



ΠΑΝΕΠΙΣΤΗΜΙΟ ΑΙΓΑΙΟΥ

ΠΟΛΥΤΕΧΝΙΚΗ ΣΧΟΛΗ

ΤΜΗΜΑ ΜΗΧΑΝΙΚΩΝ ΠΛΗΡΟΦΟΡΙΑΚΩΝ ΚΑΙ ΕΠΙΚΟΙΝΩΝΙΑΚΩΝ ΣΥΣΤΗΜΑΤΩΝ

ΠΡΟΓΡΑΜΜΑ ΜΕΤΑΠΤΥΧΙΑΚΩΝ ΣΠΟΥΔΩΝ

ΠΛΗΡΟΦΟΡΙΑΚΑ ΚΑΙ ΕΠΙΚΟΙΝΩΝΙΑΚΑ ΣΥΣΤΗΜΑΤΑ

**Ανασκόπηση της αυτόματης αξιολόγησης υποψηφίων σε θέσεις εργασίας
μέσω μηχανικής μάθησης.**

**Survey on the automatic evaluation of job candidates based on machine
learning.**

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ

της

ΜΑΡΙΝΙΔΗ ΕΛΕΝΗΣ

Επιβλέπων: κ. Κρητικός Κυριάκος

Μέλη εξεταστικής επιτροπής: κ.κ. Κωστούλας & Συμεωνίδης

Σάμος, Φεβρουάριος 2025

Η σελίδα αυτή είναι σκόπιμα λευκή.

Ευχαριστίες

Η παρούσα μεταπτυχιακή εργασία με τίτλο: «Ανασκόπηση της αυτόματης αξιολόγησης υποψηφίων σε θέσεις εργασίας μέσω μηχανικής μάθησης» εκπονήθηκε στα πλαίσια του Προγράμματος Μεταπτυχιακών Σπουδών «Πληροφορικά και Επικοινωνιακά Συστήματα» του Τμήματος Μηχανικών Πληροφοριακών και Επικοινωνιακών Συστημάτων του Πανεπιστημίου Αιγαίου.

Η εκπόνηση αυτής της εργασίας δεν θα μπορούσε να πραγματοποιηθεί χωρίς την ουσιαστική συμβολή κάποιων ανθρώπων, τους οποίους και θα ήθελα να ευχαριστήσω.

Αρχικά θα ήθελα να ευχαριστήσω τον επιβλέποντα καθηγητή της εργασίας μου, τον Καθηγητή του Πανεπιστημίου Αιγαίου κ. Κρητικό Κυριάκο, για την συστηματική του καθοδήγηση, καθ' όλη τη διάρκεια της εκπόνησης. Η προθυμία και η άμεση απόκρισή του υπήρξε καθοριστική για το αποτέλεσμα της μεταπτυχιακής μου εργασίας. Είμαι πραγματικά ευγνώμων για την βαθιά ευγένεια και ενσυναίσθηση που τον διακρίνει.

Ιδιαίτερες ευχαριστίες θα ήθελα να εκφράσω στους γονείς μου Ζωή και Γιώργο, διότι χωρίς τη συμπαράσταση, την ενθάρρυνση και τη στήριξή τους σε όλη αυτή τη διαδρομή, δεν θα πραγματοποιούνταν αυτή η εργασία.

Ελίνα Μαρινίδη

© 2025

της

Μαρινίδη Ελένης

Τμήμα Μηχανικών Πληροφοριακών και Επικοινωνιακών Συστημάτων

ΠΑΝΕΠΙΣΤΗΜΙΟ ΑΙΓΑΙΟΥ

Η σελίδα αυτή είναι σκόπιμα λευκή.

