπλοκότητα και το μεγάλο πλήθος δεδομένων έχουν ως συνέπεια ορισμένα μοντέλα μηχανικής μάθησης να απαιτούν μεγάλη υπολογιστική ισχύ και πόρους.

Ηθικές ανησυχίες: Υπάρχουν ορισμένες ηθικές ανησυχίες γύρω από τη μηχανική μάθηση που σχετίζονται με τη δυνατότητα των μοντέλων να δικαιωνίζουν κοινωνικές προκαταλήψεις, να χρησιμοποιούνται για παρακολούθηση ή άλλους επεμβατικούς σκοπούς.

2.4.3 Γενικό Συμπέρασμα

Συμπερασματικά, η Μηχανική Μάθηση είναι ένας τομέας με πολλά πιθανά οφέλη και προοπτικές σημαντικής τεχνολογικής εξέλιξης υπό την προϋπόθεση ότι τόσο οι περιορισμοί και οι απαιτήσεις του εκάστοτε μοντέλου, όσο και οι όροι ηθικής και υπεύθυνης χρήσης αυτού θα λαμβάνονται σοβαρά υπόψη.

2.5 Μοντέλο Μηχανικής Μάθησης

Ένα Μοντέλο Μηχανικής Μάθησης είναι μια μαθηματική αναπαράσταση ενός συστήματος ή μιας διαδικασίας που εκπαιδεύεται από δεδομένα ώστε να μπορεί να πραγματοποιήσει προβλέψεις ή να λάβει αποφάσεις χωρίς να προγραμματίζεται ρητά για την εκτέλεση της

εκάστοτε εργασίας. Το μοντέλο αποτελείται ουσιαστικά από ένα σύνολο αλγορίθμων που χρησιμοποιούνται για την ανάλυση των δεδομένων και την πραγματοποίηση προβλέψεων από την ανάλυση αυτή. Υπάρχουν πολλοί τύποι μοντέλων μηχανικής μάθησης, ο καθένας με τα δικά του πλεονεκτήματα αλλά και αδυναμίες. Ορισμένοι δημοφιλείς τύποι μοντέλων περιλαμβάνουν τα εξής:

Μοντέλα γραμμικής παλινδρόμησης, τα οποία χρησιμοποιούνται για την πρόβλεψη συνεχών τιμών.

Μοντέλα λογιστικής παλινδρόμησης, τα οποία χρησιμοποιούνται για την πρόβλεψη αποτελεσμάτων δυαδικών τιμών.

Δέντρα αποφάσεων, τα οποία χρησιμοποιούνται τόσο για προβλήματα κατηγοριοποίησης όσο και για προβλήματα παλινδρόμησης.

Νευρωνικά δίκτυα, τα οποία χρησιμοποιούνται για ένα ευρύ φάσμα εργασιών και είναι ιδιαίτερα αποδοτικά στην αναγνώριση εικόνας και ομιλίας.

Μοντέλα ομαδοποίησης, τα οποία χρησιμοποιούνται για την εύρεση μοτίβων και την ομαδοποίηση παρόμοιων περιπτώσεων στα δεδομένα.

Μηχανές διανυσμάτων υποστήριξης (Support Vector Machines, SVMs), οι οποίες χρησιμοποιούνται για προβλήματα κατηγοριοποίησης.

Κάθε μοντέλο εκπαιδεύεται με χρήση ενός συγκεκριμένου συνόλου αλγορίθμων και παραμέτρων και η επιλογή του καταλληλότερου μοντέλου εξαρτάται από το προς επίλυση πρόβλημα, τον τύπο των δεδομένων και τους διαθέσιμους πόρους. Όταν το μοντέλο εκπαιδευτεί, μπορεί να χρησιμοποιηθεί για την πραγματοποίηση προβλέψεων σε νέα, άγνωστα για το μοντέλο δεδομένα. Η ακρίβεια και η ποιότητα των προβλέψεων που πραγματοποιούνται από ένα μοντέλο καθορίζεται σε μεγάλο βαθμό από την αξιολόγηση της απόδοσης του μοντέλου σε ένα ξεχωριστό σύνολο δεδομένων τα οποία ονομάζονται δεδομένα δοκιμών (testing data).

2.6 Εκπαίδευση μοντέλων Μηχανικής Μάθησης

Η εκπαίδευση μοντέλων είναι η διαδικασία αξιοποίησης ενός συνόλου δεδομένων εισόδου και εξόδου, γνωστών ως δεδομένων εκπαίδευσης, για την προσαρμογή των παραμέτρων ενός μοντέλου μηχανικής μάθησης έτσι ώστε αυτό να μπορεί να πραγματοποιεί ακριβείς προβλέψεις. Η διαδικασία εκπαίδευσης ενός μοντέλου περιλαμβάνει συνήθως τα ακόλουθα βήματα:

- Επιλογή ενός κατάλληλου για την εκάστοτε εργασία μοντέλου, απόφαση που εξαρτάται κυρίως από τον τύπο του προς επίλυση προβλήματος, τον τύπο των δεδομένων καθώς και τους διαθέσιμους πόρους.

- Προετοιμασία των δεδομένων, διαδικασία που περιλαμβάνει την προεπεξεργασία των δεδομένων, καθώς και τον διαχωρισμό τους σε ένα σύνολο δεδομένων εκπαίδευσης και ένα σύνολο δεδομένων δοκιμών.
- Εκπαίδευση του μοντέλου: μόλις προετοιμαστούν τα δεδομένα κατάλληλα, το μοντέλο μπορεί να εκπαιδευτεί με χρήση ενός συνόλου αλγορίθμων και παραμέτρων. Η διαδικασία εκπαίδευσης περιλαμβάνει την ενσωμάτωση των δεδομένων εισόδου στο μοντέλο και την προσαρμογή των παραμέτρων του έτσι ώστε το μοντέλο να μπορεί να πραγματοποιεί ακριβείς προβλέψεις για τα δεδομένα εξόδου.
- Αξιολόγηση της απόδοσης του μοντέλου, διαδικασία που απαιτεί την αξιοποίηση ενός διαφορετικού συνόλου δεδομένων, των δεδομένων δοκιμών, έτσι ώστε να διασφαλιστεί ότι η εκπαίδευση του μοντέλου πάνω στα δεδομένα εκπαίδευσης έγινε με τρόπο που επιτρέπει στο μοντέλο να πραγματοποιεί εξίσου ακριβείς προβλέψεις και για διαφορετικά, άγνωστα δεδομένα.
- Βελτιστοποίηση του μοντέλου, διαδικασία που περιλαμβάνει τη ρύθμιση των παραμέτρων του μοντέλου και την επανεκπαίδευση αυτού προκειμένου να βελτιωθεί η συνολική απόδοσή του.

Συμπερασματικά, αξίζει να σημειωθεί ότι η εκπαίδευση ενός μοντέλου μπορεί να είναι μια υπολογιστικά επίπονη και χρονοβόρα διαδικασία, ανάλογη της πολυπλοκότητας του μοντέλου και του μεγέθους των δεδομένων εκπαίδευσης, ενώ η απόδοσή του εξαρτάται άμεσα από την ποιότητα των δεδομένων.

2.7 Εφαρμογές Μηχανικής Μάθησης

Η μηχανική μάθηση είναι ένας συνεχώς εξελισσόμενος τομέας του οποίου η σημαντικότητα γίνεται αντιληπτή από το μεγάλο εύρος των πεδίων στα οποία βρίσκει εφαρμογή και τις νέες δυνατότητες που προσφέρει σε κάθε ένα από αυτά. Ενδεικτικά, αναφέρονται παρακάτω ορισμένες εφαρμογές της σε ένα πλήθος διαφορετικών πεδίων.

Αναγνώριση εικόνας: σε αυτήν την εφαρμογή, ένα μοντέλο μηχανικής μάθησης εκπαιδεύεται ώστε να μπορεί να αναγνωρίζει και να κατηγοριοποιεί αντικείμενα μέσα σε μια εικόνα. Για παράδειγμα, ένα μοντέλο μπορεί να εκπαιδευτεί ώστε να αναγνωρίζει γάτες και σκύλους ή πρόσωπα σε φωτογραφίες.

Επεξεργασία φυσικής γλώσσας: τα μοντέλα μηχανικής μάθησης μπορούν να χρησιμοποιηθούν για την επεξεργασία και την κατανόηση της ανθρώπινης γλώσσας. Για παράδειγμα, ένα μοντέλο μπορεί να χρησιμοποιηθεί για την αυτόματη δημιουργία λεζάντων σε εικόνες ή για τη βελτίωση της ακρίβειας ενός λογισμικού αναγνώρισης ομιλίας και καταγραφής αυτής.

Συστήματα συστάσεων: τα συστήματα αυτά χρησιμοποιούν τη μηχανική μάθηση για να κάνουν εξατομικευμένες συστάσεις στους χρήστες κάποιας εφαρμογής ή ιστοτόπου. Για

παράδειγμα, ένα σύστημα συστάσεων μπορεί να προτείνει προϊόντα σε πελάτες σε έναν ιστότοπο ηλεκτρονικού εμπορίου με βάση το ιστορικό των αναζητήσεων τους.

Ανίχνευση απάτης: τα μοντέλα μηχανικής μάθησης μπορούν να χρησιμοποιηθούν για τον εντοπισμό παράνομων δραστηριοτήτων σε χρηματοοικονομικές συναλλαγές. Για παράδειγμα, ένα μοντέλο μπορεί να εκπαιδευτεί για να εντοπίζει μοτίβα σε ύποπτες συναλλαγές με πιστωτικές κάρτες και να τις αποτρέπει.

Αυτοοδηγούμενα αυτοκίνητα: η μηχανική μάθηση διαδραματίζει κρίσιμο ρόλο στην ανάπτυξη αυτοοδηγούμενων αυτοκινήτων καθώς τα μοντέλα εκπαιδεύονται για να κατανοούν το περιβάλλον και να λαμβάνουν αποφάσεις με τη χρήση αισθητήρων, κάτι που αναμένεται να οδηγήσει σε ταχύτερες και ασφαλέστερες μετακινήσεις.

Προγνωστική συντήρηση εξοπλισμού: η προγνωστική συντήρηση είναι ένας κλάδος της μηχανικής μάθησης που χρησιμοποιεί τα δεδομένα αισθητήρων σε κάποιον εξοπλισμό προκειμένου να προβλέψει μελλοντικές ζημιές ή αστοχίες του, επιτρέποντας τον έγκαιρο προγραμματισμό της συντήρησής του και την αποφυγή σφαλμάτων ή και διακοπών της λειτουργίας του.

Φυσικά οι εφαρμογές της μηχανικής μάθησης περιλαμβάνουν πολλά περισσότερα πεδία όπως την ιατρική, την οικονομία, τη βιοπληροφορική, τη χημειοπληροφορική κα.[9]

Κεφάλαιο 3: Ιδέα Μηχανικής Μάθησης

3.1 Συλλογή δεδομένων

Η συλλογή δεδομένων κατέχει εξέχοντα ρόλο στη μηχανική μάθηση καθώς ο σχεδιασμός μοντέλων με ακρίβεια και η μετέπειτα εξαγωγή χρήσιμων συμπερασμάτων απαιτούν μεγάλο όγκο δεδομένων.

Οι πηγές των δεδομένων αυτών ποικίλουν και μπορεί να περιλαμβάνουν δημόσιες βάσεις δεδομένων, πλατφόρμες μέσων κοινωνικής δικτύωσης, Διεπαφές Προγραμματισμού Εφαρμογών (Application Programming Interfaces, APIs) και άλλα σύνολα δεδομένων. Αποτελεί μία διαδικασία δύσκολη και συχνά χρονοβόρα καθώς είναι απαραίτητο να διασφαλιστεί ότι τα δεδομένα που συλλέγονται είναι σχετικά με το προς επίλυση πρόβλημα αλλά ταυτόχρονα και αντιπροσωπευτικά. Πρέπει λοιπόν να υπάρχει ποικιλία στα δεδομένα που συλλέγονται προκειμένου το μοντέλο να μπορεί να αντιμετωπίσει τα διάφορα σενάρια που μελλοντικά θα συναντήσει.

Ανάλογα με τη φύση του εκάστοτε προβλήματος και τους διαθέσιμους πόρους, υπάρχουν διάφορες μέθοδοι συλλογής δεδομένων. Η χειροκίνητη εισαγωγή δεδομένων είναι μία χρονοβόρα και επιρρεπής σε σφάλματα μέθοδος, πλην όμως χρήσιμη όταν τα δεδομένα δεν είναι διαθέσιμα σε ψηφιακή μορφή. Το Web scraping είναι μία μέθοδος κατά την οποία τα δεδομένα εξάγονται αυτόματα από ιστοτόπους, διαδικασία σχετικά γρήγορη και χρήσιμη για τη συλλογή μεγάλου όγκου δεδομένων. Στην περίπτωση που τα δεδομένα προκύπτουν από το φυσικό περιβάλλον, η συλλογή τους μπορεί να γίνει μέσω αισθητήρων, φωτογραφικών μηχανών ή άλλων συσκευών καταγραφής δεδομένων σε πραγματικό χρόνο.

Πρέπει να σημειωθεί ότι η ποιότητα των δεδομένων παίζει σημαντικό ρόλο στην απόδοση ενός μοντέλου μηχανικής μάθησης. Ως εκ τούτου, είναι σημαντικό να διασφαλίζεται ότι τα δεδομένα είναι ακριβή και αντιπροσωπευτικά. Τέλος, είναι απαραίτητο να λαμβάνονται υπόψη οι ηθικές επιπτώσεις της συλλογής και χρήσης των δεδομένων, διαδικασία που πρέπει πάντα να είναι σύμφωνη με τους σχετικούς νόμους και κανονισμούς.

3.2 Προεπεξεργασία δεδομένων

Μετά τη συλλογή, η προεπεξεργασία των δεδομένων είναι επίσης καθοριστικής σημασίας για τη μηχανική μάθηση και περιλαμβάνει τη μετατροπή των μη επεξεργασμένων δεδομένων σε μια μορφή που να μπορεί να αξιοποιηθεί για την εκπαίδευση του μοντέλου. Τα δεδομένα αρχικά είναι συνήθως μη δομημένα και περιέχουν ενδεχομένως και ελλιπείς τιμές, κάτι που μπορεί να οδηγήσει σε κακή απόδοση του μοντέλου. Ως εκ τούτου, η προεπεξεργασία των

δεδομένων είναι απαραίτητη για να διασφαλιστεί ότι τα δεδομένα είναι καθαρά, πλήρη και μορφοποιημένα κατάλληλα για την εκπαίδευση του μοντέλου.

Το πρώτο βήμα στην προεπεξεργασία δεδομένων είναι ο καθαρισμός των δεδομένων, ο οποίος περιλαμβάνει την αφαίρεση ή τη διόρθωση τυχόν σφαλμάτων ή ακραίων τιμών στα δεδομένα. Αυτό μπορεί να γίνει με διάφορες τεχνικές, όπως συμπληρώνοντας τις ελλειπείς τιμές με χρήση στατιστικών μεθόδων (Data Imputation) ή την αφαίρεση δεδομένων που περιέχουν σφάλματα ή ακραίες τιμές.

Το επόμενο βήμα στην προεπεξεργασία δεδομένων είναι ο μετασχηματισμός δεδομένων, ο οποίος περιλαμβάνει τη μετατροπή των δεδομένων σε μορφή κατάλληλη και αξιοποιήσιμη από τους αλγόριθμους μηχανικής μάθησης. Αυτό επιτυγχάνεται με διάφορες τεχνικές όπως η κανονικοποίηση των δεδομένων ή κάποιων μεμονωμένων χαρακτηριστικών τους για τη βελτίωση της απόδοσης του μοντέλου.

Η κωδικοποίηση δεδομένων είναι ένα άλλο σημαντικό βήμα στην προεπεξεργασία δεδομένων, η οποία περιλαμβάνει τη μετατροπή κατηγορικών μεταβλητών σε αριθμητική μορφή που μπορεί να αξιοποιηθεί από τους αλγόριθμους μηχανικής μάθησης. Μια τεχνική που επιτυγχάνει το παραπάνω δημιουργεί δυαδικές μεταβλητές για κάθε κατηγορία της κατηγορικής μεταβλητής.

Τέλος, η προεπεξεργασία δεδομένων μπορεί επίσης να περιλαμβάνει την επιλογή συγκεκριμένων χαρακτηριστικών για κάποιο συγκεκριμένο μοντέλο αντί για ολόκληρο το σύνολο των δεδομένων, μειώνοντας τον όγκο τους και βελτιώνοντας τη συνολική απόδοση του μοντέλου.

3.3 Εκπαίδευση μοντέλου

Μετά τη συλλογή και την προεπεξεργασία τους, τα δεδομένα χωρίζονται σε δύο σύνολα, τα δεδομένα εκπαίδευσης και τα δεδομένα δοκιμών. Αρχικά, τα δεδομένα εκπαίδευσης χρησιμοποιούνται προκειμένου το μοντέλο να μάθει από αυτά και να μπορεί να πραγματοποιεί ακριβείς προβλέψεις. Η εκπαίδευση του μοντέλου γίνεται με τη χρήση διαφόρων αλγορίθμων και παραμέτρων.

Στην επιβλεπόμενη μάθηση για παράδειγμα τα δεδομένα αποτελούνται από ζεύγη δεδομένων εισόδου και εξόδου. Η εκπαίδευση ενός τέτοιου μοντέλου περιλαμβάνει ένα σύνολο παραμέτρων που αντιπροσωπεύουν με ακρίβεια τη σχέση μεταξύ των δεδομένων εισόδου και των δεδομένων εξόδου. Αυτό επιτυγχάνεται με την προσαρμογή των παραμέτρων του μοντέλου επαναληπτικά μέχρι τελικά το μοντέλο να παράγει την επιθυμητή έξοδο για μια δεδομένη είσοδο.

Ανεξάρτητα από τον τύπο του μοντέλου, κατά την εκπαίδευσή του είναι σημαντικό η απόδοσή του να αξιολογείται με βάση το σύνολο των δεδομένων δοκιμών. Με τον τρόπο αυτό

διασφαλίζεται ότι το μοντέλο έχει εκπαιδευτεί κατάλληλα προκειμένου να πραγματοποιεί ακριβείς προβλέψεις.

3.4 Κλήση του μοντέλου σε άγνωστα δεδομένα

Τα μοντέλα επιβλεπόμενης μάθησης, τα οποία εκπαιδεύονται σε ζεύγη δεδομένων εισόδου και εξόδου, με την έξοδο να είναι γνωστή μπορούν όταν κληθούν να κάνουν προβλέψεις για νέα, άγνωστα δεδομένα με βάση αυτά που έχουν μάθει από τα δεδομένα εκπαίδευσης. Από την άλλη πλευρά, τα μοντέλα χωρίς επίβλεψη δεν διαθέτουν τέτοια ζεύγη δεδομένων. Αντ' αυτού, κατά την κλήση τους χρησιμοποιούνται για τον εντοπισμό μοτίβων και συγκεκριμένης δομής στα δεδομένα εισόδου κάνοντας για παράδειγμα κατηγοριοποιήσεις. Ένας ιδιαίτερος τύπος μοντέλων, τα ημιαποπτευόμενα μοντέλα, τα οποία είναι ένας συνδυασμός επιβλεπόμενης και μη επιβλεπόμενης μάθησης εκπαιδεύονται σε μια μικρή ποσότητα δεδομένων με γνωστή έξοδο και μια μεγάλη ποσότητα δεδομένων χωρίς γνωστή έξοδο. Μπορούν να χρησιμοποιηθούν για να κάνουν προβλέψεις για νέα, άγνωστα δεδομένα αλλά γενικά απαιτούν περισσότερα δεδομένα για να εκπαιδευτούν σε σχέση με τα μοντέλα επιβλεπόμενης μάθησης. Αξίζει να σημειωθεί ότι ανεξάρτητα από τον τύπο του μοντέλου, είναι πάντα σημαντικό να αξιολογείται η απόδοση του μοντέλου σε ξεχωριστό σύνολο δεδομένων δοκιμών, ώστε να διασφαλίζεται η σωστή λειτουργία του και η αποφυγή φαινομένων όπως η Υπερπροσαρμογή και η Υποπροσαρμογή.

3.5 Υπερπροσαρμογή και Υποπροσαρμογή

Στη μηχανική μάθηση οι όροι Υπερπροσαρμογή (Overfitting) και Υποπροσαρμογή (Underfitting) αναφέρονται σε δύο ανεπιθύμητα φαινόμενα που σχετίζονται με το μοντέλο μηχανικής μάθησης και την απόδοσή του στα δεδομένα εκπαίδευσης και σε άγνωστα δεδομένα.

Υπερπροσαρμογή συμβαίνει όταν ένα μοντέλο εκπαιδεύεται και προσαρμόζεται πολύ καλά στα δεδομένα εκπαίδευσης αλλά η απόδοσή του όταν κληθεί σε άγνωστα δεδομένα είναι κακή. Το παραπάνω συμβαίνει διότι το μοντέλο εντοπίζει και συμπεριλαμβάνει περισσότερο τον «θόρυβο» που περιέχουν τα δεδομένα εκπαίδευσης παρά το μοτίβο που υποβόσκει σε αυτά και την πραγματική σχέση που τα συνδέει. Αυτό μπορεί να οφείλεται είτε στην υψηλή πολυπλοκότητα (και κατ' επέκταση και υψηλή προσαρμοστικότητα) του μοντέλου, είτε στο μικρό μέγεθος των δεδομένων εκπαίδευσης. Τεχνικές με τις οποίες αντιμετωπίζουμε το φαινόμενο της Υπερπροσαρμογής είναι η κανονικοποίηση, η τεχνική early stopping και η μέθοδος cross-validation.

Υποπροσαρμογή συμβαίνει όταν ένα μοντέλο δεν είναι αρκετά πολύπλοκο ώστε να αποτυπώνει το μοτίβο ή τη σχέση που συνδέει τα δεδομένα και έχει κακή απόδοση τόσο με τα δεδομένα εκπαίδευσης, όσο και με τα δεδομένα δοκιμών. Το παραπάνω μπορεί να συμβεί λόγω της υπερβολικής απλότητας του μοντέλου ή λόγω του μικρού μεγέθους των δεδομένων εκπαίδευσης. Για την αποφυγή του φαινομένου αυτού χρησιμοποιούνται τεχνικές όπως η αύξηση της πολυπλοκότητας του μοντέλου, η προσθήκη νέων δεδομένων εκπαίδευσης και η τεχνική feature engineering.

Είναι σημαντικό να αναφερθεί ότι η εύρεση της «χρυσής τομής» μεταξύ της Υπερπροσαρμογής και της Υποπροσαρμογής ώστε να αποφευχθούν αμφότερα τα φαινόμενα απαιτεί καλή κατανόηση του προβλήματος, του μοντέλου και των δεδομένων, χρήση διαφόρων αρχιτεκτονικών μοντέλων και παραμέτρων αλλά και έλεγχο της απόδοσης του μοντέλου με κάποιο ξεχωριστό σύνολο δεδομένων δοκιμών.[11]

Κεφάλαιο 4: Κατηγορίες αλγορίθμων Μηχανικής Μάθησης

4.1 Παλινδρόμηση και Κατηγοριοποίηση

Η παλινδρόμηση και η κατηγοριοποίηση είναι δύο κύριοι τύποι αλγορίθμων επιβλεπόμενης μάθησης και χρησιμοποιούνται για την πραγματοποίηση προβλέψεων βάσει των δεδομένων εισόδου.

Οι αλγόριθμοι παλινδρόμησης χρησιμοποιούνται για την πρόβλεψη συνεχών τιμών, όπως για παράδειγμα η αξία ενός ακινήτου ή η θερμοκρασία σε μια πόλη κάποια συγκεκριμένη χρονική στιγμή. Βασίζονται στην υπόθεση ότι υπάρχει σχέση που συνδέει τις ανεξάρτητες μεταβλητές (δεδομένα εισόδου) και τις εξαρτημένες μεταβλητές (δεδομένα εξόδου) και ο στόχος είναι να εντοπισθεί η σχέση αυτή. Ορισμένοι βασικοί αλγόριθμοι παλινδρόμησης που χρησιμοποιούνται ευρέως είναι η γραμμική παλινδρόμηση (linear regression), η πολυωνυμική παλινδρόμηση (polynomial regression) και η πολλαπλή παλινδρόμηση (multiple regression).

Από την άλλη πλευρά, οι αλγόριθμοι κατηγοριοποίησης χρησιμοποιούνται για την πρόβλεψη τιμών κατηγορικών μεταβλητών, όπως για παράδειγμα το εάν ένα μήνυμα ηλεκτρονικού ταχυδρομείου είναι ανεπιθύμητο ή όχι ή εάν μια εικόνα περιέχει γάτα ή σκύλο. Βασίζονται στην υπόθεση ότι τα δεδομένα μπορούν να χωριστούν σε διαφορετικές τάξεις ή κατηγορίες και ο στόχος είναι να βρεθεί το όριο που διαχωρίζει αυτές τις τάξεις (decision boundary). Μερικοί συνηθισμένοι αλγόριθμοι κατηγοριοποίησης είναι η λογιστική παλινδρόμηση (logistic regression), ο αλγόριθμος των k-πλησιέστερων γειτόνων (k-nearest neighbors, k-NN), τα δένδρα αποφάσεων, ο αλγόριθμος Random Forest και οι Μηχανές διανυσμάτων υποστήριξης (SVMs).

Η κύρια διαφορά μεταξύ των αλγορίθμων παλινδρόμησης και κατηγοριοποίησης είναι ο τύπος της εξόδου που παράγουν. Οι αλγόριθμοι παλινδρόμησης παράγουν εξόδους με συνεχείς τιμές, ενώ οι αλγόριθμοι κατηγοριοποίησης παράγουν εξόδους με διακριτές τιμές. Επιπλέον, διαφέρουν μεταξύ τους και οι μέθοδοι που χρησιμοποιούνται για τη μέτρηση της απόδοσης των δύο τύπων αλγορίθμων. Για την παλινδρόμηση μελετώνται δείκτες όπως το μέσο τετραγωνικό σφάλμα (MSE) και το R-τετράγωνο (R^2), ενώ για την κατηγοριοποίηση μελετώνται η ακρίβεια (accuracy και precision) και οι δείκτες recall και F1-score. Τέλος, όσον αφορά την πολυπλοκότητά τους, οι περισσότεροι αλγόριθμοι παλινδρόμησης είναι απλούστεροι από τους αλγορίθμους κατηγοριοποίησης καθώς οι αλγόριθμοι παλινδρόμησης βασίζονται σε γραμμικές ή πολυωνυμικές συναρτήσεις, οι οποίες μπορούν να λυθούν αναλυτικά.[6]

4.2 Επιβλεπόμενη Μάθηση, Μη Επιβλεπόμενη Μάθηση και Ενισχυτική Μάθηση

Όπως έχει ήδη αναφερθεί, οι κατηγορίες τεχνικών μηχανικής μάθησης περιλαμβάνουν τρεις βασικούς τύπους: την επιβλεπόμενη μάθηση, τη μη επιβλεπόμενη μάθηση και την ενισχυτική μάθηση. Αυτές οι τεχνικές διαφέρουν στον τρόπο με τον οποίο χρησιμοποιούν τα δεδομένα για να εκπαιδευτούν και να κάνουν προβλέψεις.

Στην επιβλεπόμενη μάθηση το μοντέλο εκπαιδεύεται χρησιμοποιώντας επισημασμένα δεδομένα τα οποία αποτελούνται από μεταβλητές εισόδου και αντίστοιχες μεταβλητές εξόδου. Ο στόχος της επιβλεπόμενης μάθησης είναι ο εντοπισμός της σχέσης μεταξύ των μεταβλητών εισόδου και των μεταβλητών εξόδου, η οποία επιτρέπει στο μοντέλο να κάνει προβλέψεις για νέα, άγνωστα δεδομένα. Στη μη επιβλεπόμενη μάθηση το μοντέλο εκπαιδεύεται χρησιμοποιώντας μη επισημασμένα δεδομένα τα οποία αποτελούνται από μεταβλητές εισόδου χωρίς αντίστοιχες μεταβλητές εξόδου. Ο στόχος της μη επιβλεπόμενης μάθησης είναι ο εντοπισμός μοτίβων και σχέσεων που συνδέουν τα δεδομένα. Η ενισχυτική μάθηση είναι ένας τύπος μηχανικής μάθησης στον οποίο το μοντέλο μαθαίνει κάνοντας δοκιμές και υπολογίζοντας κάποιο σφάλμα. Το μοντέλο εκπαιδεύεται χρησιμοποιώντας ένα σύστημα ανταμοιβής, και, με βάση τις ενέργειές του, λαμβάνει είτε θετική είτε αρνητική ανατροφοδότηση. Ο στόχος της ενισχυτικής μάθησης είναι να μάθει τη βέλτιστη τακτική ή ακολουθία ενεργειών η οποία μεγιστοποιεί τη συνολική ανταμοιβή.[38,39]

Η επιλογή μεταξύ επιβλεπόμενης, μη επιβλεπόμενης και ενισχυτικής μάθησης εξαρτάται από το είδος του εκάστοτε προβλήματος. Η επιβλεπόμενη μάθηση λοιπόν χρησιμοποιείται κυρίως όταν υπάρχουν διαθέσιμα δεδομένα επισημασμένα και ο στόχος είναι να προβλεφθεί η τιμή μιας συγκεκριμένης μεταβλητής εξόδου. Η μη επιβλεπόμενη μάθηση χρησιμοποιείται όταν δεν υπάρχουν διαθέσιμα τέτοια επισημασμένα δεδομένα και ο στόχος είναι να βρεθούν μοτίβα και σχέσεις που συνδέουν τα δεδομένα. Η ενισχυτική μάθηση χρησιμοποιείται όταν το μοντέλο πρέπει να μάθει μέσω δοκιμών και μέτρησης σφαλμάτων με στόχο τη μεγιστοποίηση κάποιας συνάρτησης ανταμοιβής. Είναι σαφές ότι η βαθιά κατανόηση των διαφορών μεταξύ αυτών των τριών ειδών μηχανικής μάθησης είναι απαραίτητη για την επιλογή της κατάλληλης τεχνικής για ένα δεδομένο πρόβλημα.

Κεφάλαιο 5: Επιβλεπόμενη Μάθηση

5.1 Ορισμός και Παραδείγματα

Η επιβλεπόμενη μάθηση αποτελεί ακρογωνιαίο λίθο στο πεδίο της μηχανικής μάθησης ως τεχνική που επιστρατεύεται στις περιπτώσεις που το μοντέλο εκπαιδεύεται από επισημασμένα δεδομένα, επιτρέποντάς του έτσι να κάνει ακριβείς προβλέψεις για άλλα, άγνωστα δεδομένα. Όπως έχει ήδη αναφερθεί, τα δεδομένα που χρησιμοποιούνται για την εκπαίδευση περιλαμβάνουν μεταβλητές εισόδου (ή χαρακτηριστικά) και αντίστοιχες μεταβλητές εξόδου ή ετικέτες (labels). Ο πρωταρχικός στόχος είναι να δημιουργηθεί ένα προγνωστικό μοντέλο που μπορεί να συμπεράνει τη σχέση μεταξύ μεταβλητών εισόδου και μεταβλητών εξόδου. Η αξιοποίηση της γνώσης αυτής της σχέσης δίνει τη δυνατότητα στο μοντέλο να πραγματοποιεί ακριβείς προβλέψεις και για νέες, μη επισημασμένες περιπτώσεις δεδομένων.[35]

Η σημασία της επιβλεπόμενης μάθησης γίνεται καλύτερα αντιληπτή από το ευρύ φάσμα των εφαρμογών με τις οποίες σχετίζεται. Στον τομέα των χρηματοοικονομικών, για παράδειγμα, χρησιμοποιούνται τεχνικές επιβλεπόμενης μάθησης για την πρόβλεψη των τιμών των μετοχών ή τον εντοπισμό παράνομων συναλλαγών. Στον τομέα της ιατρικής, αυτή η τεχνική επιτρέπει την ακριβή διάγνωση διαφόρων νόσων με ανάλυση των δεδομένων των ασθενών σε ένα πλαίσιο επισημασμένων δεδομένων παλαιότερων ιατρικών αρχείων. Οι εφαρμογές επεξεργασίας φυσικής γλώσσας αξιοποιούν επίσης την επιβλεπόμενη μάθηση για να επιτύχουν ανάλυση συναισθήματος (Sentiment Analysis), μετάφραση γλώσσας και αναγνώριση ομιλίας.

Η επιβλεπόμενη μάθηση περιλαμβάνει δύο κύριες υποκατηγορίες, την παλινδρόμηση και την κατηγοριοποίηση, τα χαρακτηριστικά των οποίων έχουν ήδη αναφερθεί, ενώ αυτές με τη σειρά τους περιλαμβάνουν επιμέρους κατηγορίες αλγορίθμων που αναλύονται παρακάτω και καθένας εξ' αυτών επιλέγεται κάθε φορά με βάση την καταλληλότητά του για το εκάστοτε πρόβλημα. Ορισμένοι διαδεδομένοι αλγόριθμοι που θα μελετηθούν είναι η γραμμική παλινδρόμηση, η πολλαπλή γραμμική παλινδρόμηση, η πολυωνυμική παλινδρόμηση, η μέθοδος των k-πλησιέστερων γειτόνων (k-Nearest Neighbors), η παλινδρόμηση με διανύσματα υποστήριξης, η παλινδρόμηση με δένδρα αποφάσεων, η παλινδρόμηση με μέθοδο τυχαίου δάσους και οι μέθοδοι XGBoost και ADABOost.[10,16]

5.2 Κατηγορίες αλγορίθμων Επιβλεπόμενης Μάθησης

5.2.1 Γραμμική παλινδρόμηση

Η απλή γραμμική παλινδρόμηση αποτελεί έναν από τους θεμελιώδεις αλγορίθμους στον τομέα της επιβλεπόμενης μάθησης. Αφορά τον καθορισμό μιας γραμμικής σχέσης μεταξύ μιας μεταβλητής εισόδου (ανεξάρτητη μεταβλητή) και μιας συνεχούς μεταβλητής εξόδου (εξαρτημένη μεταβλητή). Ο βασικός στόχος αυτού του αλγορίθμου είναι η εύρεση της βέλτιστης γραμμής, εκείνης δηλαδή που «ταιριάζει» καλύτερα στα δεδομένα. Αυτό επιτυγχάνεται από τη γραμμή που ελαχιστοποιεί το άθροισμα των τετραγωνικών διαφορών μεταξύ των παρατηρούμενων τιμών και των προβλεφθεισών τιμών, δηλαδή το άθροισμα των τετραγώνων της απόστασης κάθε σημείου-δεδομένου από το αντίστοιχο σημείο της γραμμής αυτής. Μαθηματικά το μοντέλο μπορεί να αναπαρασταθεί ως:

$$y = b_0 + b_1x + \varepsilon$$

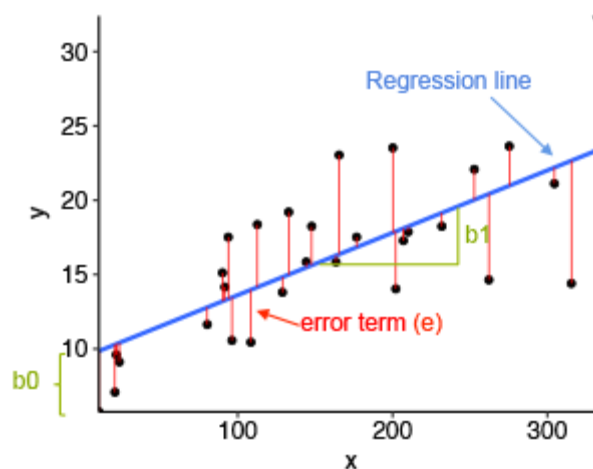
όπου y : η έξοδος (ή εξαρτημένη μεταβλητή),

x : η είσοδος (ή ανεξάρτητη μεταβλητή),

b_1 : η κλίση της γραμμής,

b_0 : ο σταθερός όρος και

ε : το σφάλμα



Σχήμα 1: Γραμμική παλινδρόμηση

Το παραπάνω επιτυγχάνεται από τον αλγόριθμο με τον προσδιορισμό των βέλτιστων τιμών για τα b_0 και b_1 έτσι ώστε να ελαχιστοποιείται η απόσταση αυτή. Η απλή γραμμική παλινδρόμηση χρησιμοποιείται σε διάφορες εφαρμογές. Στον τομέα της οικονομίας, μπορεί να χρησιμοποιηθεί για την πρόβλεψη της σχέσης μεταξύ μεταβλητών όπως το εισόδημα και οι δαπάνες. Στις περιβαλλοντικές επιστήμες, μπορεί να μοντελοποιήσει συσχετισμούς μεταξύ της θερμοκρασίας και της ανάπτυξης των φυτών. Ο αλγόριθμος αυτός βοηθάει στην κατανόηση της έννοιας της παλινδρόμησης και θέτει τα θεμέλια για την κατανόηση πιο σύνθετων

τεχνικών, όπως η πολλαπλή γραμμική παλινδρόμηση και η πολυωνυμική παλινδρόμηση που αναφέρονται παρακάτω.

5.2.2 Πολλαπλή γραμμική παλινδρόμηση

Η πολλαπλή γραμμική παλινδρόμηση επεκτείνει τις αρχές της απλής γραμμικής παλινδρόμησης και εφαρμόζεται στις περιπτώσεις με περισσότερες μεταβλητές εισόδου. Ο αλγόριθμος αυτός λοιπόν επιτρέπει την εύρεση συσχετίσεων μεταξύ πολλαπλών μεταβλητών εισόδου και μιας συνεχούς μεταβλητής εξόδου. Ο στόχος του είναι και πάλι να προσδιορίσει τους βέλτιστους συντελεστές που ελαχιστοποιούν την απόκλιση μεταξύ παρατηρούμενων και προβλεφθεισών τιμών. Μαθηματικά, το μοντέλο μπορεί να αναπαρασταθεί ως:

$$y = b_0 + b_1x_1 + b_2x_2 + \dots + b_nx_n + \varepsilon$$

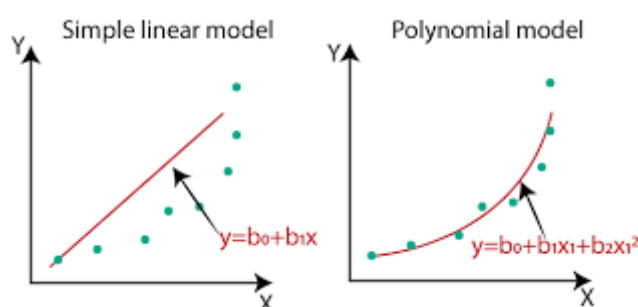
όπου η ερμηνεία κάθε όρου προκύπτει από την επεξήγηση των όρων της απλής γραμμικής παλινδρόμησης με τη διαφορά ότι εδώ έχουμε n ανεξάρτητες μεταβλητές και ισάριθμους συντελεστές για τις κλίσεις των αντίστοιχων ευθειών. Ο αλγόριθμος περιλαμβάνει την επαναλαμβανόμενη προσαρμογή των συντελεστών $b_0, b_1, \dots, b_n$ ως την εύρεση των βέλτιστων, όμοια δηλαδή με την απλή γραμμική παλινδρόμηση.[13] Οι εφαρμογές της πολλαπλής γραμμικής παλινδρόμησης είναι επίσης εκτεταμένες. Για παράδειγμα, στον τομέα της διοίκησης επιχειρήσεων, μπορεί να χρησιμοποιηθεί για την ανάλυση του αντίκτυπου διαφόρων στρατηγικών μάρκετινγκ στις πωλήσεις. Στις περιβαλλοντικές επιστήμες, μπορεί να βοηθήσει στη διάκριση των επιπτώσεων διαφόρων παραγόντων στις τάσεις αύξησης ή μείωσης της θερμοκρασίας. Γενικά, η ευελιξία του αλγορίθμου πηγάζει από την ικανότητά του να χειρίζεται πολυδιάστατα δεδομένα και να παρέχει πληροφορίες για πολύπλοκες σχέσεις μεταξύ αυτών.

5.2.3 Πολυωνυμική παλινδρόμηση

Η πολυωνυμική παλινδρόμηση αποτελεί μία ευέλικτη επέκταση της γραμμικής παλινδρόμησης. Εφαρμόζεται σε περιπτώσεις μη γραμμικών σχέσεων μεταξύ των μεταβλητών εισόδου και της μεταβλητής εξόδου, προσφέροντας μια αποτελεσματική προσέγγιση για τη μοντελοποίηση πιο σύνθετων συσχετισμών μεταξύ των δεδομένων. Η πολυωνυμική παλινδρόμηση περιλαμβάνει την προσαρμογή των συντελεστών μιας πολυωνυμικής εξίσωσης παρέχοντας μια καμπύλη γραμμή που μπορεί να συλλάβει κάποιες περίπλοκες σχέσεις

καλύτερα απ' ό,τι μια ευθεία γραμμή, όπως φαίνεται και στο Σχήμα 2. Το μαθηματικό μοντέλο αναπαρίσταται ως:

$$y = b_0 + b_1x + b_2x^2 + \dots + b_dx^d + \varepsilon$$



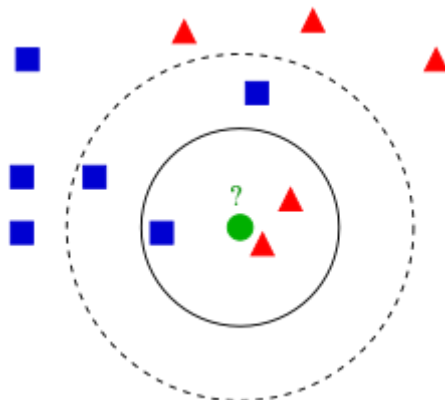
Σχήμα 2: Γραμμική παλινδρόμηση (αριστερά), Πολυωνυμική παλινδρόμηση (δεξιά)

Ο βαθμός του πολυωνύμου καθορίζει την καμπυλότητα της γραμμής (ένας υψηλότερος βαθμός συνεπάγεται πιο σύνθετες καμπύλες, ενδεχομένως «ταιριάζοντας» στα δεδομένα με μεγαλύτερη ακρίβεια).[14] Η πολυωνυμική παλινδρόμηση εφαρμόζεται σε πληθώρα πεδίων όπως η φυσική, καθώς μπορεί να μοντελοποιήσει την τροχιά ενός βλήματος σε ποικίλα βαρυτικά πεδία ή στα οικονομικά, καθώς μπορεί να συσχετίσει περίπλοκες τάσεις στη συμπεριφορά των καταναλωτικών δαπανών με την πάροδο του χρόνου. Ενώ η πολυωνυμική παλινδρόμηση προσφέρει μεγάλη ευελιξία, η επιλογή του κατάλληλου πολυωνυμικού βαθμού είναι ζωτικής σημασίας για την απόδοση του αλγορίθμου. Ένας υπερβολικά υψηλός βαθμός του πολυωνύμου μπορεί να οδηγήσει στο φαινόμενο της υπερπροσαρμογής και αντίστοιχα ένας υπερβολικά χαμηλός βαθμός μπορεί να οδηγήσει στο φαινόμενο της υποπροσαρμογής.

5.2.4 Μέθοδος των k-πλησιέστερων γειτόνων

Η μέθοδος των k-πλησιέστερων γειτόνων (k-Nearest Neighbors, k-NN) είναι ένας αλγόριθμος μη παραμετρικής, επιβλεπόμενης μηχανικής μάθησης που χρησιμοποιείται τόσο για κατηγοριοποίηση όσο και για παλινδρόμηση. Ο αλγόριθμος βασίζεται στην ιδέα ότι συχνά, σημεία δεδομένων με παρόμοια χαρακτηριστικά, που βρίσκονται δηλαδή σε κοντινές αποστάσεις στο χώρο των χαρακτηριστικών (feature space), έχουν παραπλήσιες τιμές εξόδου ή ανήκουν στην ίδια τάξη. Η λειτουργία του αλγορίθμου περιλαμβάνει την εύρεση ενός αριθμού k πλησιέστερων σημείων του συνόλου δεδομένων σε ένα νέο σημείο (στο χώρο χαρακτηριστικών) και, βάσει αυτών, να προβλεφθεί η τάξη ή η τιμή του σημείου αυτού. Στην περίπτωση της κατηγοριοποίησης, ο αλγόριθμος εκχωρεί στο νέο σημείο την επικρατούσα

ετικέτα μεταξύ των k γειτόνων βάσει πλειοψηφίας (ένα τέτοιο παράδειγμα φαίνεται στο Σχήμα 3), ενώ στην περίπτωση της παλινδρόμησης, ο αλγόριθμος υπολογίζει τη μέση τιμή εξόδου για τους k γείτονες και την εκχωρεί στο νέο σημείο.