Περίληψη

Η ταχύτατη ανάπτυξη της τεχνολογίας έχει οδηγήσει σε σημαντικές εξελίξεις σε πολλούς τομείς και διαδικασίες, συμπεριλαμβανομένων και αυτών της πρόσληψης προσωπικού. Οι παραδοσιακές μέθοδοι είναι συχνά χρονοβόρες καθώς και επιρρεπείς σε σφάλματα και προκαταλήψεις. Η Μηχανική Μάθηση, μέρος της Τεχνητής Νοημοσύνης, προσφέρει τη δυνατότητα αυτοματοποίησης αυτών των διαδικασιών, εξασφαλίζοντας μεγαλύτερη αντικειμενικότητα, ταχύτητα και ακρίβεια. Ωστόσο, ζητούμενο παραμένει η επίτευξη υψηλής ακρίβειας στην αξιολόγηση των υποψηφίων, η οποία βασίζεται σε δεδομένα ουσιαστικών και τυπικών προσόντων.

Η παρούσα διπλωματική εργασία εστιάζει στο να διερευνήσει και να αξιολογήσει τις υφιστάμενες προσεγγίσεις αυτοματοποίησης στη διαδικασία πρόσληψης με τη χρήση αλγορίθμων Μηχανικής Μάθησης. Εξετάζονται τα πλεονεκτήματα και τα μειονεκτήματα της κάθε τέτοιας προσέγγισης ενώ επίσης αυτές συγκρίνονται με βάση ένα καλά σχεδιασμένο σύνολο από κριτήρια. Επιπλέον, βάσει και των αποτελεσμάτων της σύγκρισης, παρουσιάζονται κατευθυντήριες γραμμές για τη βελτίωση της ακρίβειας και της αποτελεσματικότητας των προσεγγίσεων αυτών. Τα ευρήματα αναδεικνύουν τη δυναμική της τεχνολογίας στη δημιουργία διαφανών, αντικειμενικών και οικονομικά αποδοτικών διαδικασιών πρόσληψης, παρέχοντας πολύτιμες γνώσεις για μελλοντική έρευνα και ανάπτυξη, τόσο στον ιδιωτικό, όσο και στον δημόσιο τομέα.

Λέξεις κλειδιά: *Επεξεργασία Φυσικής Γλώσσας (NLP), Αντιστοίχιση Δεξιοτήτων, Ανάλυση & Έλεγχος Δεδομένων Υποψηφίων, Βιογραφικά, Κατάταξης Υποψηφίων, Προσλήψεις, Αυτόματη Πρόβλεψη για Πρόσληψη, Μηχανική Μάθηση, Αλγόριθμοι Μηχανικής Μάθησης, Συνεντεύξεις Εργασίας, Συμπεράσματα Προσωπικότητας, Ουσιαστικά Προσόντα, Τυπικά Προσόντα, Επιλογή Προσωπικού.*

ABSTRACT

The rapid evolution of technology has led to significant developments in many areas and processes, including those of personnel recruitment. Traditional methods are often time-consuming, prone to errors and biases. Machine Learning (ML), as a part of Artificial Intelligence, offers the possibility of automating these processes, ensuring greater objectivity, speed and accuracy. However, the challenge remains to achieve high accuracy in the assessment of candidates, which is based on data of substantive and formal qualifications.

This thesis focuses on investigating and evaluating existing automation approaches in the recruitment process using Machine Learning algorithms. The advantages and disadvantages of these approaches are examined while they are compared according to a well-designed criteria set. Furthermore, based on the comparison results, guidelines are presented for improving the accuracy and effectiveness of these approaches. The findings highlight the potential of ML technology in creating transparent, objective and cost-effective recruitment processes, providing valuable insights for future research and development, both in the private and public sectors.

Key Words: *Natural Language Processing (NLP), Skill Matching, Candidate Data Analysis & Verification, Resumes, Candidate Ranking, Recruitment, Automatic Hiring Prediction, Machine Learning, Machine Learning Algorithms, Job Interviews, Personality Inference, Essential Qualifications, Formal Qualifications, Personnel Selection.*