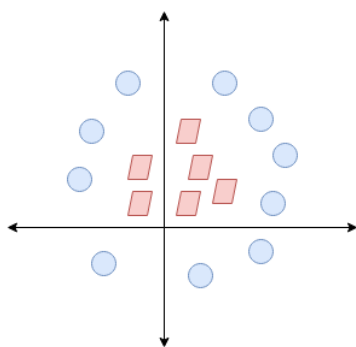
Σχήμα 3: K-Πλησιέστεροι Γείτονες

Η επιλογή μιας κατάλληλης τιμής k είναι πρωταρχικής σημασίας για την απόδοση του αλγορίθμου. Πολύ μικρές τιμές ενδέχεται να εισάγουν ευαισθησία στον τοπικό θόρυβο, οδηγώντας σε ανακριβείς προβλέψεις. Αντίθετα, πολύ μεγάλες τιμές ενδεχομένως να γενικεύσουν υπερβολικά τις προβλέψεις, αποκρύπτοντας μοτίβα που εμφανίζονται τοπικά.[15] Η επιλογή της μεθόδου μέτρησης της απόστασης μεταξύ των σημείων (πχ Ευκλείδεια απόσταση, απόσταση Manhattan ή άλλες) επηρεάζει επίσης τη συμπεριφορά του αλγορίθμου, αφού αυτή καθορίζει και ποιοι θα είναι οι γείτονες για το εκάστοτε σημείο ενδιαφέροντος. Ο αλγόριθμος k -NN βρίσκει εφαρμογή σε ένα πλήθος πεδίων. Για παράδειγμα, σε ό,τι αφορά την ιατρική διάγνωση, μπορεί να βοηθήσει στον προσδιορισμό της πιθανότητας ένας ασθενής να πάσχει από μια συγκεκριμένη πάθηση με βάση παρόμοια ιατρικά αρχεία άλλων ασθενών ενώ στα συστήματα συστάσεων μπορεί να βοηθήσει στην πρόταση προϊόντων ή υπηρεσιών σε χρήστες με προτιμήσεις παρόμοιες με εκείνες άλλων χρηστών. Γενικά, ενώ ο τρόπος λειτουργίας του αλγορίθμου k -NN είναι εύκολο να κατανοηθεί και μπορεί να εφαρμοστεί σχετικά γρήγορα, ο αλγόριθμος δεν είναι τόσο αποδοτικός με δεδομένα σε χώρους πολλών διαστάσεων. Παρ' όλα αυτά, αποτελεί ένα σημαντικό εργαλείο στις τεχνικές προγνωστικής μοντελοποίησης, προσφέροντας χρήσιμες προσεγγίσεις για την αντιμετώπιση προβλημάτων που χαρακτηρίζονται από τοπικές σχέσεις μεταξύ των μεταβλητών εισόδου και εξόδου.

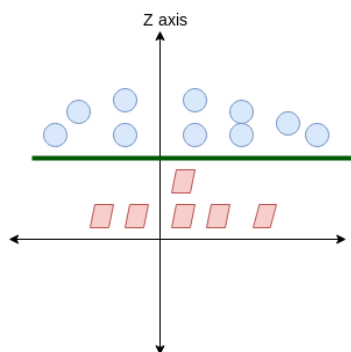
5.2.5 Παλινδρόμηση με διανύσματα υποστήριξης

Η μέθοδος της παλινδρόμησης με διανύσματα υποστήριξης (Support Vector Regression, SVR) είναι ένας αλγόριθμος παλινδρόμησης που προκύπτει από το μοντέλο των μηχανών

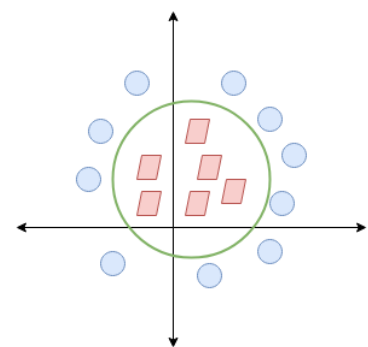
διανυσμάτων υποστήριξης (SVMs). Ο αλγόριθμος αυτός είναι σχεδιασμένος για να χειρίζεται μη γραμμικές σχέσεις και ο βασικός του στόχος είναι να βρει το υπερεπίπεδο (hyperplane) το οποίο «ταιριάζει» βέλτιστα στα δεδομένα, περιλαμβάνει δηλαδή όσο το δυνατόν περισσότερα σημεία δεδομένων εντός μιας ορισμένης απόστασης αλλά και περιορίζοντας παράλληλα το σφάλμα εντός καθορισμένων ορίων γύρω από τη γραμμή παλινδρόμησης (Σχήμα 7). Ο τρόπος που αυτό επιτυγχάνεται είναι με την αντιστοίχιση των σημείων δεδομένων (Σχήμα 4) σε έναν χώρο υψηλότερων διαστάσεων (Σχήμα 5) με τη βοήθεια ενός μετασχηματισμού, συνήθως μέσω των συναρτήσεων kernel. Ένας τέτοιος μετασχηματισμός επιτρέπει στον αλγόριθμο να εντοπίζει περίπλοκα, μη γραμμικά μοτίβα (Σχήμα 6). Συνήθεις συναρτήσεις kernel είναι η γραμμική, η πολυωνυμική και η συνάρτηση ακτινικής βάσης (Radial Basis Function, RBF), η καθεμία κατάλληλη για διαφορετικές δομές δεδομένων. Ο αλγόριθμος SVR επιδιώκει λοιπόν να διατηρήσει μια ισορροπία μεγιστοποιώντας το περιθώριο μεταξύ της γραμμής παλινδρόμησης και των διανυσμάτων υποστήριξης αλλά και ελαχιστοποιώντας ταυτόχρονα το σφάλμα εντός αυτού του περιθωρίου. Η επίτευξη αυτής της ισορροπίας απαιτεί την προσαρμογή των συντελεστών που καθορίζουν το υπερεπίπεδο και την κατάλληλη επιλογή της συνάρτησης kernel και των σχετικών παραμέτρων οι οποίες θα καθορίσουν και την απόδοση του αλγορίθμου. Ο αλγόριθμος SVR εφαρμόζεται σε ποικίλους τομείς όπως, για παράδειγμα στα χρηματοοικονομικά για την πρόβλεψη των τιμών των μετοχών, στη μηχανική για την πρόβλεψη της απόδοσης των μηχανών και στη βιοπληροφορική για την πρόβλεψη των δομών των πρωτεϊνών. Η παλινδρόμηση SVM χρησιμοποιείται επίσης στην επεξεργασία φυσικής γλώσσας και στην ανάλυση συναισθήματος. Επιπλέον, χρησιμοποιείται στην επεξεργασία εικόνας για την αναγνώριση αντικειμένων και στην ιατρική. Γενικά, ο SVR είναι ένας ευέλικτος αλγόριθμος που μπορεί να χρησιμοποιηθεί σε πολλούς τομείς για την πραγματοποίηση ακριβών προβλέψεων.



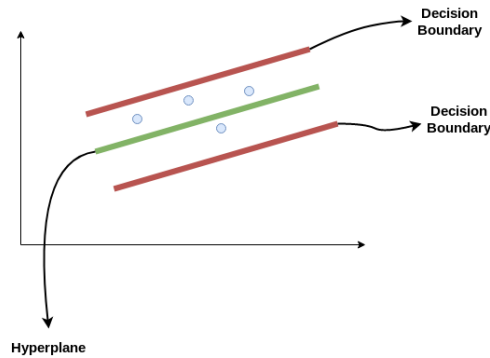
Σχήμα 4: Σημεία δεδομένων στον χώρο χαμηλής διάστασης



Σχήμα 5: Σημεία δεδομένων σε χώρο υψηλότερης διάστασης



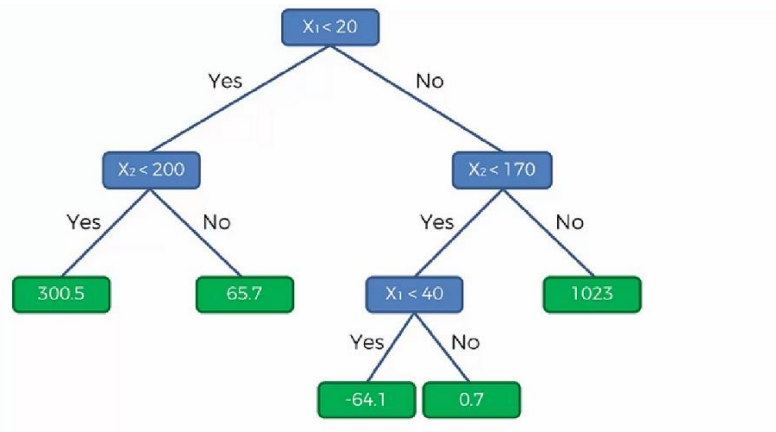
Σχήμα 6: Μη γραμμικό μοτίβο στον χώρο χαμηλής διάστασης



Σχήμα 7: Παλινδρόμηση με διανύσματα υποστήριξης

5.2.6 Παλινδρόμηση με δένδρα αποφάσεων

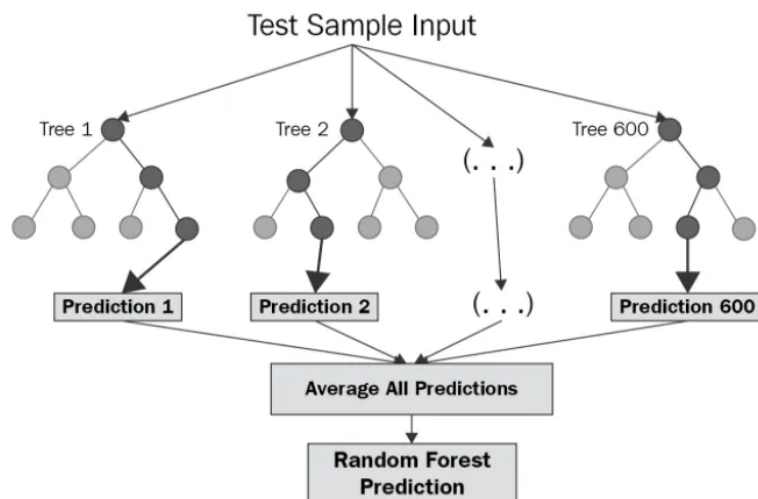
Η παλινδρόμηση με δένδρα αποφάσεων (Decision Tree Regression) είναι ένας αλγόριθμος παλινδρόμησης που βασίζεται στο μοντέλο των δένδρων αποφάσεων, μία μέθοδο επιβλεπόμενης μάθησης που γενικά χρησιμοποιείται τόσο για προβλήματα κατηγοριοποίησης όσο και για προβλήματα παλινδρόμησης. Η παλινδρόμηση με δένδρα αποφάσεων δημιουργεί ένα μοντέλο αποφάσεων καθώς και των αποτελεσμάτων τους το οποίο συμπεριλαμβάνει και την πρόβλεψη της τιμής μιας συνεχούς μεταβλητής εξόδου. Ένα δένδρο αποφάσεων κατασκευάζεται χρησιμοποιώντας μια μέθοδο επαναλαμβανόμενης διαίρεσης κατά την οποία τα δεδομένα χωρίζονται συνεχώς σε υποσύνολα βάσει των χαρακτηριστικών τους (μεταβλητές εισόδου). Σε κάθε εσωτερικό κόμβο του δένδρου δηλαδή πραγματοποιείται μια δοκιμή-ερώτηση για ένα από τα χαρακτηριστικά του αντικειμένου και τα δεδομένα χωρίζονται σε υποσύνολα με βάση το αποτέλεσμα της ερώτησης αυτής. Η διαδικασία επαναλαμβάνεται σε κάθε υποσύνολο μέχρις ότου να πληρούται ένα κριτήριο διακοπής. Το τελικό αποτέλεσμα είναι ένα δένδρο του οποίου κάθε φύλλο αντιστοιχεί σε μια πρόβλεψη για τη μεταβλητή εξόδου (Σχήμα 8). Η παλινδρόμηση με δένδρα αποφάσεων είναι γενικά εύκολο να κατανοηθεί ωστόσο καθοριστική για την αποδοτικότητα του αλγορίθμου είναι η επιλογή των κατάλληλων κάθε φορά παραμέτρων όπως το μέγιστο βάθος του δένδρου και ο ελάχιστος αριθμός δειγμάτων που απαιτούνται για τη διαίρεση ενός εσωτερικού κόμβου.[17] Είναι επίσης ως μέθοδος ευαίσθητη στην παρουσία ακραίων τιμών στα δεδομένα. Για την αντιμετώπιση αυτών των ζητημάτων μπορούν να εφαρμοστούν μέθοδοι όπως η Random Forest και η Gradient Boosting, οι οποίες μπορούν να παρέχουν πιο ακριβή μοντέλα συνδυάζοντας τις τιμές των προβλέψεων πολλαπλών δένδρων αποφάσεων. Οι εφαρμογές της μεθόδου καλύπτουν μεγάλο πλήθος πεδίων. Ενδεικτικά, στα χρηματοοικονομικά, χρησιμοποιείται στις προβλέψεις πιστωτικής ικανότητας λαμβάνοντας υπόψη πολλαπλά οικονομικά χαρακτηριστικά, ενώ στην οικολογία, χρησιμοποιείται σε προβλέψεις για την αύξηση ή τη μείωση του πληθυσμού διαφόρων ειδών.



Σχήμα 8: Παλινδρόμηση με δένδρα αποφάσεων

5.2.7 Παλινδρόμηση με τον αλγόριθμο τυχαίου δάσους

Η παλινδρόμηση με τον αλγόριθμο τυχαίου δάσους (Random Forest Regression) δεν είναι παρά μία επέκταση της μεθόδου παλινδρόμησης με δένδρα αποφάσεων. Αποτελεί έναν αλγόριθμο επιβλεπόμενης μάθησης κατά τον οποίο πολλαπλά δένδρα αποφάσεων συνδυάζονται με σκοπό τη βελτίωση της ακρίβειας της τελικής πρόβλεψης, την αντιμετώπιση του φαινομένου της υπερπροσαρμογής και τη διαχείριση σύνθετων σχέσεων μεταξύ των μεταβλητών εισόδου και των (αριθμητικών) μεταβλητών εξόδου. Επομένως, η βασική ιδέα του αλγορίθμου παλινδρόμησης με τη μέθοδο Random Forest περιλαμβάνει την κατασκευή ενός πλήθους δένδρων αποφάσεων, το καθένα εκ των οποίων εκπαιδεύεται σε ένα τυχαίο υποσύνολο των δεδομένων εκπαίδευσης. Αυτά τα δένδρα λειτουργούν ανεξάρτητα και το κάθε ένα πραγματοποιεί τη δική του ξεχωριστή πρόβλεψη. Η τελική πρόβλεψη προκύπτει από τον υπολογισμό της μέσης τιμής των προβλέψεων όλων των διαφορετικών δένδρων (Σχήμα 9), κάτι που διασφαλίζει τη μεγαλύτερη ακρίβειά της σε σχέση με μία μεμονωμένη πρόβλεψη ενός μόνο δένδρου.[18] Σημαντικό προτέρημα του αλγορίθμου είναι, όπως αναφέρθηκε, η ικανότητά του να αποφεύγει φαινόμενα υπερπροσαρμογής, χαρακτηριστικό εγγενές ενός δένδρου αποφάσεων, εισάγοντας ποικιλομορφία μεταξύ των μεμονωμένων δένδρων. Το παραπάνω επιτυγχάνεται μέσω τυχαίας δειγματοληψίας στα δεδομένα και τυχαίας επιλογής χαρακτηριστικών εισόδου σε κάθε κόμβο κάθε δένδρου. Η χρησιμότητα του αλγορίθμου είναι εμφανής σε πλήθος πεδίων καθώς χρησιμοποιείται σε τραπεζικά συστήματα, σε χρηματιστηριακές συναλλαγές, στο ηλεκτρονικό εμπόριο, στην ιατρική και σε διάφορους άλλους τομείς. Τέλος, είναι σημαντικό να αναφερθεί ότι ο αλγόριθμος Random Forest απαιτεί προσεκτική επιλογή παραμέτρων όπως ο αριθμός των δένδρων και το μέγιστο βάθος κάθε δένδρου καθώς με ένα πολύ μεγάλο πλήθος δένδρων ο αλγόριθμος μπορεί να καταστεί μη ερμηνεύσιμος (uninterpretable) αλλά και υπολογιστικά κοστοβόρος.



Σχήμα 9: Παλινδρόμηση με τον αλγόριθμο τυχαίου δάσους

5.2.8 Οι Αλγόριθμοι XGBoost και AdaBoost

Οι αλγόριθμοι XGBoost (Extreme Gradient Boosting) και AdaBoost (Adaptive Boosting) αποτελούν δύο σημαντικές τεχνικές στον τομέα της επιβλεπόμενης μηχανικής μάθησης οι οποίες συνδυάζουν πολλαπλές προβλέψεις άλλων αλγορίθμων, συνήθως δένδρων αποφάσεων, για την κατασκευή ενός ακριβούς και υψηλής απόδοσης μοντέλου και την εξαγωγή μίας τελικής πρόβλεψης από αυτό.

Ο αλγόριθμος XGBoost είναι ένας ευέλικτος αλγόριθμος μηχανικής μάθησης ο οποίος μπορεί να χρησιμοποιηθεί τόσο για παλινδρόμηση όσο και για κατηγοριοποίηση και μπορεί να χειριστεί διαφόρους τύπους δεδομένων. Χρησιμοποιείται ευρέως στον τομέα της επιστήμης των δεδομένων ενώ έχει αποδειχθεί αποτελεσματικός σε αρκετούς διαγωνισμούς μηχανικής μάθησης. Η λειτουργία του περιλαμβάνει την επαναλαμβανόμενη εκπαίδευση δένδρων αποφάσεων, εστιάζοντας κάθε φορά στις περιπτώσεις στις οποίες η πρόβλεψη ήταν λανθασμένη με στόχο την ελαχιστοποίηση αυτού του σφάλματος σε κάθε βήμα. Η διαδικασία αυτή επιτρέπει στο μοντέλο να βελτιώνει σταδιακά την προγνωστική του ικανότητα και ακρίβεια και συνεχίζεται μέχρις ότου να επιτευχθεί ένας προκαθορισμένος αριθμός δένδρων ή «γύρων ενίσχυσης» (boosting rounds). Το βασικό πλεονέκτημα του αλγορίθμου XGBoost είναι ότι περιλαμβάνει σχετικά λίγα δένδρα αποφάσεων με μικρό βάθος, κάτι που καθιστά τον αλγόριθμο εύκολα κατανοητό (ερμηνεύσιμο), αποτελεσματικό και ακριβή.[19] Για να επιτευχθούν τα παραπάνω και να αποφευχθεί το φαινόμενο της υπερπροσαρμογής σημαντικό ρόλο παίζει η επιλογή του κριτηρίου διακοπής της επαναλαμβανόμενης διαδικασίας της ενίσχυσης (boosting). Χρησιμοποιείται σε πλήθος διαφόρων πεδίων ενώ αξίζει να αναφερθεί ότι έχει αναγνωριστεί από το CERN (Ευρωπαϊκό Συμβούλιο Πυρηνικής Έρευνας) ως μία από τις καλύτερες μεθόδους για την κατηγοριοποίηση των σημάτων από τον Μεγάλο Επιταχυντή Αδρονίων (Large Hadron Collider, LHC).

Ο Αλγόριθμος AdaBoost, όμοια με τον αλγόριθμο XGBoost, μπορεί να χρησιμοποιηθεί τόσο για παλινδρόμηση, όσο και για κατηγοριοποίηση καθώς και με διάφορους τύπους δεδομένων. Η βασική ιδέα των δύο αλγορίθμων είναι παρόμοια δεδομένου ότι και ο αλγόριθμος AdaBoost περιλαμβάνει μία επαναλαμβανόμενη διαδικασία κατά την οποία σε κάθε βήμα επικεντρώνεται στα δείγματα που δεν κατηγοριοποιήθηκαν σωστά. Χρησιμοποιεί μικρού βάθους δένδρα, συχνά και με ένα μόνο επίπεδο (decision stumps) με αποτέλεσμα να απαιτεί λιγότερη προεπεξεργασία των δεδομένων και να είναι λιγότερο επιρρεπής στην εμφάνιση του φαινομένου της υπερπροσαρμογής. Είναι εξίσου διαδεδομένος με τον XGBoost και χρησιμοποιείται σε πεδία όπως η όραση υπολογιστών, η επεξεργασία φυσικής γλώσσας και η βιοπληροφορική.

Κεφάλαιο 6: Μη Επιβλεπόμενη Μάθηση

6.1 Ορισμός και Παραδείγματα

Η μη επιβλεπόμενη μηχανική μάθηση χρησιμοποιεί αλγόριθμους μηχανικής μάθησης για την ανάλυση και την ομαδοποίηση μη επισημασμένων συνόλων δεδομένων. Αυτοί οι αλγόριθμοι ανακαλύπτουν σε αυτά τα σύνολα κρυμμένα μοτίβα ή υποομάδες δεδομένων που σχετίζονται με βάση κάποιο χαρακτηριστικό χωρίς την ανάγκη ανθρώπινης παρέμβασης. Η ικανότητα των αλγορίθμων της να ανακαλύπτουν ομοιότητες και διαφορές στις πληροφορίες καθιστά τη μη επιβλεπόμενη μάθηση ιδανική λύση για πληθώρα προβλημάτων όπως η διερευνητική ανάλυση δεδομένων (exploratory data analysis, EDA), ο διαχωρισμός πελατών (customer segmentation) και η αναγνώριση εικόνας. Τα μοντέλα μη επιβλεπόμενης μηχανικής μάθησης χρησιμοποιούνται για τρεις βασικές εργασίες: την ομαδοποίηση (clustering) και τη μείωση διαστάσεων (dimensionality reduction) που αναφέρονται παρακάτω, καθώς και τη συσχέτιση (association).[12]

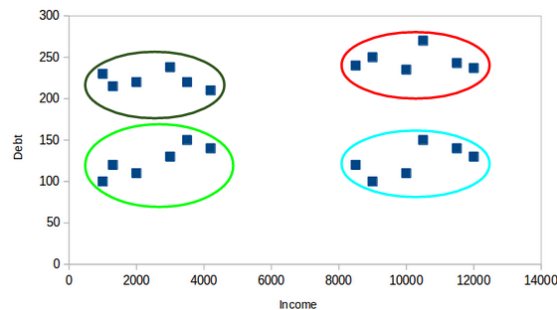
6.2 Ομαδοποίηση

Η ομαδοποίηση (clustering) είναι μια τεχνική η οποία διαχωρίζει σε ομάδες μη επισημασμένα δεδομένα με βάση τις ομοιότητες ή τις διαφορές τους. Οι αλγόριθμοι ομαδοποίησης λοιπόν χρησιμοποιούνται για τον διαχωρισμό ακατέργαστων, μη κατηγοριοποιημένων δεδομένων σε ομάδες (clusters) που προκύπτουν από τη δομή ή από μοτίβα στο σύνολο των δεδομένων. Οι αλγόριθμοι ομαδοποίησης μπορούν να είναι διαφόρων τύπων όπως αποκλειστικοί (exclusive), επικαλυπτόμενοι (overlapping), ιεραρχικοί (hierarchical) και πιθανολογικοί (probabilistic).[40] Μερικοί από τους σημαντικότερους αλγορίθμους ομαδοποίησης αναφέρονται παρακάτω.

6.2.1 Η μέθοδος ομαδοποίησης K-means

Η K-means είναι ένα παράδειγμα μιας αποκλειστικού τύπου μεθόδου ομαδοποίησης στην οποία τα σημεία δεδομένων κατατάσσονται σε K ομάδες, με βάση την απόστασή τους από το κέντρο (centroid) κάθε ομάδας, το οποίο και την ορίζει. Τα σημεία δεδομένων που βρίσκονται πλησιέστερα σε ένα από τα κέντρα των ομάδων κατατάσσονται στην ομάδα που αντιστοιχεί στο κέντρο αυτό (Σχήμα 10). Μια μεγαλύτερη τιμή του K έχει σαν αποτέλεσμα περισσότερες ομάδες μικρότερου μεγέθους άρα κάθε ομάδα θα μπορούσε να χαρακτηριστεί και περισσότερο λεπτομερής, ενώ μια μικρότερη τιμή του K θα έχει σαν αποτέλεσμα λιγότερες ομάδες, μεγαλύτερου μεγέθους όντας έτσι περισσότερο γενικές.[20] Η ομαδοποίηση K-means

χρησιμοποιείται σε τομείς όπως η κατάτμηση της αγοράς (market segmentation), η ομαδοποίηση εγγράφων (document clustering), η κατάτμηση εικόνας (image segmentation) και η συμπίεση εικόνας (image compression).



Σχήμα 10: K-means ομαδοποίηση

6.2.2 Ο αλγόριθμος ομαδοποίησης k -πλησιέστερων γειτόνων

Εκτός από την επιβλεπόμενη μάθηση, όπως αναφέρθηκε στο αντίστοιχο κεφάλαιο, ο αλγόριθμος k -NN μπορεί να χρησιμοποιηθεί και για ομαδοποίηση στη μη επιβλεπόμενη μάθηση. Στην περίπτωση αυτή δεν αντιστοιχεί το υπό εξέταση σημείο σε κάποια ετικέτα, καθώς τα δεδομένα είναι μη επισημασμένα, επιστρέφει όμως τα k κοντινότερα σε αυτό σημεία από το σύνολο των δεδομένων, συνήθως με τη βοήθεια κάποιου αλγορίθμου όπως οι αλγόριθμοι BallTree, KDTree ή Brute Force.

6.2.3 Ιεραρχική ομαδοποίηση

Η ιεραρχική ομαδοποίηση είναι ένας τύπος αλγορίθμων ομαδοποίησης μη επιβλεπόμενης μάθησης και περιλαμβάνει δύο επιμέρους κατηγορίες αλγορίθμων, τους συσσωρευτικούς (agglomerative) και τους διαχωριστικούς (divisive).

Η συσσωρευτική ομαδοποίηση θεωρείται μια «από κάτω προς τα πάνω» προσέγγιση κατά την οποία τα σημεία δεδομένων είναι αρχικά πλήρως απομονωμένα ως ξεχωριστές ομαδοποιήσεις το καθένα και στη συνέχεια συγχωνεύονται με μία επαναληπτική διαδικασία με βάση την ομοιότητά τους μέχρι να σχηματίσουν μία ομάδα. Για τη μέτρηση της ομοιότητας χρησιμοποιούνται συνήθως τέσσερις διαφορετικές μέθοδοι: η μέθοδος Ward's linkage (βασίζεται στο άθροισμα των τετραγώνων), η μέθοδος average linkage (βασίζεται στη μέση απόσταση μεταξύ δύο σημείων της ίδιας ομάδας), η μέθοδος complete (ή maximum) linkage (βασίζεται στη μέγιστη απόσταση μεταξύ δύο σημείων της ίδιας ομάδας) και η μέθοδος single (ή minimum) linkage (βασίζεται στην ελάχιστη απόσταση μεταξύ δύο σημείων της ίδιας

ομάδας). Για τον υπολογισμό αυτών των αποστάσεων χρησιμοποιείται συνήθως η Ευκλείδεια απόσταση και σπανιότερα η απόσταση Manhattan.

Η διαχωριστική ομαδοποίηση μπορεί να οριστεί ως το αντίθετο της συσσωρευτικής ομαδοποίησης. Είναι μία «από πάνω προς τα κάτω» προσέγγιση στην οποία μία ενιαία αρχική ομάδα δεδομένων διαιρείται με βάση τις διαφορές μεταξύ των σημείων δεδομένων που την αποτελούν επαναληπτικά.[21] Η διαιρετική ομαδοποίηση χρησιμοποιείται σπανιότερα, αλλά αξίζει να αναφερθεί στο πλαίσιο της ιεραρχικής ομαδοποίησης.

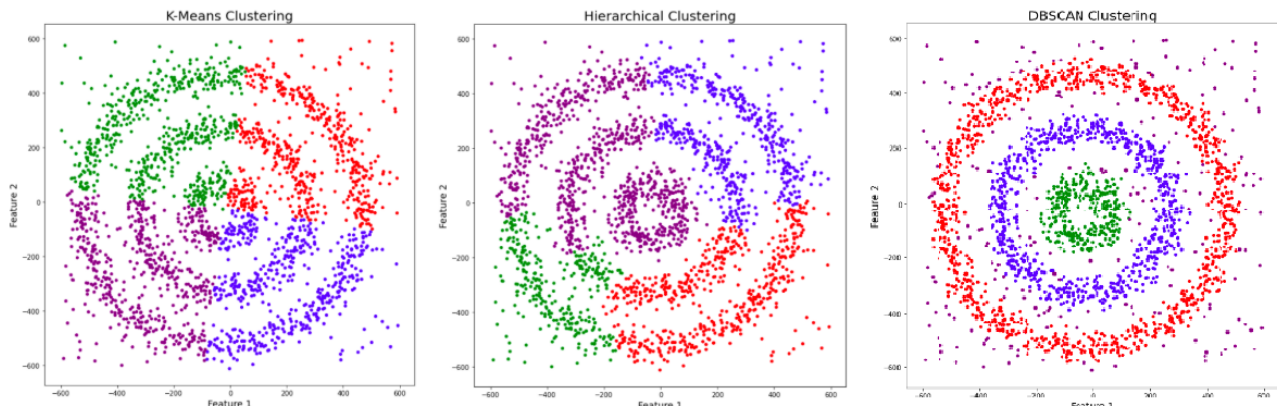
Οι παραπάνω διαδικασίες ομαδοποίησης συνήθως απεικονίζονται μέσω ενός δενδρογραφήματος, ενός διαγράμματος που μοιάζει με δένδρο στο οποίο αποτυπώνεται η συγχώνευση ή η διαίρεση των σημείων δεδομένων σε κάθε επανάληψη (Σχήμα 11).



Σχήμα 11: Ιεραρχική ομαδοποίηση

6.2.4 Ο αλγόριθμος ομαδοποίησης DBSCAN

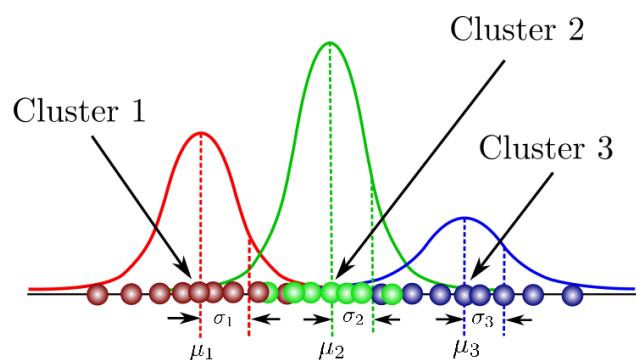
Ο αλγόριθμος DBSCAN (Χωρική ομαδοποίηση εφαρμογών με θόρυβο βάσει πυκνότητας, Density-Based Spatial Clustering of Applications with Noise) είναι ένας δημοφιλής αλγόριθμος ομαδοποίησης που χρησιμοποιείται για την ανάλυση δεδομένων και την αναγνώριση προτύπων. Ομαδοποιεί τα σημεία δεδομένων με βάση την πυκνότητά τους, εντοπίζοντας τις ομάδες σε περιοχές υψηλής πυκνότητας και θεωρώντας τις ακραίες τιμές που βρίσκονται εκτός των ομάδων αυτών ως θόρυβο.[22] Ο αλγόριθμος DBSCAN είναι αποτελεσματικός στην ανακάλυψη ομαδοποιήσεων σε σημεία δεδομένων αυθαίρετου σχήματος υπερέχοντας σε τέτοιες περιπτώσεις του αλγορίθμου k-means ή αλγορίθμων ιεραρχικής ομαδοποίησης (Σχήμα 12). Χρησιμοποιείται ευρέως στην εξόρυξη δεδομένων (data mining), την ανάλυση χωρικών δεδομένων (spatial data analysis) και σε διάφορες εφαρμογές μηχανικής μάθησης.



Σχήμα 12: K-means ομαδοποίηση, Ιεραρχική ομαδοποίηση, DBSCAN (από τα αριστερά προς τα δεξιά)

6.2.5 Μοντέλα Gaussian Mixture

Τα Μοντέλα Gaussian Mixture (Gaussian Mixture Models, GMMs) είναι μία από τις πιο διαδεδομένες μεθόδους πιθανολογικής ομαδοποίησης. Ένα πιθανολογικό μοντέλο είναι μια τεχνική μη επιβλεπόμενης μάθησης στην οποία τα σημεία δεδομένων ομαδοποιούνται με βάση την πιθανότητα να ανήκουν σε κάποια συγκεκριμένη κατανομή. Τα GMMs αποτελούνται από ένα πλήθος συναρτήσεων κατανομής πιθανότητας (Σχήμα 13) και χρησιμοποιούνται κυρίως για τον προσδιορισμό της Γκαουσιανής ή της κανονικής κατανομής στην οποία ανήκει ένα συγκεκριμένο σημείο. Εάν παράμετροι όπως ο μέσος όρος (μ) και η διακύμανση (σ^2) ήταν γνωστές, τότε θα μπορούσαμε να προσδιορίσουμε σε ποια κατανομή ανήκει ένα σημείο δεδομένων. Ωστόσο, στα GMMs, αυτές οι παράμετροι δεν είναι γνωστές και για την τιμή τους γίνεται μία εκτίμηση από το μοντέλο. Συχνά στα GMMs χρησιμοποιείται ο αλγόριθμος μεγιστοποίησης προσδοκίων (Expectation-Maximization, EM) για την εκτίμηση της πιθανότητας ένα συγκεκριμένο σημείο δεδομένων να ανήκει σε μία ομάδα δεδομένων. Χρησιμοποιούνται σε τομείς όπως η ανίχνευση απάτης και παράνομης δραστηριότητας και η ομαδοποίηση εικόνων.



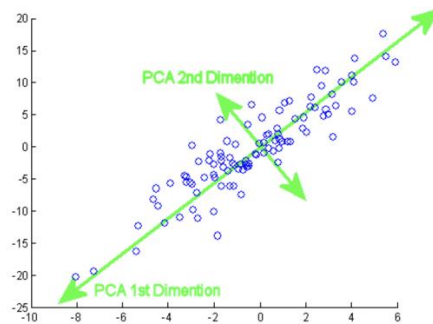
Σχήμα 13: Μοντέλο Gaussian Mixture

6.3 Μείωση Διαστάσεων

Ενώ συνήθως περισσότερα δεδομένα οδηγούν σε ακριβέστερα αποτελέσματα, ένα μεγάλο πλήθος δεδομένων είναι πιθανόν να μειώσει την απόδοση των αλγορίθμων μηχανικής μάθησης (για παράδειγμα με την εμφάνιση του φαινομένου της υπερπροσαρμογής) καθώς και να καταστήσει δυσκολότερη την απεικόνιση του συνόλου δεδομένων σε αντίθεση με ένα απλούστερο και εύκολα κατανοητό γράφημα. Η μείωση διαστάσεων (Dimensionality Reduction) είναι μια τεχνική που χρησιμοποιείται σε τέτοιες περιπτώσεις, όταν δηλαδή ο αριθμός των χαρακτηριστικών (ή διαστάσεων) σε ένα σύνολο δεδομένων είναι πολύ υψηλός.[23,41] Μειώνει λοιπόν τον αριθμό των δεδομένων εισόδου σε ένα διαχειρίσιμο μέγεθος, ενώ παράλληλα διατηρεί την απαραίτητη πληροφορία που εμπεριέχεται στο σύνολο των δεδομένων. Ταυτόχρονα μειώνει και την πολυπλοκότητα του μοντέλου. Εφαρμόζεται συνήθως στο στάδιο της προεπεξεργασίας δεδομένων με διάφορες μεθόδους, μερικές από τις οποίες αναφέρονται παρακάτω.

6.3.1 Ανάλυση κύριων συνιστωσών

Η Ανάλυση κύριων συνιστωσών (Principal component analysis, PCA) είναι μια μέθοδος του πεδίου της στατιστικής που επιτρέπει τη μείωση των χαρακτηριστικών (ή διαστάσεων) ενός συνόλου δεδομένων με έναν γραμμικό μετασχηματισμό και τη μεταφορά σε ένα νέο σύστημα συντεταγμένων. Στο νέο αυτό σύστημα το μεγαλύτερο μέρος της ποικιλίας ανάμεσα στα δεδομένα και των σχέσεων που τα συνδέουν μπορεί να αποτυπωθεί από λιγότερες διαστάσεις σε σχέση με τα αρχικά δεδομένα.[24] Συχνά στη μέθοδο αυτή χρησιμοποιούνται οι δύο πρώτες κύριες συνιστώσες, κάτι που επιτρέπει την απεικόνιση των δεδομένων σε δύο διαστάσεις και κατ' επέκταση τον ευκολότερο εντοπισμό (οπτικά) πιθανών ομαδοποιήσεων των σημείων δεδομένων που συνδέονται με βάση κάποιο χαρακτηριστικό τους (Σχήμα 14). Γενικά, οι κύριες συνιστώσες μιας συλλογής σημείων δεδομένων σε έναν πραγματικό χώρο συντεταγμένων είναι μια ακολουθία μοναδιαίων διανυσμάτων p , όπου το i διάνυσμα αντιστοιχεί στην κατεύθυνση μιας γραμμής που «ταιριάζει» καλύτερα στα δεδομένα και είναι κάθετη στα $i - 1$ προηγούμενα διανύσματα. Εδώ, η γραμμή που «ταιριάζει» καλύτερα ορίζεται ως αυτή που ελαχιστοποιεί τη μέση τετραγωνική κάθετη απόσταση από τα σημεία στη γραμμή. Αυτές οι κατευθύνσεις αποτελούν μια ορθοκανονική βάση στην οποία οι διαφορετικές διαστάσεις των δεδομένων είναι γραμμικά ανεξάρτητες μεταξύ τους. Η μέθοδος αυτή έχει εφαρμογές σε πολλούς τομείς όπως η γενετική πληθυσμών (population genetics), η μελέτη μικροβιοκοινοτήτων (microbiome) και οι ατμοσφαιρικές επιστήμες.



Σχήμα 14: Ανάλυση κύριων συνιστωσών

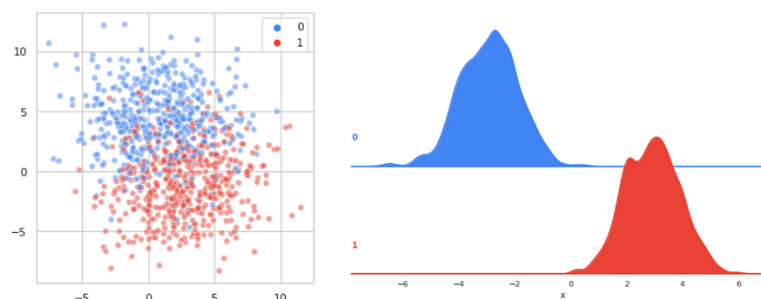
6.3.2 Ανάλυση ανεξάρτητων συνιστωσών

Η Ανάλυση ανεξάρτητων συνιστωσών (Independent component analysis, ICA) είναι επίσης μία μέθοδος μη επιβλεπόμενης μάθησης με σκοπό τη μείωση διαστάσεων και η οποία στην ουσία επιχειρεί να αποσυνθέσει ένα σήμα πολλών μεταβλητών σε ένα άθροισμα ανεξάρτητων, μη Γκαουσιανών σημάτων.[25] Για παράδειγμα, ο ήχος είναι συνήθως ένα σήμα που αποτελείται από το αλγεβρικό άθροισμα σημάτων από διάφορες πηγές. Το ερώτημα λοιπόν είναι αν είναι δυνατόν να διαχωριστούν αυτές οι συνιστώσες από το παρατηρούμενο συνολικό σήμα. Όταν ισχύει η υπόθεση της στατιστικής ανεξαρτησίας, η μέθοδος ICA εφαρμοζόμενη σε ένα μικτό σήμα δίνει πολύ καλά αποτελέσματα. Μια απλή εφαρμογή της ICA είναι το “cocktail party problem” στο οποίο σε ένα δείγμα δεδομένων που αποτελείται από πολλά άτομα που μιλούν ταυτόχρονα σε ένα δωμάτιο, τα διαφορετικά σήματα ομιλίας απομονώνονται. Συνήθως το πρόβλημα απλοποιείται υποθέτοντας ότι δεν υπάρχουν χρονικές καθυστερήσεις ή ηχώ. Το γεγονός ότι ο διαχωρισμός των μικτών σημάτων από την ICA δίνει πολύ καλά αποτελέσματα βασίζεται σε δύο υποθέσεις: τα σήματα πηγής είναι ανεξάρτητα το ένα από το άλλο και οι τιμές σε κάθε σήμα πηγής δεν ανήκουν σε Γκαουσιανές κατανομές. Πρέπει να σημειωθεί ότι, αν και οι τεχνικές PCA και ICA είναι παρόμοιες, στην πραγματικότητα αποτελούν δύο διαφορετικές προσεγγίσεις και εκτελούν διαφορετικά καθήκοντα. Η μέθοδος ICA χρησιμοποιείται σε πληθώρα πεδίων όπως τα συστήματα αναγνώρισης προσώπου, οι τηλεπικοινωνίες, τα χρηματοοικονομικά, και η αστρονομία.

6.3.3 Γραμμική διακριτική ανάλυση

Η Γραμμική διακριτική ανάλυση (Linear Discriminant Analysis, LDA) είναι μία μέθοδος μείωσης διαστάσεων που χρησιμοποιείται σε προβλήματα επιβλεπόμενης μάθησης για εργασίες κατηγοριοποίησης. Σχετίζεται όμως στενά με τη μέθοδο PCA, καθώς και οι δύο αναζητούν γραμμικούς συνδυασμούς μεταβλητών που εξηγούν καλύτερα τις σχέσεις μεταξύ των δεδομένων. Η LDA επιχειρεί λεπτομερώς να μοντελοποιήσει τις διαφορές μεταξύ των

κατηγοριών στις οποίες κατηγοριοποιούνται τα δεδομένα, ενώ η PCA δε λαμβάνει υπόψη τις διαφορές αυτές. Είναι λοιπόν μια τεχνική που χρησιμοποιείται για την εύρεση ενός γραμμικού συνδυασμού χαρακτηριστικών που διαχωρίζει καλύτερα τις τάξεις σε ένα σύνολο δεδομένων.[26] Λειτουργεί προβάλλοντας τα δεδομένα σε ένα χώρο χαμηλότερης διάστασης η οποία μεγιστοποιεί τον διαχωρισμό μεταξύ των διαφόρων τάξεων. Με άλλα λόγια, βρίσκει τις κατευθύνσεις στο χώρο χαρακτηριστικών που διαχωρίζουν καλύτερα τις διαφορετικές κατηγορίες δεδομένων. Η LDA υποθέτει ότι τα δεδομένα ακολουθούν Γκαουσιανή κατανομή και ότι οι πίνακες συνδιακύμανσης (covariance) των διαφορετικών τάξεων είναι ίσοι. Υποθέτει επίσης ότι τα δεδομένα είναι γραμμικά διαχωρίσιμα, πράγμα που σημαίνει ότι ένα γραμμικό όριο απόφασης μπορεί να διαχωρίσει με ακρίβεια τις διαφορετικές τάξεις (Σχήμα 15). Παρά το γεγονός ότι απαιτεί όλες τις παραπάνω υποθέσεις με αποτέλεσμα να μην μπορεί να εφαρμοστεί πάντα, είναι εν γένει ένας απλός και αποτελεσματικός αλγόριθμος ο οποίος μπορεί να λειτουργήσει καλά ακόμα και όταν ο αριθμός των χαρακτηριστικών είναι πολύ μεγαλύτερος από τον αριθμό των δειγμάτων εκπαίδευσης. Εφαρμόζεται σε τομείς όπως τα συστήματα αναγνώρισης προσώπου, οι γεωεπιστήμες και οι επιστήμες βιοιατρικής.