ΛΙΣΤΑ ΕΙΚΟΝΩΝ	8
ΛΙΣΤΑ ΠΙΝΑΚΩΝ	9
ΑΚΡΩΝΥΜΙΑ	10
1. ΕΙΣΑΓΩΓΗ	12
2. ΜΗΧΑΝΙΚΗ ΜΑΘΗΣΗ	15
2.1 Εισαγωγή & Ορισμός	15
2.2 Είδη Μηχανικής Μάθησης	17
2.2.1 Επιβλεπόμενη Μάθηση	17
2.2.2 Μη Επιβλεπόμενη Μάθηση	18
2.2.3 Ενισχυτική Μάθηση	19
2.3 Πλεονεκτήματα & Μειονεκτήματα	21
2.4 Εφαρμογές Μηχανικής Μάθησης	22
2.4.1 Γενικές Εφαρμογές	22
2.4.2 Εφαρμογές στη Διαδικασία Πρόσληψης Υποψηφίων	23
3. ΔΙΑΔΙΚΑΣΙΑ ΠΡΟΣΛΗΨΗΣ	25
3.1 Παραδοσιακή Διαδικασία Πρόσληψης	26
3.2 Ιδιαιτερότητες Διαδικασίας Πρόσληψης στο Δημόσιο Τομέα – Α.Σ.Ε.Π.	27
3.3 Βελτιώσεις Διαδικασίας Πρόσληψης μέσω Μηχανικής Μάθησης	30
4. ΜΕΘΟΔΟΛΟΓΙΑ ΒΙΒΛΙΟΓΡΑΦΙΚΗΣ ΑΝΑΣΚΟΠΗΣΗΣ	32
4.1 Στόχος Ανασκόπησης	32
4.2 Ερευνητικά Ερωτήματα	32
4.3 Τρόπος Αναζήτησης Άρθρων	33
4.4 Φιλτράρισμα Άρθρων	34
4.5 Ποσοτική Ανάλυση	35
4.5.1 Διαγράμματα Πίτας	36
4.5.2 Γραφήματα	37
5. Ανάλυση Βιβλιογραφίας	38
5.1 Ιεραρχική Κατηγοριοποίηση Άρθρων	38
5.2 Ανάλυση Άρθρων	45
5.2.1 Ανάλυση Άρθρων Εστιαζόμενων στα Ουσιαστικά Προσόντα	45
5.2.2 Ανάλυση Άρθρων Εστιαζόμενων στα Τυπικά Προσόντα	62
5.3 Συγκριτική Αξιολόγηση Άρθρων	80
5.3.1 Κριτήρια Σύγκρισης	80
5.3.2 Αποτελέσματα Σύγκρισης	84
5.3.3 Συμπεράσματα Σύγκρισης	87
5.3.3.1 Συμπεράσματα Άρθρων που Εστιάζουν στα Ουσιαστικά Προσόντα.	87
5.3.3.2 Συμπεράσματα Άρθρων που Εστιάζουν στα Τυπικά Προσόντα.	89
5.3.3.3 Καθολικά Συμπεράσματα	91
6. ΣΥΜΠΕΡΑΣΜΑΤΑ	92
6.1 Συνεισφορά Διπλωματικής & Σημαντικότερα Αποτελέσματα	92
6.2 Κρίσιμα Ερευνητικά Κενά & Κατευθύνσεις	93
6.3 Προσωπική Εμπειρία	97
7. ΒΙΒΛΙΟΓΡΑΦΙΚΕΣ ΑΝΑΦΟΡΕΣ	98

ΛΙΣΤΑ ΕΙΚΟΝΩΝ

Εικόνα 1 - Γράφημα πίτας για τα ποσοστά άρθρων ανά είδος	33
Εικόνα 2 - Γράφημα πίτας για τα ποσοστά άρθρων ανά εκδότη	34
Εικόνα 3 - Αριθμός άρθρων ανά έτος δημοσίευσης	35
Εικόνα 4 - Οι κατηγορίες των ουσιαστικών προσόντων και οι υποκατηγορίες τους	36
Εικόνα 5 - Οι κατηγορίες των τυπικών προσόντων	37
Εικόνα 6 - Παράθεση των εργαλείων και εφαρμογών που χρησιμοποιήθηκαν στα άρθρα με τα ουσιαστικά προσόντα	38
Εικόνα 7 - Παράθεση των εργαλείων και εφαρμογών που χρησιμοποιήθηκαν στα άρθρα με τα τυπικά προσόντα	39
Εικόνα 8 - Οι κατηγορίες των αλγορίθμων MM και NN	40
Εικόνα 9 - Τα νευρωνικά δίκτυα Μηχανικής Μάθησης που χρησιμοποιήθηκαν στις έρευνες	41
Εικόνα 10 - Οι κατηγορίες των συνεντεύξεων και από ποιους αξιολογούνται οι υποψήφιοι	42

ΛΙΣΤΑ ΠΙΝΑΚΩΝ

Πίνακας 1 - Πίνακας σύγκρισης ουσιαστικών προσόντων	87
Πίνακας 2 - Πίνακας σύγκρισης ουσιαστικών προσόντων	87
Πίνακας 3 - Πίνακας σύγκρισης τυπικών προσόντων	88

ΑΚΡΩΝΥΜΙΑ

ΑΚΡΩΝΥΜΙΑ	ΟΡΙΣΜΟΣ
ANN	Artificial Neural Networks
API	Application Programming Interface
AS	Automated Scoring
ASR	Automatic Speech Recognition
AUC	Area Under the Curve
BERT	Bidirectional Encoder Representations from Transformers
BoW	Bag of Words
CNN	Convolution Neural Networks
CSS	Cascading Style Sheets
CV	Curriculum Vitae
DBSCAN	Density-Based Spatial Clustering of Applications with Noise
DFT	Discrete Fourier Transform
EDA	Exploratory Data Analysis
ETS	Educational Testing Service
GDPR	General Data Protection Regulation
GENSIM	Generate Similar
GNB	Gaussian Naive Bayes
GPU	Graphics Processing Unit
GRU	Gated Recurrent Unit
HR	Human Resources
HTML	Hyper Text Markup Language
IDFT	Inverse Discrete Fourier Transform
KNN	K-Nearest Neighbor
LASSO	Least Absolute Shrinkage and Selection Operator
LDA	Latent Dirichlet Allocation
LIWC	Linguistic Inquiry Word Count
LLE	Locally Linear Embedding
LLM	Large Language Models
MCTS	Monte Carlo Tree Search
MIT	Massachusetts Institute of Technology
MLP	Multi-Layer Perceptron
NER	Named Entity Recognition
NLP	Natural Language Processing
NLTK	Natural Language Toolkit
NN	Neural Network
PCA	Principal Component Analysis
PRAAT	Προέρχεται από την ολλανδική λέξη "ομιλία"

RF	Random Forest
RNN	Recurrent Neural Networks
RR	Ridge Regression
SSP	Social Signal Processing
SVM	Support Vector Machine
TF-IDF	Term Frequency-Inverse Document Frequency
UBIAI	Universal Business Intelligence and Artificial Intelligence
VADER	Valence Aware Dictionary and sentiment Reasoner
ΑΣΕΠ	Ανώτατο Συμβούλιο Επιλογής Προσωπικού
MM	Μηχανική Μάθηση
ΝΔ	Νευρωνικά Δίκτυα

1. ΕΙΣΑΓΩΓΗ

Η εξέλιξη της τεχνολογίας υπήρξε ραγδαία τα τελευταία χρόνια. Η Μηχανική Μάθηση, κομμάτι της Τεχνητής Νοημοσύνης, έχει βελτιώσει πολλούς τομείς της ζωής μας και αναμένεται να αλλάξει ακόμα περισσότερο την οργάνωση, τη στρατηγική, τον σχεδιασμό, τις επιχειρηματικές διεργασίες παγκοσμίως και οπουδήποτε αλλού εφαρμοστεί. Η ανάγκη βελτιστοποίησης της διοίκησης των επιχειρήσεων και των υπηρεσιών, ώθησε την τεχνολογία να εξερευνήσει και να εξελίξει αυτό το πεδίο. Όσο η τεχνολογία εξελίσσεται, δημιουργούνται νέες τεχνικές, νέες εφαρμογές και νέες μέθοδοι.