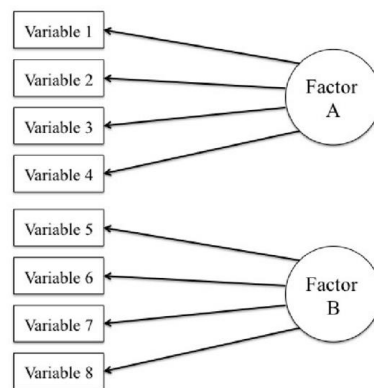


Σχήμα 15: Γραμμική διακριτική ανάλυση

6.3.4 Ανάλυση παραγόντων

Η Ανάλυση παραγόντων (Factor Analysis, FA) είναι μία μέθοδος μη επιβλεπόμενης μάθησης που χρησιμοποιείται επίσης για τη μείωση διαστάσεων. Πρόκειται για ένα γραμμικό μοντέλο στατιστικής που έχει σαν στόχο να περιγράψει την ποικιλία μεταξύ των παρατηρούμενων μεταβλητών μέσα από έναν δυνητικά μικρότερο αριθμό μη παρατηρούμενων μεταβλητών που ονομάζονται παράγοντες. Για παράδειγμα, είναι πιθανό οι διακυμάνσεις μεταξύ οκτώ παρατηρούμενων μεταβλητών να αντικατοπτρίζονται κυρίως στις διακυμάνσεις μεταξύ δύο μη παρατηρούμενων μεταβλητών (Σχήμα 16). Οι παρατηρούμενες μεταβλητές μοντελοποιούνται ως γραμμικοί συνδυασμοί των παραγόντων αφού εισαχθούν και κάποιοι όροι σφάλματος. Η μοντελοποίηση αυτή επιτρέπει την ποσοτικοποίηση του βαθμού στον οποίο μία μεταβλητή σχετίζεται με έναν συγκεκριμένο παράγοντα.[27] Η λογική πίσω από τις μεθόδους ανάλυσης παραγόντων είναι ότι οι πληροφορίες που αποκτώνται σχετικά με τις

αλληλεξαρτήσεις μεταξύ των παρατηρούμενων μεταβλητών μπορούν να χρησιμοποιηθούν αργότερα για να μειωθεί το σύνολο των διαστάσεων σε ένα σύνολο δεδομένων. Σχετίζεται επίσης με τις μεθόδους PCA και LDA, με τη διαφορά ότι η FA δημιουργεί τους γραμμικούς συνδυασμούς των μεταβλητών εστιάζοντας περισσότερο στις διαφορές παρά στις ομοιότητες. Η ανάλυση παραγόντων χρησιμοποιείται συνήθως στην ψυχομετρία, την ψυχολογία της προσωπικότητας, τη βιολογία, το μάρκετινγκ, τη διαχείριση προϊόντων, την έρευνα επιχειρήσεων και τη χρηματοδότηση.



Σχήμα 16: Ανάλυση παραγόντων

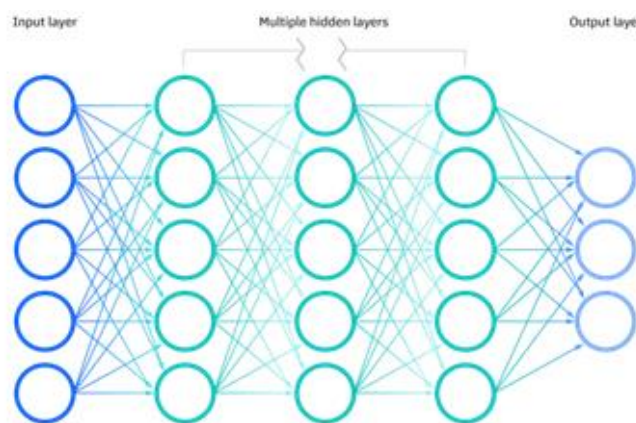
Κεφάλαιο 7: Νευρωνικά Δίκτυα

7.1 Εισαγωγή

Ένα νευρωνικό δίκτυο είναι ένας τύπος μοντέλου μηχανικής μάθησης εμπνευσμένος από τη δομή και τη λειτουργία του ανθρώπινου εγκεφάλου. Είναι ένα δίκτυο τεχνητών νευρώνων οι οποίοι συνδέονται μεταξύ τους και μπορούν να εκπαιδευτούν για τη διεκπεραίωση διαφόρων εργασιών όπως η αναγνώριση εικόνας, η επεξεργασία φυσικής γλώσσας και η πραγματοποίηση προβλέψεων. Το δίκτυο αυτό αποτελείται από επίπεδα (ή στρώματα) κόμβων (τεχνητοί νευρώνες), οι οποίοι συνδέονται μεταξύ τους με ακμές που χαρακτηρίζονται από κάποιο βάρος. Κάθε νευρώνας λαμβάνει είσοδο από άλλους νευρώνες, εκτελεί μια μαθηματική λειτουργία στα δεδομένα εισόδου και στη συνέχεια παράγει μία έξοδο, στέλνοντάς την σε άλλους νευρώνες στο επόμενο επίπεδο. Η διαδικασία συνεχίζεται μέχρι να παραχθεί η τελική έξοδος. Το βασικό χαρακτηριστικό ενός νευρωνικού δικτύου είναι η εγγενής ικανότητα μάθησης μέσω της εκπαίδευσης, δηλαδή μιας επαναληπτικής διαδικασίας προσαρμογής των παραμέτρων του δικτύου όπως τα βάρη, με στόχο την ελαχιστοποίηση του σφάλματος των προβλέψεων. Υπάρχουν διάφοροι τύποι νευρωνικών δικτύων όπως τα τροφοδοτούμενα προς τα εμπρός (feedforward), τα αναδρομικά (recurrent) και τα συνελκτικά (convolutional) και κάθε τύπος δικτύου είναι κατάλληλος για διαφορετικού τύπου εργασίες και δεδομένα.[28] Για παράδειγμα, τα τροφοδοτούμενα προς τα εμπρός χρησιμοποιούνται για απλές εργασίες όπως παλινδρόμηση και κατηγοριοποίηση, τα αναδρομικά χρησιμοποιούνται σε χρονικά μεταβαλλόμενα μοντέλα και σε εργασίες όπως η επεξεργασία φυσικής γλώσσας ενώ τα συνελκτικά χρησιμοποιούνται σε εργασίες όπως η επεξεργασία εικόνας και βίντεο. Εν κατακλείδι, τα νευρωνικά δίκτυα έχουν εφαρμογές σε μεγάλο πλήθος πεδίων καθώς είναι σε θέση να μαθαίνουν από μεγάλες ποσότητες δεδομένων και να χειρίζονται μη γραμμικές και πολύπλοκες σχέσεις μεταξύ εισόδων και εξόδων.

7.2 Δομή νευρωνικού δικτύου

Η δομή ενός νευρωνικού δικτύου περιλαμβάνει τρεις τύπους επιπέδων: ένα επίπεδο εισόδων, τα κρυμμένα (ενδιάμεσα) επίπεδα και ένα επίπεδο εξόδων (Σχήμα 17). Το επίπεδο εισόδου λαμβάνει τα δεδομένα εισόδου και τα μεταφέρει στα κρυμμένα επίπεδα. Ο αριθμός των νευρώνων στο επίπεδο εισόδου είναι ίσος με τον αριθμό όλων των χαρακτηριστικών (μεταβλητών) στα δεδομένα εισόδου. Τα κρυμμένα επίπεδα είναι τα επίπεδα μεταξύ των επιπέδων εισόδου και εξόδου και ονομάζονται έτσι επειδή οι εσωτερικοί υπολογισμοί τους δεν είναι άμεσα ορατοί στον χρήστη. Ένα νευρωνικό δίκτυο μπορεί να έχει ένα ή περισσότερα κρυμμένα επίπεδα. Τέλος, το επίπεδο εξόδου είναι το επίπεδο που παράγει την τελική πρόβλεψη ή έξοδο του νευρωνικού δικτύου. Ο αριθμός των νευρώνων στο στρώμα εξόδου είναι ίσος με τον αριθμό των πιθανών αποτελεσμάτων ή τάξεων. Γενικά, η δομή του νευρωνικού δικτύου καθορίζεται από τον αριθμό των επιπέδων, τον αριθμό των νευρώνων σε κάθε επίπεδο και τον τύπο των συνδέσεων μεταξύ των επιπέδων. Ο σχεδιασμός του νευρωνικού δικτύου καθορίζεται από το εκάστοτε πρόβλημα καθώς και την ποσότητα των διαθέσιμων δεδομένων για εκπαίδευση.



Σχήμα 17: Δομή νευρωνικού δικτύου

7.3 Απλά και βαθιά νευρωνικά δίκτυα

Τα νευρωνικά δίκτυα διακρίνονται σε απλά (shallow) νευρωνικά δίκτυα και σε βαθιά (deep) νευρωνικά δίκτυα. Ένα απλό νευρωνικό δίκτυο αποτελείται από έναν μικρό αριθμό επιπέδων, με συνήθως ένα ή δύο κρυμμένα επίπεδα. Τα νευρωνικά δίκτυα αυτά είναι απλούστερα στη δομή τους και περιλαμβάνουν λιγότερες παραμέτρους, γεγονός που τα καθιστά υπολογιστικά απλούστερα καθώς και ευκολότερα στην εκπαίδευση σε σχέση με τα βαθιά νευρωνικά δίκτυα. Από την άλλη πλευρά, τα βαθιά νευρωνικά δίκτυα είναι νευρωνικά δίκτυα με μεγάλο αριθμό επιπέδων, συνήθως δύο ή περισσότερα κρυμμένα επίπεδα. Τα βαθιά νευρωνικά δίκτυα είναι

πιο περίπλοκα στη δομή τους και περιλαμβάνουν περισσότερες παραμέτρους όντας έτσι υπολογιστικά περισσότερο κοστοβόρα. Ωστόσο, είναι ικανά να μοντελοποιήσουν πιο σύνθετες και μη γραμμικές σχέσεις μεταξύ εισόδων και εξόδων επιτρέποντας έτσι τη διεκπεραίωση πιο σύνθετων εργασιών.[29] Κατά την εκπαίδευση, τα πρώτα επίπεδα αναγνωρίζουν χαρακτηριστικά χαμηλού επιπέδου, για παράδειγμα γραμμές, τα επόμενα επίπεδα αναγνωρίζουν χαρακτηριστικά υψηλότερου επιπέδου όπως σχήματα και τα τελικά επίπεδα αναγνωρίζουν πιο αφηρημένα χαρακτηριστικά όπως μέρη αντικειμένων ή και ολόκληρα αντικείμενα. Γενικά τα απλά νευρωνικά δίκτυα είναι προτιμότερα σε απλά προβλήματα με μικρά σύνολα δεδομένων ενώ τα βαθιά νευρωνικά δίκτυα είναι καταλληλότερα σε σύνθετα προβλήματα με μεγάλα σύνολα δεδομένων. Η επιλογή ανάμεσά τους εξαρτάται από την πολυπλοκότητα του προβλήματος, τον όγκο των δεδομένων και τους διαθέσιμους υπολογιστικούς πόρους.

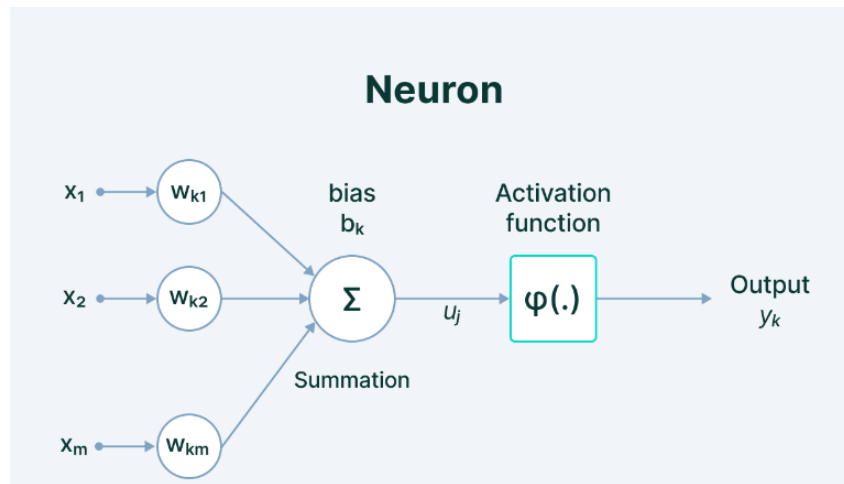
7.4 Δομή νευρώνα

Όπως έχει αναφερθεί, βασικό δομικό στοιχείο ενός νευρωνικού δικτύου είναι οι νευρώνες που το αποτελούν. Ένας νευρώνας είναι μια υπολογιστική μονάδα που εκτελεί μια συγκεκριμένη λειτουργία στην είσοδό της και παράγει μια έξοδο. Βασικά συστατικά ενός νευρώνα είναι ένας σταθμισμένος μέσος όρος (weighted average) και μια συνάρτηση ενεργοποίησης (activation function).

Ο νευρώνας τροφοδοτείται με δεδομένα (εισόδους) που σχετίζονται το καθένα με κάποιον συντελεστή βαρύτητας, οι οποίοι, σε συνδυασμό με έναν όρο προκατάληψης (bias) που προστίθεται σε αυτά, αποτελούν και τις παραμέτρους του νευρωνικού δικτύου.[30] Η τιμή των παραμέτρων αυτών διαμορφώνεται κατά τη διάρκεια της εκπαίδευσης του μοντέλου.

Τέλος, η συνάρτηση ενεργοποίησης είναι μια μαθηματική συνάρτηση που εφαρμόζεται στην έξοδο του νευρώνα. Επιτρέπει στον νευρώνα να μοντελοποιεί πιο σύνθετες σχέσεις μεταξύ των εισόδων και των εξόδων αφού εισάγει ουσιαστικά την ιδιότητα της μη γραμμικότητας στο μοντέλο. Διαδεδομένες συναρτήσεις ενεργοποίησης περιλαμβάνουν τη σιγμοειδή (sigmoid), τη συνάρτηση διορθωμένης γραμμικής μονάδας (Rectified Linear Unit, ReLU) και την υπερβολική εφαπτομένη (tanh).

Συνοπτικά, ένας νευρώνας ή μία νευρωνική μονάδα (neural unit) είναι μια υπολογιστική μονάδα που παίρνει ένα διάνυσμα εισόδου, εφαρμόζει ένα σύνολο βαρών σε αυτό και, στη συνέχεια, εφαρμόζει μια συνάρτηση ενεργοποίησης στο αποτέλεσμα. Η έξοδος της νευρωνικής μονάδας μεταφέρεται στη συνέχεια σε άλλες μονάδες του δικτύου για περαιτέρω επεξεργασία. Οι συντελεστές βαρύτητας και προκατάληψης διαμορφώνονται κατά τη διάρκεια της διαδικασίας της εκπαίδευσης, με τρόπο τέτοιο ώστε η νευρωνική μονάδα να μπορεί τελικά να αναγνωρίζει μοτίβα στα δεδομένα και να παράγει ακριβείς προβλέψεις.



Σχήμα 18: Δομή νευρώνα

7.5 Συνάρτηση κόστους

Μια συνάρτηση κόστους, είναι μια μαθηματική συνάρτηση που υπολογίζει τη διαφορά μεταξύ της προβλεφθείσας εξόδου ενός μοντέλου και της πραγματικής εξόδου. Ο στόχος της εκπαίδευσης ενός μοντέλου μηχανικής μάθησης είναι να βρεθεί το σύνολο των παραμέτρων που ελαχιστοποιούν την τιμή αυτής της συνάρτησης κόστους, έτσι ώστε το μοντέλο να μπορεί να κάνει ακριβείς προβλέψεις για νέα δεδομένα.

Η βασική ιδέα είναι ότι η συνάρτηση κόστους χρησιμοποιείται για τη μέτρηση της απόδοσης του μοντέλου σε κάποιο σύνολο δεδομένων με στόχο την ελαχιστοποίηση της τιμής της. Δηλαδή όσο μικρότερη είναι η τιμή της συνάρτησης κόστους, τόσο καλύτερο είναι το μοντέλο στο να κάνει προβλέψεις.

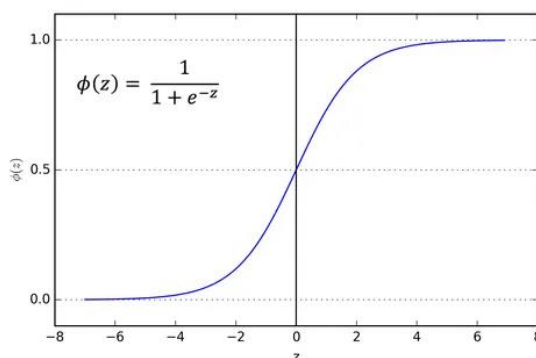
Υπάρχουν διάφοροι τύποι συναρτήσεων κόστους, ο καθένας κατάλληλος για διαφορετικούς τύπους προβλημάτων.[31] Μερικοί από τους σημαντικότερους είναι οι εξής: η συνάρτηση μέσου τετραγωνικού σφάλματος (Mean Squared Error, MSE) η οποία χρησιμοποιείται σε προβλήματα παλινδρόμησης και υπολογίζεται ως ο μέσος όρος των τετραγωνικών διαφορών μεταξύ των προβλεφθεισών και των πραγματικών τιμών, η συνάρτηση Binary Cross-Entropy η οποία χρησιμοποιείται σε προβλήματα δυαδικής κατηγοριοποίησης και η Categorical Cross-Entropy η οποία χρησιμοποιείται για προβλήματα κατηγοριοποίησης σε περισσότερες διαφορετικές τάξεις. Η επιλογή της καταλληλότερης συνάρτησης κόστους εξαρτάται από την εκάστοτε εργασία και τη φύση των δεδομένων.

7.6 Συνάρτηση ενεργοποίησης

Όπως έχει αναφερθεί, μια συνάρτηση ενεργοποίησης είναι μια μαθηματική συνάρτηση που εφαρμόζεται στην έξοδο του νευρώνα σε ένα νευρωνικό δίκτυο. Η σημασία των συναρτήσεων ενεργοποίησης για το νευρωνικό δίκτυο έγκειται στο γεγονός ότι δίχως αυτές το νευρωνικό δίκτυο δε θα ήταν παρά ένα απλό γραμμικό μοντέλο και επομένως δεν θα ήταν σε θέση να χειριστεί πολύπλοκες, μη γραμμικές σχέσεις μεταξύ των δεδομένων.[32]

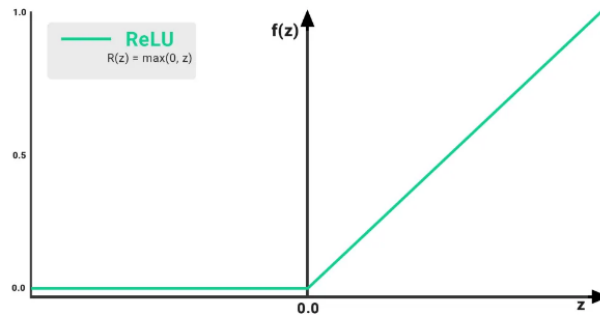
Υπάρχουν διάφοροι τύποι συναρτήσεων ενεργοποίησης και η επιλογή του καταλληλότερου εξαρτάται από το εκάστοτε πρόβλημα και τη φύση των δεδομένων. Στις περισσότερες περιπτώσεις ένας συνδυασμός διαφορετικών συναρτήσεων ενεργοποίησης χρησιμοποιείται στα διαφορετικά επίπεδα του νευρωνικού δικτύου για να επιτευχθεί καλύτερη απόδοση. Ορισμένες από τις σημαντικότερες συναρτήσεις ενεργοποίησης είναι οι εξής:

Σιγμοειδής: Είναι μία συνάρτηση συνεχής και διαφορίσιμη και η γραφική της παράσταση μοιάζει με το γράμμα «S», απ' όπου προκύπτει και η ονομασία της (Sigmoid). Λαμβάνει τιμές μεταξύ του 0 και του 1 και, ως εκ τούτου, χρησιμοποιείται συχνά σε μοντέλα πρόβλεψης κάποιας πιθανότητας (Σχήμα 19).



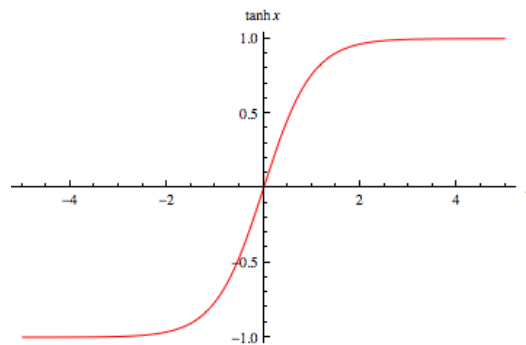
Σχήμα 19: Σιγμοειδής συνάρτηση

ReLU: Είναι η πλέον διαδεδομένη συνάρτηση ενεργοποίησης σε ό,τι αφορά τη βαθιά μάθηση. Για αρνητικές εισόδους η έξοδος λαμβάνει μηδενική τιμή, ενώ για μη αρνητικές η έξοδος ταυτίζεται με την είσοδο. Το κύριο πλεονέκτημα της είναι ότι δεν ενεργοποιεί όλους τους νευρώνες ταυτόχρονα, καθώς όταν η είσοδος είναι αρνητική, η μηδενική έξοδος έχει σαν αποτέλεσμα ουσιαστικά τη μη ενεργοποίηση του αντίστοιχου νευρώνα καθιστώντας έτσι το πρόβλημα υπολογιστικά απλούστερο (Σχήμα 20).



Σχήμα 20: Συνάρτηση ReLU

Υπερβολική εφαπτομένη: είναι παρόμοια με τη σιγμοειδή καθώς είναι συνεχής και διαφορίσιμη με τη διαφορά ότι λαμβάνει τιμές μεταξύ -1 και 1 ενώ η γραφική της παράσταση είναι συμμετρική ως προς την αρχή των αξόνων. Τέλος, ο μέσος όρος των εξόδων είναι 0 (Σχήμα 21).



Σχήμα 21: Υπερβολική εφαπτομένη

7.7 Μέθοδος οπισθοδρόμησης

Η μέθοδος οπισθοδρόμησης (Backpropagation) είναι ένας αλγόριθμος επιβλεπόμενης μάθησης τεχνητών νευρωνικών δικτύων ο οποίος χρησιμοποιεί τη μέθοδο gradient descent, έναν μαθηματικό αλγόριθμο για την εύρεση τοπικών ελαχίστων σε μία διαφορίσιμη συνάρτηση. Δεδομένου ενός τεχνητού νευρωνικού δικτύου και μιας συνάρτησης κόστους, η μέθοδος υπολογίζει την κλίση της συνάρτησης κόστους. Καλείται οπισθοδρόμηση καθώς ο υπολογισμός της κλίσης προχωρά προς τα πίσω στα επίπεδα του νευρωνικού δικτύου, με την κλίση που αντιστοιχεί στο τελευταίο επίπεδο των βαρών να υπολογίζεται πρώτη και την κλίση που αντιστοιχεί στο πρώτο επίπεδο βαρών να υπολογίζεται τελευταία.[33] Στον αλγόριθμο αυτόν, υπολογισμοί ενός επιπέδου επαναχρησιμοποιούνται στον υπολογισμό της κλίσης του προηγούμενου επιπέδου. Αυτή η αντίστροφη ροή των πληροφοριών σφάλματος αποτελεί μία αποτελεσματική μέθοδο υπολογισμού της κλίσης που αντιστοιχεί σε κάθε

επίπεδο. Ουσιαστικά ο στόχος του αλγορίθμου είναι η ελαχιστοποίηση της συνάρτησης κόστους με την κατάλληλη προσαρμογή των συντελεστών βαρύτητας και προκατάληψης από τους οποίους και εξαρτάται. Χρησιμοποιείται ευρέως στη μάθηση βαθιών νευρωνικών δικτύων σε τομείς όπως η αναγνώριση εικόνας και η αναγνώριση ομιλίας.

7.8 Εκπαίδευση βαθιών νευρωνικών δικτύων

Η εκπαίδευση βαθιών νευρωνικών δικτύων (Deep Neural Networks, DNNs) περιλαμβάνει διάφορα βήματα μετά τη συλλογή και προεπεξεργασία των δεδομένων, διαδικασία που έχει προαναφερθεί. Σημαντική είναι η επιλογή της κατάλληλης αρχιτεκτονικής στο μοντέλο, δηλαδή του τύπου του μοντέλου (πχ. τροφοδοτούμενο προς τα εμπρός, αναδρομικό, συνελικτικό), του πλήθους των ενδιάμεσων επιπέδων και των συναρτήσεων ενεργοποίησης ενώ ακολουθεί η αρχικοποίηση των συντελεστών σε κάποια τιμή. Έπειτα, η διαδικασία περιλαμβάνει την επαναλαμβανόμενη εισαγωγή δεδομένων στο επίπεδο εισόδων, τον υπολογισμό της εξόδου του επιπέδου αυτού, την προώθηση της εκάστοτε εξόδου στα επόμενα επίπεδα μέχρι και το τελικό επίπεδο εξόδων, τον υπολογισμό του σφάλματος και τη μεταβολή των συντελεστών ξεκινώντας από το τελευταίο επίπεδο και καταλήγοντας στο αρχικό διατρέχοντας ένα-ένα τα επίπεδα (οπισθοδρόμηση). Με βάση τη συνάρτηση κόστους και ορισμένων προδιαγραφών η εκπαίδευση του μοντέλου ολοκληρώνεται με στόχο την επίτευξη μιας ικανοποιητικής ακρίβειας αλλά και την ταυτόχρονη αποφυγή του φαινομένου της υπερπροσαρμογής.

7.9 Αξιολόγηση απόδοσης βαθιών νευρωνικών δικτύων

Η αξιολόγηση απόδοσης των DNNs είναι η διαδικασία μέτρησης της ακρίβειας και της αποτελεσματικότητας ενός εκπαιδευμένου μοντέλου DNN και για την επίτευξή της χρησιμοποιούνται διάφορες τεχνικές και δείκτες. Για παράδειγμα σε μοντέλα κατηγοριοποίησης, μελετώνται δείκτες που προκύπτουν από έναν πίνακα (Confusion Matrix) όπως το ποσοστό σωστών προβλέψεων του μοντέλου ως προς το σύνολο των προβλέψεων που πραγματοποίησε (accuracy), το ποσοστό των επαληθευμένων προβλέψεων για μία κατηγορία ως προς το σύνολο των προβλέψεων για την κατηγορία αυτή (precision), το ποσοστό των επαληθευμένων προβλέψεων για μία κατηγορία ως προς το σύνολο όλων των δεδομένων που ανήκουν πράγματι στην κατηγορία αυτή (recall) καθώς και ο δείκτης F1 ο οποίος συνδυάζει τους δύο προαναφερθέντες δείκτες. Από την άλλη, στα μοντέλα παλινδρόμησης δείκτες αξιολόγησης απόδοσης του νευρωνικού δικτύου αποτελούν το μέσο απόλυτο σφάλμα (Mean Absolute Error, MAE), το μέσο τετραγωνικό σφάλμα (Mean Squared Error, MSE) ή η τετραγωνική ρίζα αυτού (RMSE) καθώς και ο δείκτης R^2 . Είναι σημαντικό να αναφερθεί ότι η αξιολόγηση της απόδοσης ενός μοντέλου DNN παρέχει πληροφορίες για τα

πλεονεκτήματα και τις αδυναμίες του μοντέλου και μπορεί να αξιοποιηθεί για τη βελτίωση του μοντέλου ανάλογα πάντα με τη λειτουργία που καλείται να επιτελέσει.

Κεφάλαιο 8: Πειραματικό μέρος

8.1 Εισαγωγή

Η ναυτιλιακή βιομηχανία, ένα σημαντικό κομμάτι της παγκόσμιας οικονομίας, αντιμετωπίζει προκλήσεις σχετικά με τη διαχείριση του στόλου, τη μείωση της κατανάλωσης καυσίμου και γενικότερα τη μείωση του λειτουργικού της κόστους. Για την αντιμετώπιση αυτών των ζητημάτων, η χρήση προηγμένων τεχνολογιών όπως το Διαδίκτυο των Πραγμάτων (Internet of Things, IoT) έχει καταστεί απαραίτητη. Οι αλγόριθμοι μηχανικής μάθησης, ιδιαίτερα τα τεχνητά νευρωνικά δίκτυα (Artificial Neural Networks, ANNs), έχουν αναδειχθεί ως ισχυρά εργαλεία για την επεξεργασία μεγάλων συνόλων δεδομένων. Ωστόσο, οι παραδοσιακές μέθοδοι βελτιστοποίησης εμφανίζουν περιορισμούς και απαιτούν περαιτέρω εξερεύνηση σε ό,τι αφορά τις ML τεχνικές.

Η εκπαίδευση ML λειτουργεί, όπως έχει αναφερθεί, με τρεις βασικούς τρόπους: επιβλεπόμενη μάθηση με επισημασμένα δεδομένα, μη επιβλεπόμενη μάθηση με ανίχνευση μοτίβων και (βαθιά) ενισχυτική μάθηση που περιλαμβάνει αλληλεπίδραση μέσω ανταμοιβών και ποινών. Στον ναυτιλιακό τομέα, οι προκλήσεις της συλλογής δεδομένων από διαφορετικές, γεωγραφικά διασκορπισμένες πηγές απαιτούν καινοτόμες λύσεις.

Η Ομοσπονδιακή Μάθηση (Federated Learning, FL), μια βασική έννοια στο υπό συζήτηση πλαίσιο, επαναπροσδιορίζει την προσέγγιση της μηχανικής μάθησης. Στην FL, αντί όλα τα δεδομένα να συγκεντρώνονται σε έναν μόνο κεντρικό διακομιστή (server), η μάθηση συμβαίνει τοπικά σε μεμονωμένες συσκευές ή κόμβους. Αυτοί οι κόμβοι εκπαιδεύουν συνεργατικά ένα κοινό παγκόσμιο (global) μοντέλο χωρίς να απαιτείται να ανταλλάσσουν μεταξύ τους όλα τα δεδομένα. Αυτή η προσέγγιση διασφαλίζει την απαιτούμενη ιδιωτικότητα και ασφάλεια, ειδικά σε σενάρια όπου εμπλέκονται ευαίσθητες πληροφορίες καθώς επιτρέπει σε μεμονωμένες συσκευές, όπως πλοία ή αισθητήρες σε θαλάσσια περιβάλλοντα, να συμμετέχουν στη διαδικασία μάθησης χωρίς να αποκαλύπτουν τα ευαίσθητα δεδομένα τους.

Η FL βρίσκει εφαρμογή σε διάφορους τομείς της ναυτιλίας. Στις θαλάσσιες μεταφορές, διευκολύνει τα συστήματα διαχείρισης της ενεργειακής απόδοσης που βασίζονται σε δεδομένα, βελτιστοποιώντας τις παραμέτρους των πλοίων για εξοικονόμηση ενέργειας. Οι συνεργατικές αυτές τεχνικές ελέγχου της FL ενισχύουν την αποδοτικότητα και την ασφάλεια στη ναυσιπλοΐα. Η ανίχνευση σφαλμάτων και η προγνωστική συντήρηση στα πλοία διευκολύνονται από την FL, καθώς επιτρέπει σε πολλά σκάφη να συμμετέχουν στην εκπαίδευση μειώνοντας έτσι και τον συνολικό απαιτούμενο χρόνο. Τέλος, σε υποβρύχιες εφαρμογές, η FL βελτιστοποιεί τις επικοινωνίες και τον σχεδιασμό της διαδρομής των αυτόνομων υποβρύχιων οχημάτων (Autonomous underwater vehicles, AUVs).

8.2 Πειραματικά αποτελέσματα

Στην ενότητα αυτή παρουσιάζεται μια πρακτική εφαρμογή Ομοσπονδιακής Μάθησης (FL) με αξιοποίηση ενός συνόλου δεδομένων που παρέχεται από μια εμπορική ναυτιλιακή επιχείρηση. Ο στόχος των πειραματικών αποτελεσμάτων που παρουσιάζονται είναι η εξακρίβωση του βαθμού στον οποίο ένα αποκεντρωμένο σύστημα βασισμένο σε ένα μοντέλο FL είναι ανώτερο σε σχέση με προσεγγίσεις κεντρικών συστημάτων βασισμένα σε μάθηση μεταφοράς (Transfer Learning, TL) και σε ένα σύστημα συνδυαστικής μάθησης (Ensemble Learning). Πιο συγκεκριμένα, εξετάζεται ένα σύστημα παλινδρόμησης πολλών μεταβλητών για τη μοντελοποίηση και την πρόβλεψη της κατανάλωσης καυσίμου κατά τη μετακίνηση μεγάλων φορτηγών πλοίων. Δεδομένου ότι η κατανάλωση καυσίμου είναι ανάλογη της κατανάλωσης της κύριας ισχύος του κινητήρα (Main Engine Power, MEP) του πλοίου, η τελευταία χρησιμοποιήθηκε ως μεταβλητή-στόχος. Οι προβλέψεις για την κατανάλωση της MEP έχουν πολύ μεγάλη σημασία στον τομέα των θαλάσσιων μεταφορών, καθώς η ακριβής πρόβλεψη της κατανάλωσης καυσίμου μπορεί να επιτρέψει την υιοθέτηση μιας πολιτικής πρόληψης και έγκαιρων διορθωτικών ενεργειών για τη μείωση τόσο του λειτουργικού κόστους του πλοίου όσο και των περιβαλλοντικών επιπτώσεων. Ένα μοντέλο βαθιάς μάθησης με πλήρως διασυνδεδεμένα τεχνητά νευρωνικά δίκτυα (ANNs) επιλέχθηκε για την εκτίμηση της MEP βάσει πληθώρας χαρακτηριστικών-μεταβλητών που σχετίζονται με το πλοίο και τον καιρό. Καθότι η σχέση μεταξύ των μεταβλητών αυτών και της μεταβλητής-στόχου δεν είναι γνωστή εκ των προτέρων, τα ANNs επιλέχθηκαν ως μοντέλα παλινδρόμησης λόγω της ικανότητάς τους να προβλέπουν με ακρίβεια τόσο γραμμικές όσο και μη γραμμικές σύνθετες σχέσεις.[42] Οι ακόλουθες συνθήκες προσομοίωσης εφαρμόστηκαν σε γλώσσα Python 3.8 με χρήση των βιβλιοθηκών Sci-kit learn και Tensorflow (έκδοση 2.4). Η εκπαίδευση όλων των αλγορίθμων διεκπεραιώθηκε σε έναν υπολογιστή (CPU i7-8700, 3,2 GHz, RAM 8 GB, χωρίς χρήση GPU).

8.2.1 Περιγραφή συνόλου δεδομένων

Η ακατέργαστη μορφή του συνόλου δεδομένων περιείχε πληροφορίες για δύο πανομοιότυπα φορτηγά πλοία που παρέχονται από μια ναυτιλιακή επιχείρηση. Δέκα χρονοσειρές για κάθε πλοίο καταγράφονται με περίοδο δειγματοληψίας 1 λεπτό και αφορούν διάφορα ταξίδια των πλοίων ενώ συμπεριλαμβάνουν δεδομένα που σχετίζονται τόσο με το πλοίο όσο και με τον καιρό. Οι καταγραφές είχαν διάρκεια 1 έτους (από 01-07-2021 έως 30-06-2022) και είχαν σαν αποτέλεσμα 514.525 και 525.005 συνολικά δείγματα ανά χρονοσειρά για τα πλοία 1 και 2, αντίστοιχα. Πιο συγκεκριμένα, για κάθε φορτηγό πλοίο είναι διαθέσιμες οι ακόλουθες μεταβλητές:

1. Ταχύτητα πάνω από το έδαφος (Speed Over Ground, SOG): είναι η ταχύτητα του σκάφους σε σχέση με την επιφάνεια της γης, μετρούμενη σε κόμβους.

2. Ταχύτητα μέσω νερού (Speed Through Water, STW): είναι η ταχύτητα του σκάφους σε σχέση με τα ρεύματα νερού, μετρούμενη σε κόμβους.
3. Κατεύθυνση (Heading): είναι η κατεύθυνση προς την οποία δείχνει το πλοίο και εκφράζεται ως η γωνιακή απόσταση σε σχέση με τον Βορρά (0°), δεξιόστροφα έως 359° .
4. Συνεχής ταχύτητα ανέμου (Continuous Wind Speed, CWS): είναι το πλάτος της ταχύτητας του ανέμου που εκφράζεται σε μονάδες m/s.
5. Διακριτοποιημένη ταχύτητα ανέμου (Discretized Wind Speed, DWS): είναι η κβαντισμένη εκδοχή της CWS, εκπεφρασμένη σε κλίμακα Beaufort (bft), και παίρνει τις κατηγορικές τιμές 0-12 που αντιστοιχούν στους 13 τύπους ανέμων (από ασθενείς σε κατηγορίες ισχυρών ανέμων).
6. Κατεύθυνση ανέμου (Wind Direction, WD): είναι η κατεύθυνση του ανέμου σε σχέση με την κατεύθυνση του πλοίου. Λαμβάνει τιμές από 0° - 360° (άνεμος στην πλώρη στους 0° , άνεμος κάθετα στο πλοίο στους 90° , άνεμος στην πρύμνη στους 180°).
7. Βύθισμα προς τα εμπρός (Draft Forward, DF): το βύθισμα προς τα εμπρός (πλώρη) μετριέται (βάθος σε μέτρα) κάθετα στα προκαθορισμένα επίπεδα βάθους της πλώρης.
8. Βύθισμα προς τα πίσω (Draft aft, DA): το βύθισμα προς τα πίσω (πρύμνη) μετριέται (βάθος σε μέτρα) κάθετα στα προκαθορισμένα επίπεδα βάθους της πρύμνης.
9. Διαγωγή πλοίου (Trim): Είναι η διαφορά μεταξύ του DF και του DA (βάθος σε μέτρα), σε σχέση με τη σχεδιασμένη ίσαλο γραμμή που βρίσκεται στη μέση του πλοίου. Καθορίζει το ελάχιστο βάθος νερού στο οποίο μπορεί να πλέει με ασφάλεια ένα πλοίο.
10. Ισχύς κύριου κινητήρα (MEP): είναι η συνολική ισχύς που παρέχεται από τον κύριο κινητήρα που είναι εγκατεστημένος σε πλοίο για την πρόωσή του και εκφράζεται σε kW. Προφανώς, συσχετίζεται θετικά με τις απαιτήσεις ηλεκτρικής ενέργειας του πετρελαιοκινητήρα του πλοίου, με αποτέλεσμα να είναι ανάλογη με την κατανάλωση καυσίμου.

8.2.2 Προεπεξεργασία και Μείωση Διαστάσεων

Εδώ, παρατηρούμε τα βήματα προεπεξεργασίας που πραγματοποιήθηκαν για την αύξηση της αναλογίας σήματος προς θόρυβο των ακατέργαστων δεδομένων. Δεδομένης της χαμηλής περιόδου δειγματοληψίας (1 λεπτό), πολλά διαδοχικά δείγματα των χρονοσειρών παρουσιάζουν μηδενικές ή αργές διαφορές. Αυτό επιτρέπει την άμεση αντικατάσταση τιμών που λείπουν με παρεμβολή. Για κάθε πλοίο και κάθε μεταβλητή ξεχωριστά, οι τιμές που λείπουν εντοπίστηκαν και αντικαταστάθηκαν από τον μέσο όρο των αμέσως προηγούμενων και των αμέσως επόμενων διαθέσιμων δειγμάτων. Οι ελλείπουσες τιμές δεν υπερέβαιναν το 8% του συνολικού μεγέθους του δείγματος για όλες τις μεταβλητές. Προκειμένου να

διατηρηθούν μόνο τα δείγματα με το σκάφος εν κινήσει, τα δείγματα που αντιστοιχούν σε στιγμές ακινησίας των πλοίων ($SOG < 0,2$ κόμβοι) απορρίφθηκαν. Στη συνέχεια, εφαρμόστηκε η μέθοδος μείωσης ρυθμού δειγματοληψίας (downsampling) με συντελεστή 15 (1 δείγμα ανά 15 λεπτά) σε ολόκληρο το σύνολο δεδομένων με στόχο τον περιορισμό της πλεοναστικότητας των δεδομένων αλλά και τη μείωση του χρόνου εκτέλεσης του μοντέλου. Πριν από τη μείωση αυτή, εφαρμόστηκε η μέθοδος κινητού μέσου όρου για την εξομάλυνση των καμπυλών και τον περιορισμό της παραμόρφωσης λόγω του φαινομένου aliasing. Έπειτα πραγματοποιήθηκε οπτικός έλεγχος των κατανομών των παραμέτρων για τον εντοπισμό ακραίων τιμών (δείγματα που διαφέρουν από τον μέσο όρο κατά 3 τυπικές αποκλίσεις). Τέλος, όλα τα χαρακτηριστικά και οι μεταβλητές-στόχοι κανονικοποιήθηκαν ώστε να λάβουν τιμές στο διάστημα $[0,1]$, σύμφωνα με τον ακόλουθο τύπο:

$$x'_{i,j} = \frac{x_{i,j} - \min\{X_j\}}{\max\{X_j\} - \min\{X_j\}},$$

$$\forall i \in [1, 2, \dots, |X_j|], j \in [1, 2, \dots, 10]$$

όπου $x'_{i,j}$ και $x_{i,j}$ υποδηλώνουν το δείγμα i κλίμακας (scaled) και μη (non-scaled) αντίστοιχα της μεταβλητής j και X_j είναι το διάνυσμα-στήλη της μεταβλητής j . Ο τελεστής $|\cdot|$ χρησιμεύει ως μετρητής των διανυσμάτων. Ως μία πρώτη προσπάθεια μείωσης της πολυπλοκότητας του μοντέλου παλινδρόμησης και αποφυγής πλεοναστικότητας χαρακτηριστικών, εφαρμόστηκαν κριτήρια για τη διατήρηση ή και τη μείωση των διαστάσεων του μοντέλου (αντιστοιχεί στον αριθμό νευρώνων εισόδου). Πρώτον, η μεταβλητή Heading αγνοήθηκε ως παράγοντας εντελώς άσχετος με την MEP του πλοίου. Οι μεταβλητές DF και DA επίσης αγνοήθηκαν, καθώς αντικατοπτρίζονται άμεσα στη μεταβλητή Trim. Τέλος, μόνο η μεταβλητή DWS λήφθηκε υπόψη σχετικά με το πλάτος της ταχύτητας του ανέμου. Αξίζει να σημειωθεί ότι τόσο η μεταβλητή SOG όσο και η STW λήφθηκαν υπόψη, δεδομένου ότι η STW περιέχει έμμεσα πληροφορίες σχετικά με τα θαλάσσια ρεύματα και ροές και είναι πιθανό να διαφέρει από την SOG (π.χ. όταν το σκάφος πλέει σε ρεύμα 5 κόμβων, τότε $SOG=5$ kn και $STW=0$). Το παραπάνω επιβεβαιώθηκε επίσης με τον υπολογισμό του συντελεστή συσχέτισης μεταξύ SOG και STW (Pearson's $R_{coef} = 0.23$). Περαιτέρω εξάλειψη του αριθμού των μεταβλητών που συμπεριλαμβάνονται στα MEP μοντέλα παρουσιάζεται παρακάτω.

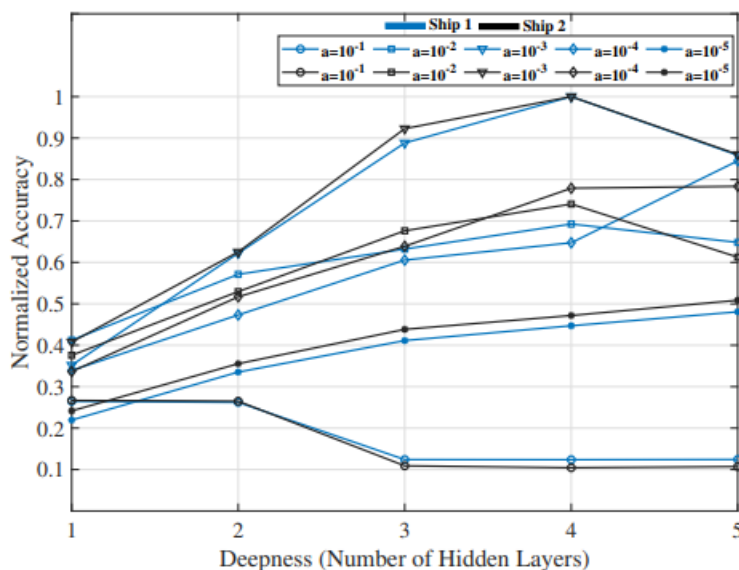
8.2.3 Ρύθμιση υπερπαραμέτρων μάθησης

Κατά την εκπαίδευση ενός μοντέλου ANN πολλών επιπέδων, η επιλογή των υπερπαραμέτρων πρέπει να γίνει προσεκτικά, καθώς οι τιμές τους ενδέχεται να οδηγήσουν σε εσφαλμένη πρόβλεψη. Μεταξύ διαφόρων υπερπαραμέτρων μάθησης, διαφορετικές τιμές του ρυθμού

μάθησης α (ο οποίος σχετίζεται με τη μέθοδο οπισθοδρόμησης και την προσαρμογή των συντελεστών βαρύτητας του ANN) και του βάθους του δικτύου (δηλαδή του αριθμός των κρυμμένων επιπέδων ο οποίος επηρεάζει την πολυπλοκότητα του μοντέλου) παρατηρήθηκε ότι προκαλούν σημαντικές διακυμάνσεις στην απόδοση του μοντέλου. Για το λόγο αυτό, πραγματοποιήθηκαν εκτεταμένοι πειραματισμοί για τη ρύθμιση αυτών των παραμέτρων. Ως συνάρτηση ενεργοποίησης στους νευρώνες του δικτύου χρησιμοποιήθηκε η συνάρτηση ReLu, ενώ για την εφαρμογή της στοχαστικής μεθόδου Gradient Descent στη μέθοδο οπισθοδρόμησης εφαρμόστηκε ο αλγόριθμος βελτιστοποίησης Adam. Αρχικά, το σύνολο δεδομένων χωρίστηκε σε δεδομένα εκπαίδευσης και δεδομένα ελέγχου με ποσοστιαία αναλογία 90% και 10% αντίστοιχα. Το μέσο απόλυτο σφάλμα (Mean Absolute Error, MAE) εξήχθη από τα δεδομένα ελέγχου. Η κανονικοποιημένη ακρίβεια (Normalized Accuracy) A_{norm} κάθε παραγόμενου μοντέλου ήταν αντιστρόφως ανάλογη με το MAE και συσχετίστηκε με το μοντέλο ελαχίστου MAE ($A_{norm} = 1$ για το καλύτερα ρυθμισμένο μοντέλο), σύμφωνα με τον τύπο:

$$A_{norm}^i = \frac{\min_{j \in M} \{ \sum_{n=1}^{N_{val}} |y_{pred}^{n,j} - y_{real}^{n,j}| \}}{\sum_{n=1}^{N_{val}} |y_{pred}^{n,i} - y_{real}^{n,i}|},$$

όπου A_{norm}^i είναι η κανονικοποιημένη ακρίβεια του μοντέλου i , $y_{pred}^{n,i}$ και $y_{real}^{n,i}$ είναι οι προβλεφθείσες και οι πραγματικές τιμές αντίστοιχα, του δείγματος ελέγχου n και N_{val} ο μετρητής των δεδομένων ελέγχου. Οι δείκτες i και j διατρέχουν το σύνολο M των μοντέλων που προκύπτουν για διάφορες τιμές των υπερπαραμέτρων (ο αριθμός των μοντέλων είναι $5 \cdot 5 = 25$ μοντέλα). Έτσι, η ακρίβεια όλων των μοντέλων συσχετίζεται με την απόδοση του βέλτιστου μοντέλου, δεδομένου ότι ο παρονομαστής αντικατοπτρίζει το MAE που προκύπτει από το μοντέλο i , ενώ ο αριθμητής το ελάχιστο MAE μεταξύ όλων των μοντέλων. Στο Σχήμα 22 παρουσιάζεται η κανονικοποιημένη ακρίβεια για 25 προεκπαιδευμένα μοντέλα για πέντε διαφορετικές τιμές του α και του βάθους του δικτύου για τα πλοία 1 και 2. Όπως φαίνεται, τα τοπικά μοντέλα ANN και για τα δύο πλοία είναι βέλτιστα διαμορφωμένα για $\alpha = 0.001$ και 4 κρυμμένα επίπεδα. Οι νευρώνες που περιέχονται σε κάθε κρυφό επίπεδο ήταν αντίστοιχα [1000, 800, 600, 400].



Σχήμα 22: Απόδοση του μοντέλου συναρτήσει του βάθους του νευρωνικού δικτύου για τις διάφορες τιμές του ρυθμού μάθησης α

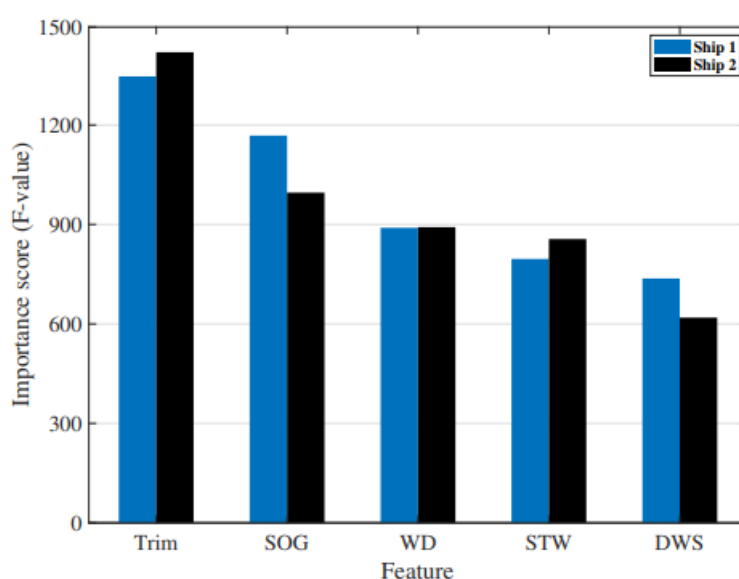
8.2.4 Αναδρομική εξάλειψη χαρακτηριστικών

Κατά την προσπάθεια περαιτέρω μείωσης της πολυπλοκότητας του μοντέλου, υιοθετήθηκε μια μέθοδος αναδρομικής εξάλειψης χαρακτηριστικών (Recursive Feature Elimination, RFE). Για την ταξινόμηση των 5 χαρακτηριστικών (SOG, SWT, DWS, WD και Trim) με βάση τη σημασία τους στην πρόβλεψη της MEP, η RFE περιελάμβανε έναν έλεγχο σημαντικότητας χαρακτηριστικών βασισμένο στη μέθοδο XGBoost η οποία εφαρμόστηκε στα δεδομένα εκπαίδευσης για κάθε μοντέλο πλοίου ξεχωριστά. Για την ταξινόμηση αυτή προέκυψαν οι τιμές-F (F-values) από στατιστικούς ελέγχους συσχέτισης μεταξύ μεταβλητών-χαρακτηριστικών και μεταβλητής-στόχου). Η μέθοδος XGBoost προτιμήθηκε έναντι άλλων μεθόδων που σχετίζονται με τη σημαντικότητα των χαρακτηριστικών λόγω του μικρού χρόνου εκτέλεσής της και της υπεροχής της έναντι των γραμμικών μοντέλων. Και τα δύο μοντέλα κατέληξαν στην ίδια διάταξη των χαρακτηριστικών [Trim, SOG, WD, STW, DWS], σε φθίνουσα σειρά σημαντικότητας. Όπως αναμενόταν, η μεταβλητή Trim (η οποία εξαρτάται από το φορτίο) και η μεταβλητή SOG (η οποία εξαρτάται από την κατανάλωση καυσίμου) έχουν υψηλό βαθμό συσχέτισης με την MEP, με το WD να είναι το αμέσως επόμενο χαρακτηριστικό βάσει σημαντικότητας (Σχήμα 23). Έχοντας τη λίστα χαρακτηριστικών ταξινομημένη, εφαρμόστηκαν αναδρομικά τα ακόλουθα βήματα:

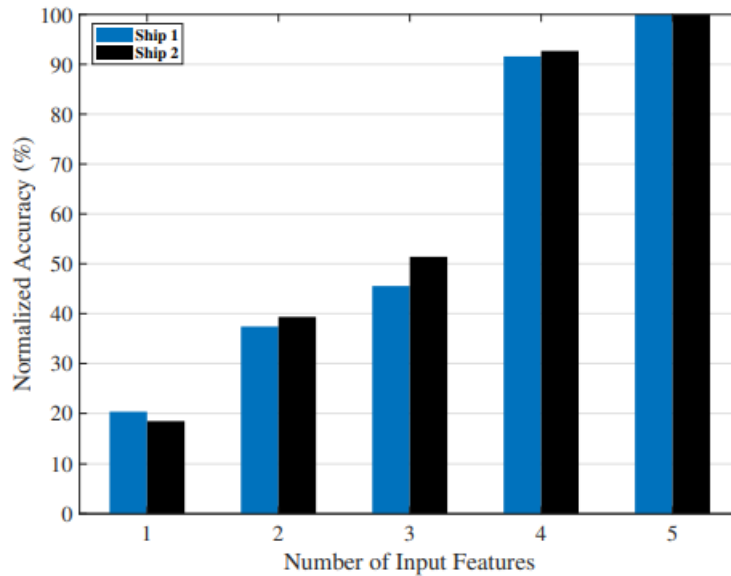
1. Ρύθμιση των δύο μοντέλων πλοίων στις βέλτιστες υπερπαραμέτρους.
2. Ορισμός του αριθμού των χαρακτηριστικών $K = 5$.
3. Επιλογή των K σημαντικότερων χαρακτηριστικών ως εισόδους του μοντέλου.
4. Εκπαίδευση μοντέλου και εξαγωγή του MAE ελέγχου.

5. Μείωση του αριθμού των χαρακτηριστικών $K = K - 1$.
6. Επανάληψη βημάτων 3-5 μέχρις ότου $K = 1$.
7. Συσχέτιση της ακρίβεια όλων των μοντέλων με το βέλτιστο.

Στο Σχήμα 24 απεικονίζονται σε γράφημα ράβδων οι κανονικοποιημένες επί τοις εκατό ακρίβειες για διάφορες τιμές του πλήθους χαρακτηριστικών-εισόδων. Η μέτρηση έγινε επί τοις εκατό για να καταδειχθεί η ποσοστιαία μείωση που επιφέρει η εξάλειψη χαρακτηριστικών στην ακρίβεια ($A_{norm}(\%)$). Το παραπάνω εξηγείται από το γεγονός ότι κάθε χαρακτηριστικό που σχετίζεται, περισσότερο ή λιγότερο, με την MEP μπορεί να βελτιώσει την απόδοση του μοντέλου με κόστος, βέβαια την αύξηση της πολυπλοκότητας του μοντέλου. Δεδομένου ότι η επιλογή 4 χαρακτηριστικών δεν μειώνει σημαντικά την ακρίβεια του μοντέλου (μικρότερη από 10% η μείωση της ακρίβεια σε σχέση με την επιλογή και των 5 χαρακτηριστικών), τελικά επιλέχθηκε η εξάλειψη ενός χαρακτηριστικού και η τελική επιλογή των 4 σημαντικότερων (Trim, SOG, WD, STW), έτσι ώστε να μειωθεί και η πολυπλοκότητα του μοντέλου.



Σχήμα 23: Ταξινομημένη F-Value των χαρακτηριστικών όπως προέκυψε από τον αλγόριθμο XGBoost



Σχήμα 24: Απόδοση μοντέλου για διαφορετικό πλήθος μεταβλητών εισόδου

8.2.5 Σύγκριση απόδοσης-πολυπλοκότητας

Μετά τον καθορισμό των τιμών των υπερπαραμέτρων και τη διαδικασία εξάλειψης των μεταβλητών-χαρακτηριστικών με στόχο τη διατήρηση ισορροπίας μεταξύ απόδοσης και πολυπλοκότητας, ακολουθεί μία σύγκριση των προβλέψεων διαφόρων μοντέλων για τη μεταβλητή MER. Για να διερευνηθούν ποσοτικά τα πλεονεκτήματα και τα μειονεκτήματα διαφορετικών κεντρικών και αποκεντρωμένων συστημάτων, εξετάζονται και αντιπαραβάλλονται οι εξής τέσσερις διαφορετικές μέθοδοι πρόβλεψης ως προς την ακρίβεια και την πολυπλοκότητά τους:

- Κεντρική μάθηση και κλήση μοντέλου (Centralized Learning and Inference, CLI): Αποτελεί μία υπολογιστικά ακριβή μέθοδο, στην οποία όλα τα δείγματα δεδομένων (και από τα δύο πλοία) συγκεντρώνονται στον κεντρικό διακομιστή (server) και συγκροτούν ένα μεγάλο και ισχυρό μοντέλο. Η διαδικασία της κλήσης του μοντέλου γίνεται με την κάθετη επικοινωνία μεταξύ των τοπικών πελατών (clients) πλοίων και του κεντρικού server. Το κεντρικό μοντέλο δεν περιλαμβάνει τα πλοία και έτσι απαιτεί συνεχή ανταλλαγή δεδομένων με τους τοπικούς clients.
- Τοπική μάθηση και κοινή χρήση μοντέλων (Local Learning and Model Sharing, LLMS): Βασίζεται στις αρχές της μάθησης μεταφοράς (TL) κατά την οποία ένα μοντέλο εκπαιδεύεται τοπικά για κάποιο πλοίο και στη συνέχεια τα υπόλοιπα πλοία «κληρονομούν» το προεκπαιδευμένο μοντέλο για μελλοντική κλήση του. Δοκιμάστηκε τόσο η περίπτωση κατά την οποία το πλοίο 1 «κληρονομεί» από το πλοίο 2 ($LLMS_{1,2}$) όσο και η αντίστροφη ($LLMS_{2,1}$). Ο κεντρικός server χρησιμοποιείται μόνο για κοινή χρήση μοντέλων μεταξύ όλων των τοπικών clients-πλοίων.

- Τοπική μάθηση και συλλογική κλήση του μοντέλου (Local Learning and Collaborative Inference, LLCI): Παρόμοια με την τεχνική μεθόδων Ensemble Learning, σε αυτή τη μέθοδο εκπαιδεύονται 2 τοπικά μοντέλα και στη συνέχεια, κατά τη διαδικασία της κλήσης του μοντέλου, κάθε πλοίο λαμβάνει τη μέση πρόβλεψη όλων των μοντέλων. Για να λειτουργήσει αυτή η μέθοδος, κάθε τοπικός client πρέπει να αποθηκεύει όλα τα διαθέσιμα μοντέλα τοπικά για την εξαγωγή της μέσης πρόβλεψης ώστε να αποφευχθεί επιπρόσθετη επιβάρυνση που θα προκαλούσε η μεταξύ τους επικοινωνία.
- Federated Learning (FL): αυτή η μέθοδος ακολουθεί τον αλγόριθμο FedAvg.[51] Εν συντομία, κάθε client εκπαιδεύει το τοπικό μοντέλο αποκεντρωμένα και μόνο οι παράμετροι του μοντέλου αποστέλλονται σε ένα κέντρο. Εκεί, συνδυάζονται όλες οι παράμετροι για να συγκροτήσουν το global μοντέλο, πριν το επιστρέψει σε όλους τους τοπικούς clients. Ο τύπος που χρησιμοποιείται για τη διαδικασία της συγκρότησης (aggregation) του μοντέλου είναι ο ακόλουθος:

$$W_{Global}^{t+1} = \sum_{m=1}^{N_m} \frac{n_m}{n_{total}} \cdot W_{Local,m}^t$$

όπου W_{Global}^{t+1} είναι οι παράμετροι του global μοντέλου κατά τον $t + 1$ γύρο της διαδικασίας της συγκρότησης και $W_{Local,m}^t$ οι παράμετροι του τοπικού μοντέλου m κατά τον t γύρο, n_m είναι ο μετρητής των δειγμάτων του τοπικού μοντέλου m και n_{total} ο μετρητής των δειγμάτων όλων των τοπικών Federated Learning μοντέλων.

Η ακρίβεια A_i για κάθε μέθοδο i υπολογίζεται με βάση τα μέχρι πρότινος άγνωστα στο μοντέλο δεδομένα ελέγχου, ενώ η πολυπλοκότητα σε ποσοτικοποιημένη μορφή προκύπτει σύμφωνα με τον ακόλουθο τύπο:

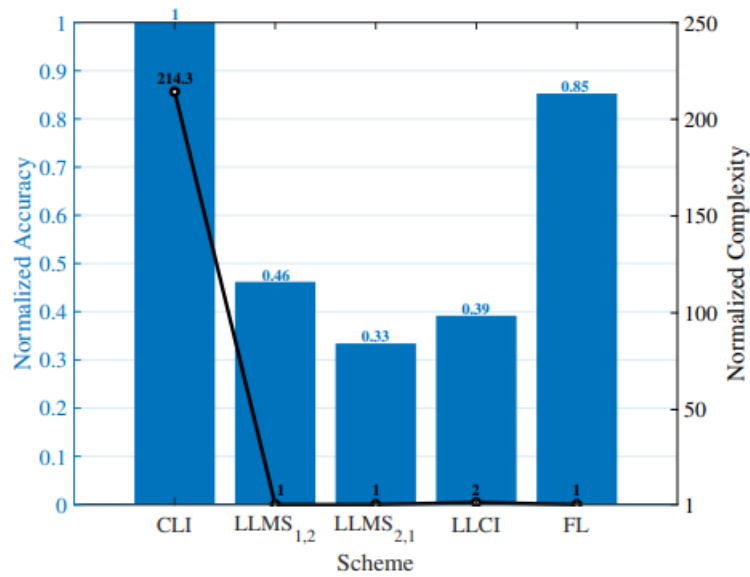
$$C_i = \sum_{m=1}^{N_m} (E_m \cdot U_m + S_m),$$

όπου η συνολική πολυπλοκότητα C_i της μεθόδου i , το οποίο περιλαμβάνει N_m μοντέλα ANNs ισούται με το άθροισμα της πολυπλοκότητας κάθε μοντέλου m η οποία προκύπτει από το άθροισμα του γινομένου των κρυμμένων ακμών E_m και των κρυμμένων νευρώνων U_m με τις απαιτήσεις S_m του νευρωνικού δικτύου σε KB.

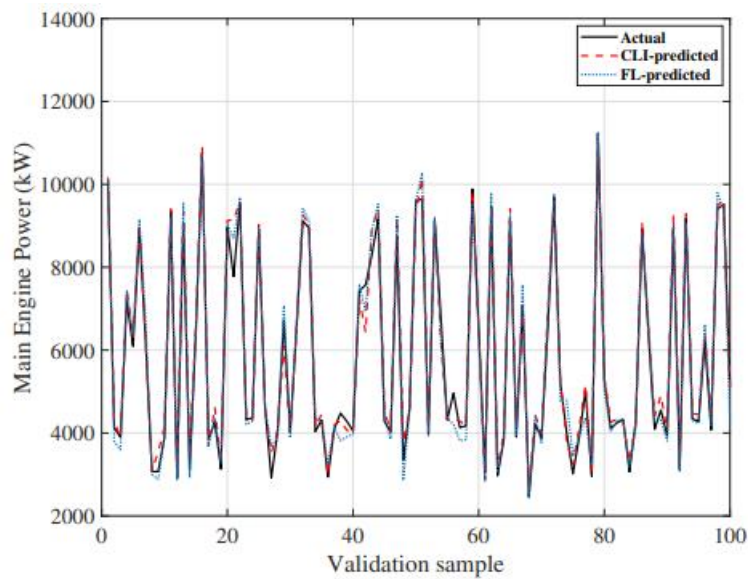
Για τον έλεγχο των μεθόδων LLMS, LLCI και FL, υιοθετήθηκε η βέλτιστη διαμόρφωση υπερπαραμέτρων και επιλογή χαρακτηριστικών για τα τοπικά μοντέλα πλοίων, όπως παρουσιάστηκε στα κεφάλαια 8.2.3 και 8.2.4, ενώ για τη μέθοδο CLI η ελάχιστη τιμή του μέσου σφάλματος στην πρόβλεψη της MEP εμφανίζεται για $\alpha = 0.001$ και για 5 κρυμμένα επίπεδα όπως προέκυψε μετά από εκτεταμένες προσομοιώσεις εκπαίδευσης.

Στο Σχήμα 25 παρουσιάζεται η απόδοση (όσον αφορά την ακρίβεια) και η πολυπλοκότητα των μεθόδων. Τα δεδομένα ελέγχου που χρησιμοποιήθηκαν κατά την αξιολόγηση ακρίβειας

και πολυπλοκότητας όλων των μεθόδων αφορούσαν δείγματα δεδομένων τόσο από το πλοίο 1 όσο και από το πλοίο 2 τα οποία δε χρησιμοποιήθηκαν κατά τη διάρκεια της εκπαίδευσης. Για την κατάδειξη των συγκριτικών διαφορών ανάμεσα στις διάφορες μεθόδους σε ό,τι αφορά την απόδοση και την πολυπλοκότητά τους, χρησιμοποιήθηκαν οι κανονικοποιημένες μορφές τους. Πιο συγκεκριμένα, η ακρίβεια κάθε μεθόδου ($1 / MAE_i$) διαιρέθηκε με τη μέγιστη ακρίβεια που παρατηρήθηκε ($1 / MAE_{min}$) (η ακρίβεια της CLI), ενώ η πολυπλοκότητα κάθε μεθόδου διαιρέθηκε με την ελάχιστη πολυπλοκότητα που παρατηρήθηκε (η πολυπλοκότητα της LLMS ή της FL). Όπως φαίνεται στο Σχήμα 25, η μέθοδος CLI παρουσιάζει τη μεγαλύτερη ακρίβεια πρόβλεψης, με τη μέθοδο FL να παρουσιάζει ακρίβεια μειωμένη κατά 15% σε σχέση με την CLI. Επιπλέον, στην περίπτωση που το μοντέλο του πρώτου πλοίου «κληρονομεί» από το μοντέλο του δεύτερου πλοίου ($LLMS_{1,2}$), παρατηρείται σημαντική μείωση της ακρίβειας, κάτι που οφείλεται στην υψηλότερη προσαρμοστικότητα του μοντέλου του δεύτερου πλοίου στα δικά του δεδομένα. Το ίδιο ισχύει και για το αντίστροφο σχήμα ($LLMS_{2,1}$) το οποίο παρουσιάζει ακόμη μικρότερη ακρίβεια από το $LLMS_{1,2}$ (μείωση 54% και 67% σε σχέση με την ακρίβεια της CLI). Η εξήγηση για την αναποτελεσματικότητα αυτών των μεθόδων έγκειται στο γεγονός ότι δε διαθέτουν global πληροφορίες και για τα δύο πλοία, σε αντίθεση με τη μέθοδο CLI που εκπαιδεύτηκε στο σύνολο δεδομένων αμφοτέρων των πλοίων. Παρατηρείται επίσης ότι η μέθοδος $LLCI$ έχει ακρίβεια μεταξύ των ακριβειών των $LLMS_{1,2}$ και $LLMS_{2,1}$ καθώς ο μέσος όρος των προβλέψεων μειώνει το σφάλμα που εισάγεται με αξιοποίηση μόνο του πρώτου μοντέλου και αυξάνει το σφάλμα που εισάγεται με αξιοποίηση μόνο του δεύτερου μοντέλου. Είναι σημαντικό να αναφερθεί ότι η ακρίβεια των μεθόδων LLMS και LLCI θα βελτιωνόταν με κόστος μεγαλύτερη διάρκεια στη διαδικασία της εκπαίδευσης και γενικότερα περισσότερο χρονοβόρες διαδικασίες από τη μέθοδο FL. Για καλύτερη κατανόηση της μη κανονικοποιημένης ακρίβειας στην πρόβλεψη της MEP, στο σχήμα 26 παρουσιάζονται οι πραγματικές και οι προβλεφθείσες καμπύλες της MEP για τα δείγματα ελέγχου για τις δύο καλύτερες (ως προς την ακρίβεια πρόβλεψης) μεθόδους (CLI και FL). Όσον αφορά την πολυπλοκότητα, παρατηρήθηκε ότι η CLI είναι 214,3 φορές πολυπλοκότερη από τις μεθόδους LLMS και FL, δεδομένης της υψηλής πυκνότητας και του μεγάλου μεγέθους του ANN ($153.6 \cdot 10^{12}$ ακμές, 3.005 νευρώνες και 36,9 MB) για την επίτευξη βέλτιστης ακρίβειας. Αντίθετα, οι μέθοδοι LLMS και FL παρουσιάζουν τη χαμηλότερη πολυπλοκότητα επειδή απαιτούν ANNs μειωμένων διαστάσεων ($768 \cdot 10^9$ ακμές, 2.805 νευρώνες και 35,4 MB). Τέλος, αξίζει να σημειωθεί ότι η μέθοδος LLCI έχει διπλάσια πολυπλοκότητα από τις μεθόδους LLMS και FL καθώς περιλαμβάνει 2 ANNs για κάθε πλοίο για την εξαγωγή της μέσης πρόβλεψης ενώ στις LLMS και FL απαιτείται 1 ANN ανά πλοίο.



Σχήμα 25: Κανονικοποιημένη ακρίβεια (αριστερός κάθετος άξονας) και πολυπλοκότητα (δεξιός κάθετος άξονας) για τις διάφορες μεθόδους



Σχήμα 26: Πραγματική (actual) τιμή και προβλεφθείσες σύμφωνα με την CLI και FL μέθοδο της μεταβλητής MEP

8.2.6 Ζητήματα σχετικά με την εκπαίδευση και την ανάπτυξη των συστημάτων

Όπως φάνηκε από τις προηγούμενες αναλύσεις σχετικά με την ισορροπία μεταξύ απόδοσης και πολυπλοκότητας του μοντέλου, η μέθοδος CLI παρουσιάζει την υψηλότερη ακρίβεια έχοντας όμως και τη μεγαλύτερη πολυπλοκότητα, ενώ η μέθοδος FL παρουσιάζει ελαφρώς χαμηλότερη ακρίβεια συνδυάζοντάς την όμως με σημαντικά χαμηλότερη πολυπλοκότητα.

Έτσι γίνεται σαφές ότι για ένα σύστημα που η ακρίβεια είναι το ζητούμενο ανεξαρτήτως του κόστους, η CLI αποτελεί την καταλληλότερη επιλογή, ενώ για ένα σύστημα που λαμβάνει υπόψη και την πολυπλοκότητα η FL θα ήταν σαφώς καλύτερη επιλογή. Κατά την επιλογή της καταλληλότερης μεθόδου θα πρέπει να ληφθούν υπόψη και τα ζητήματα σχετικά με την εκπαίδευση και ανάπτυξη των συστημάτων. Αρχικά, για να καταστεί εφικτή η εκπαίδευση στη μέθοδο CLI, ο κεντρικός server λαμβάνει απευθείας τα δεδομένα και από τα δύο πλοία μέσω πρωτοκόλλου επικοινωνίας ανερχόμενης ζεύξης (uplink communication). Μετά την ολοκλήρωση της εκπαίδευσης, το κεντρικό μοντέλο μεταφέρεται και στα δύο πλοία για μελλοντική κλήση του μοντέλου. Όσον αφορά την εκπαίδευση στη μέθοδο FL, οι αλλαγές στο μοντέλο για κάθε πλοίο ξεχωριστά εκτελούνται τοπικά και μόνο οι παράμετροι του μοντέλου αποστέλλονται στον server για την εφαρμογή του αλγορίθμου FedAvg, με αποτέλεσμα να διαφυλάσσεται το απόρρητο των δεδομένων. Έτσι, και τα δύο πλοία έχουν τελικά διαθέσιμο το global μοντέλο για την εξαγωγή του μέσου όρου των προβλέψεων. Η εκπαίδευση των μοντέλων τοπικά γίνεται επίσης και στις μεθόδους LLMs και LLCL, με τον κεντρικό server να χρησιμοποιείται κατά τη μετάδοση μόνο των μοντέλων και όχι των δεδομένων. Κατά συνέπεια, είναι σαφής η υπεροχή της μεθόδου FL όταν λαμβάνονται ταυτόχρονα υπόψη τόσο η ακρίβεια και η πολυπλοκότητα των μεθόδων, όσο και η εφαρμοστικότητα τους με παράλληλη διαφύλαξη προσωπικών δεδομένων.

8.3 Ανοικτά ζητήματα

Από την ανάλυση των προηγούμενων παραγράφων, μπορούν να προκύψουν ορισμένα ανοικτά ζητήματα και μελλοντικές ερευνητικές προκλήσεις σχετικά με εφαρμογές της FL σε ναυτιλιακά περιβάλλοντα, τα οποία περιγράφονται σε αυτή την ενότητα.

8.3.1 Υλοποίηση με δεδομένα μεγάλης κλίμακας με 6G δίκτυα

Δεδομένου ότι η συζήτηση για τη μετάβαση σε μια νέα αποκεντρωμένη αρχιτεκτονική έχει ήδη ξεκινήσει, αναμένεται ότι τα θαλάσσια δίκτυα θα προσαρμοστούν στην εποχή του 6G, ενσωματώνοντας διάφορες σημαντικές τεχνολογίες, όπως η εικονικοποίηση (virtualization) της λειτουργίας των δικτύων και η δικτύωση βάσει λογισμικού (software-defined networking, SDN).[52] Ως εκ τούτου, οι προσεγγίσεις FL θα πρέπει να εφαρμόζονται και για τη διαχείριση μεγάλου όγκου ετερογενών δεδομένων.

8.3.2 Μεγάλες αποστάσεις διάδοσης και ποικίλα χαρακτηριστικά καναλιών

Όπως έχει αναφερθεί, ένα ναυτιλιακό περιβάλλον περιλαμβάνει μεγαλύτερες αποστάσεις διάδοσης και περισσότερο πολύπλοκους μηχανισμούς μετάδοσης σε σύγκριση με τα τυπικά ασύρματα κυψελωτά δίκτυα, κυρίως λόγω της σχετικής κίνησης της επιφάνειας της θάλασσας. Ως εκ τούτου, απαιτείται κατάλληλη προεπεξεργασία δεδομένων προκειμένου να μειωθούν οι επιπτώσεις των διαφορών στις αποστάσεις και τις συνθήκες διάδοσης. Επιπλέον, αυτόνομα οχήματα επιφανείας (autonomous surface vehicles, ASV) θα πρέπει να χρησιμοποιηθούν για τον μετριάσμό των επιπτώσεων των περιβαλλοντικών παραμέτρων στην απόδοση και σχετικά με την όποια καθυστέρηση διάδοσης.[53]

8.3.3 Μη πανομοιότυπα δεδομένα εκπαίδευσης προερχόμενα από διαφορετικές πηγές

Η διαχείριση μη πανομοιότυπων δεδομένων είναι ένα πολύ σημαντικό ζήτημα κατά την ενσωμάτωση δεδομένων από διαφορετικές πηγές. Σε αυτήν την περίπτωση, ενδέχεται να υπάρχουν ασυμβατότητες στα ληφθέντα δεδομένα, ειδικά σε περιβάλλοντα υψηλής συσχέτισης (highly correlated environments).

8.3.4 Αξιολόγηση απόδοσης πλήρους κλίμακας

Μέχρι τώρα, οι ερευνητικές δραστηριότητες σχετικά με την ενσωμάτωση της FL σε θαλάσσια δίκτυα επικοινωνιών εξετάζουν περιορισμένες τοπολογίες δικτύου (για παράδειγμα με δύο μόνο πλοία) και γενικά δίκτυα με περιορισμένο αριθμό συνιστωσών του δικτύου και IoT συσκευών. Ωστόσο, για την εξαγωγή αντιπροσωπευτικών συμπερασμάτων σχετικά με τις εφαρμογές της FL στα θαλάσσια δίκτυα επικοινωνιών, απαιτείται μια πιο σύνθετη ανάλυση που μπορεί να συμπεριλαμβάνει πολλαπλά, υψηλότερων διαστάσεων και ετερογενή δίκτυα και συνιστώσες. Στο ίδιο πλαίσιο, η ακριβής μοντελοποίηση καναλιών είναι επίσης ένα άλλο δύσκολο πεδίο, καθώς αναμένεται να υποστηρίζονται πολλαπλές ζώνες συχνοτήτων στα 6G δίκτυα με διαφορετικά χαρακτηριστικά διάδοσης. Τέλος, πρέπει επίσης να σημειωθεί σε αυτό το σημείο ότι μέχρι στιγμής δεν έχουν αναφερθεί σχετικές πειραματικές δραστηριότητες.

8.3.5 Κεντρικά έναντι κατανεμημένων συστημάτων Ομοσπονδιακής Μάθησης

Η υλοποίηση των διαφόρων ML μοντέλων σε μια κεντρική δομή μπορεί να οδηγήσει σε ένα σημείο αποτυχίας (single point of failure, SPOF), δηλαδή η αδυναμία λειτουργίας ενός μέρους του συστήματος που οδηγεί σε αδυναμία και ολόκληρο το υπόλοιπο σύστημα. Επομένως, απαιτούνται πρόσθετοι ενδιάμεσοι κόμβοι επεξεργασίας κατά τη διάρκεια της εκπαίδευσης.

8.3.6 Ερμηνευσιμότητα

Το FL δίκτυο που προκύπτει μέσω της μάθησης αποτελεί ένα μαύρο κουτί του οποίου οι εσωτερικοί μηχανισμοί είναι δύσκολο να κατανοηθούν. Σε αυτήν την περίπτωση, είναι απαραίτητη μια ευκολότερα ερμηνεύσιμη προσέγγιση, για την αποφυγή κακόβουλων επιθέσεων στον κεντρικό server, σε περίπτωση που αυτός που επιχειρεί την επίθεση γνωρίζει τις παραμέτρους των άλλων clients. Σε αυτό το πλαίσιο, όπως επισημαίνουν και οι συγγραφείς [54], η υλοποίηση της δομής του Federated δικτύου σε ένα πολυεπίπεδο Federated κέντρο απαιτεί περαιτέρω διερεύνηση μελλοντικά.

Βιβλιογραφία

- [1] Alpaydin, E. (2020). *Introduction to machine learning*. MIT press.
- [2] Bishop, C. M., & Nasrabadi, N. M. (2006). *Pattern recognition and machine learning* (Vol. 4, No. 4, p. 738). New York: springer.
- [3] Carbonell, J. G., Michalski, R. S., & Mitchell, T. M. (1983). An overview of machine learning. *Machine learning*, 3-23.
- [4] Flach, P. (2012). *Machine learning: the art and science of algorithms that make sense of data*. Cambridge university press.
- [5] Hastie, T., Tibshirani, R., Friedman, J. H., & Friedman, J. H. (2009). *The elements of statistical learning: data mining, inference, and prediction* (Vol. 2, pp. 1-758). New York: springer.
- [6] Shalev-Shwartz, S., & Ben-David, S. (2014). *Understanding machine learning: From theory to algorithms*. Cambridge university press.
- [7] Khanzode, K. C. A., & Sarode, R. D. (2020). Advantages and disadvantages of artificial intelligence and machine learning: A literature review. *International Journal of Library & Information Science (IJLIS)*, 9(1), 3.
- [8] Mitchell, T. M. (1997). *Machine learning*.
- [9] Shinde, P. P., & Shah, S. (2018, August). A review of machine learning and deep learning applications. In *2018 Fourth international conference on computing communication control and*
- [10] Nasteski, V. (2017). An overview of the supervised machine learning methods. *Horizons*, 4, 51-62.
- [11] Koehrsen, W. (2018). Overfitting vs. underfitting: A complete example. *Towards Data Science*, 405.
- [12] Ghahramani, Z. (2003). Unsupervised learning. In *Summer school on machine learning* (pp. 72-112). Berlin, Heidelberg: Springer Berlin Heidelberg.
- [13] Maulud, D., & Abdulazeez, A. M. (2020). A review on linear regression comprehensive in machine learning. *Journal of Applied Science and Technology Trends*, 1(4), 140-147.
- [14] Peckov, A. (2012). A machine learning approach to polynomial regression. *Ljubljana, Slovenia*, URL: http://kt.ijs.si/theses/phd_aleksandar_peckov.pdf.

- [15] Guo, G., Wang, H., Bell, D., Bi, Y., & Greer, K. (2003). KNN model-based approach in classification. In *On The Move to Meaningful Internet Systems 2003: CoopIS, DOA, and ODBASE: OTM Confederated International Conferences, CoopIS, DOA, and ODBASE 2003, Catania, Sicily, Italy, November 3-7, 2003. Proceedings* (pp. 986-996). Springer Berlin Heidelberg.
- [16] Mahesh, B. (2020). Machine learning algorithms-a review. *International Journal of Science and Research (IJSR)*. [Internet], 9(1), 381-386.
- [17] Somvanshi, M., Chavan, P., Tambade, S., & Shinde, S. V. (2016, August). A review of machine learning techniques using decision tree and support vector machine. In *2016 international conference on computing communication control and automation (ICCUBE)* (pp. 1-7). IEEE.
- [18] Segal, M. R. (2004). Machine learning benchmarks and random forest regression.
- [19] Chen, T., & Guestrin, C. (2016, August). Xgboost: A scalable tree boosting system. In *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining* (pp. 785-794).
- [20] Likas, A., Vlassis, N., & Verbeek, J. J. (2003). The global k-means clustering algorithm. *Pattern recognition*, 36(2), 451-461.
- [21] Heller, K. A., & Ghahramani, Z. (2005, August). Bayesian hierarchical clustering. In *Proceedings of the 22nd international conference on Machine learning* (pp. 297-304).
- [22] Deng, D. (2020, September). DBSCAN clustering algorithm based on density. In *2020 7th international forum on electrical engineering and automation (IFEAA)* (pp. 949-953). IEEE.
- [23] Sorzano, C. O. S., Vargas, J., & Montano, A. P. (2014). A survey of dimensionality reduction techniques. *arXiv preprint arXiv:1403.2877*.
- [24] George, A., & Vidyapeetham, A. (2012). Anomaly detection based on machine learning: dimensionality reduction using PCA and classification using SVM. *International Journal of Computer Applications*, 47(21), 5-8.
- [25] Reza, M. S., & Ma, J. (2016, November). ICA and PCA integrated feature extraction for classification. In *2016 IEEE 13th International Conference on Signal Processing (ICSP)* (pp. 1083-1088). IEEE.
- [26] Wei, X., & Croft, W. B. (2006, August). LDA-based document models for ad-hoc retrieval. In *Proceedings of the 29th annual international ACM SIGIR conference on Research and development in information retrieval* (pp. 178-185).
- [27] Lan, A. S., Waters, A. E., Studer, C., & Baraniuk, R. G. (2014). Sparse factor analysis for learning and content analytics. *The Journal of Machine Learning Research*, 15(1), 1959-2008.
- [28] Lawrence, J. (1993). *Introduction to neural networks*. California Scientific Software.

- [29] Bianchini, M., & Scarselli, F. (2014). On the complexity of neural network classifiers: A comparison between shallow and deep architectures. *IEEE transactions on neural networks and learning systems*, 25(8), 1553-1565.
- [30] Krogh, A. (2008). What are artificial neural networks?. *Nature biotechnology*, 26(2), 195-197.
- [31] Ray, S. (2019, February). A quick review of machine learning algorithms. In *2019 International conference on machine learning, big data, cloud and parallel computing (COMITCon)* (pp. 35-39). IEEE.
- [32] Mercioni, M. A., & Holban, S. (2020, May). The most used activation functions: Classic versus current. In *2020 International Conference on Development and Application Systems (DAS)* (pp. 141-145). IEEE.
- [33] Boden, M. (2002). A guide to recurrent neural networks and backpropagation. *the Dallas project*, 2(2), 1-10.
- [34] Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated learning: Challenges, methods, and future directions. *IEEE signal processing magazine*, 37(3), 50-60.
- [35] Cunningham, P., Cord, M., & Delany, S. J. (2008). Supervised learning. In *Machine learning techniques for multimedia: case studies on organization and retrieval* (pp. 21-49). Berlin, Heidelberg: Springer Berlin Heidelberg.
- [36] Caruana, R., & Niculescu-Mizil, A. (2006, June). An empirical comparison of supervised learning algorithms. In *Proceedings of the 23rd international conference on Machine learning* (pp. 161-168).
- [37] Skianis, K., Giannopoulos, A., Spantideas, S., Hatzaki, M., Karditsa, A., & Trakadas, P. SWIRL: Statistical downscaling for Wind Pattern Reconstruction using Machine Learning.
- [38] Sutton, R. S., & Barto, A. G. (2018). *Reinforcement learning: An introduction*. MIT press.
- [39] Wiering, M. A., & Van Otterlo, M. (2012). Reinforcement learning. *Adaptation, learning, and optimization*, 12(3), 729.
- [40] Madhulatha, T. S. (2012). An overview on clustering methods. *arXiv preprint arXiv:1205.1117*.
- [41] Reddy, G. T., Reddy, M. P. K., Lakshmana, K., Kaluri, R., Rajput, D. S., Srivastava, G., & Baker, T. (2020). Analysis of dimensionality reduction techniques on big data. *Ieee Access*, 8, 54776-54788.
- [42] Giannopoulos, A., Gkonis, P., Bithas, P., Nomikos, N., Ntroulias, G., & Trakadas, P. (2023). Federated Learning for Maritime Environments: Use Cases, Experimental Results, and Open Issues.

- [43] Torrey, L., & Shavlik, J. (2010). Transfer learning. In *Handbook of research on machine learning applications and trends: algorithms, methods, and techniques* (pp. 242-264). IGI global.
- [44] Dietterich, T. G. (2002). Ensemble learning. *The handbook of brain theory and neural networks*, 2(1), 110-125.
- [45] Skianis, K., Giannopoulos, A., Gkonis, P., & Trakadas, P. (2023). Data Aging Matters: Federated Learning-Based Consumption Prediction in Smart Homes via Age-Based Model Weighting. *Electronics*, 12(14), 3054.
- [46] Giannopoulos, A., Nomikos, N., Ntroulias, G., Syriopoulos, T., & Trakadas, P. (2023, June). Maritime Federated Learning for Decentralized On-Ship Intelligence. In *IFIP International Conference on Artificial Intelligence Applications and Innovations* (pp. 195-206). Cham: Springer Nature Switzerland.
- [47] Kaloxylos, A., Gavras, A., Camps, D., Ghorraishi, M., & Hrasnica, H. (2021). AI and ML-Enablers for beyond 5G Networks.
- [48] Trakadas, P., Masip-Bruin, X., Facca, F. M., Spantideas, S. T., Giannopoulos, A. E., Kapsalis, N. C., ... & Lyridis, D. V. (2022). A Reference Architecture for Cloud-Edge Meta-Operating Systems Enabling Cross-Domain, Data-Intensive, ML-Assisted Applications: Architectural Overview and Key Concepts. *Sensors*, 22(22), 9003.
- [49] Giannopoulos, A., Spantideas, S., Kapsalis, N., Gkonis, P., Sarakis, L., Capsalis, C., ... & Trakadas, P. (2022). Supporting intelligence in disaggregated open radio access networks: Architectural principles, AI/ML workflow, and use cases. *IEEE Access*, 10, 39580-39595.
- [50] Giannopoulos, A., Spantideas, S., Kapsalis, N., Karkazis, P., & Trakadas, P. (2021). Deep reinforcement learning for energy-efficient multi-channel transmissions in 5G cognitive hetnets: Centralized, decentralized and transfer learning based solutions. *IEEE Access*, 9, 129358-129374.
- [51] McMahan, B., Moore, E., Ramage, D., Hampson, S., & y Arcas, B. A. (2017, April). Communication-efficient learning of deep networks from decentralized data. In *Artificial intelligence and statistics* (pp. 1273-1282). PMLR.
- [52] Moubayed, A., & Shami, A. (2020). Softwarization, virtualization, and machine learning for intelligent and effective vehicle-to-everything communications. *IEEE Intelligent Transportation Systems Magazine*, 14(2), 156-173.
- [53] [Guan and Wang, 2023] Guan, W., & Wang, K. (2023). Autonomous Collision Avoidance of Unmanned Surface Vehicles Based on Improved A-Star and Dynamic Window Approach Algorithms. *IEEE Intelligent Transportation Systems Magazine*.
- [54] Wang, S., & Zhang, Y. (2022). Multi-Level Federated Network Based on Interpretable Indicators for Ship Rolling Bearing Fault Diagnosis. *Journal of Marine Science and Engineering*, 10(6), 743.

- [55] Aledhari, M., Razzak, R., Parizi, R. M., & Saeed, F. (2020). Federated learning: A survey on enabling technologies, protocols, and applications. *IEEE Access*, 8, 140699-140725.
- [56] Han, C., & Yang, T. (2021, July). Privacy protection technology of maritime multi-agent communication based on part-federated learning. In *2021 IEEE/CIC International Conference on Communications in China (ICCC Workshops)* (pp. 266-271). IEEE.
- [57] Liu, W., Xu, X., Wu, L., Qi, L., Jolfaei, A., Ding, W., & Khosravi, M. R. (2022). Intrusion detection for maritime transportation systems with batch federated aggregation. *IEEE Transactions on Intelligent Transportation Systems*, 24(2), 2503-2514.
- [58] Niknam, S., Dhillon, H. S., & Reed, J. H. (2020). Federated learning for wireless communications: Motivation, opportunities, and challenges. *IEEE Communications Magazine*, 58(6), 46-51.
- [59] Angelopoulos, A., Giannopoulos, A., Spantideas, S., Kapsalis, N., Trochoutsos, C., Voliotis, S., & Trakadas, P. (2022, June). Allocating orders to printing machines for defect minimization: A comparative machine learning approach. In *IFIP International Conference on Artificial Intelligence Applications and Innovations* (pp. 79-88). Cham: Springer International Publishing.
- [60] Spantideas, S. T., Giannopoulos, A. E., Kapsalis, N. C., Angelopoulos, A., Voliotis, S., & Trakadas, P. (2022). Towards Zero-Defect Manufacturing: Machine Selection through Unsupervised Learning in the Printing Industry. *Proceedings <http://ceur-ws.org> ISSN, 1613, 0073*.
- [61] Giannopoulos, A. E., Spantideas, S. T., Nikolopoulos, C. D., Baklezos, A. T., & Capsalis, C. N. (2022, May). Dipole Fitting in Unit-Level Spacecraft Equipment with Deep Neural Networks. In *2022 ESA Workshop on Aerospace EMC (Aerospace EMC)* (pp. 1-5). IEEE.