Η συνήθης διαδικασία πρόσληψης προσωπικού στον ιδιωτικό και στο δημόσιο τομέα, τα χρόνια πριν την ανάπτυξη της τεχνολογίας, ήταν απλή. Ο υποψήφιος συμπλήρωνε μία αίτηση και την προσκόμιζε στην αρμόδια υπηρεσία. Ακολουθώντας, οι αρμόδιοι εξέταζαν τα τυπικά προσόντα του υποψήφιου και αποφάσιζαν ποιοί ήταν οι καταλληλότεροι για τη συμπλήρωση της θέσης. Οι τελευταίοι καλούνταν για συνέντευξη δια ζώσης, ή στα παλαιότερα χρόνια που δεν ήταν εύκολες οι μετακινήσεις, δια τηλεφώνου, για να αξιολογηθούν και τα ουσιαστικά τους προσόντα (Langer et al, 2020). Είναι πλέον αποδεκτό ότι δεν αρκεί η αξιολόγηση των τυπικών προσόντων, για να κριθεί η καταλληλότητα ή μη κάποιου, για μια θέση εργασίας. Χαρακτηριστικά όπως η ομαδικότητα, η ηγετικότητα, η νοημοσύνη, η κρίση, η αντίληψη, η προσαρμοστικότητα, η δημιουργικότητα, η πρωτοτυπία, η μεθοδικότητα, η συναδελφικότητα και η αλληλεγγύη συνεκτιμώνται για την τελική απόφαση της πρόσληψης. Χαρακτηριστικά που μπορούν να εκτιμηθούν μόνο μέσω συνέντευξης.

Καθώς η τεχνολογία εξελισσόταν, οι αιτήσεις μπορούσαν να υποβάλλονται ηλεκτρονικά και η αρχική διαλογή από το τμήμα ανθρώπινου δυναμικού να γινόταν ευκολότερα και γρηγορότερα με τη βοήθεια συγκεντρωτικών ηλεκτρονικών πινάκων (π.χ. Excel) (Naik & Dhotre, 2022). Σε αυτή τη φάση της διαδικασίας υπήρχε το ενδεχόμενο να γίνουν λάθη. Η παράλειψη κάποιου σημαντικού στοιχείου του βιογραφικού σημειώματος ή η λανθασμένη καταχώρισή του ήταν συχνό φαινόμενο. Στη φάση της συνέντευξης υπήρχε, και πάντα θα υπάρχει, η υποκειμενική κρίση και η προκατάληψη. Η Μηχανική Μάθηση όμως ήρθε να διορθώσει λάθη του παρελθόντος. Εκτιμάται ότι σύντομα οι προσλήψεις θα γίνονται αυτόματα, χωρίς την συμβολή του ανθρώπινου παράγοντα. Ήδη οι εταιρείες HireVue και Precire, στις

Η.Π.Α. και την Γερμανία αντίστοιχα, χρησιμοποιούν αλγόριθμους Μηχανικής Μάθησης για την αξιολόγηση των υποψηφίων. Η πρώτη χρησιμοποιεί ηχογραφήσεις βίντεο-συνεντεύξεων και η δεύτερη φωνητικές τηλεφωνικές ηχογραφήσεις (Langer et al, 2020). Ο αριθμός εταιρειών που χρησιμοποιούν παρόμοιες αυτόματες διαδικασίες πρόσληψης εργαζόμενων κάθε χρόνο αυξάνεται (Anju Raj & Farzana, 2023). Ήδη εταιρείες κάνουν την πρώτη διαλογή των υποψηφίων δια τηλεφώνου όπου “διαδραστικοί” υπεύθυνοι προσλήψεων τεχνητής νοημοσύνης μέσω ερωτήσεων συγκεντρώνουν βασικές πληροφορίες και αποφασίζουν εάν κάποιος θα περάσει στην επόμενη φάση της αξιολόγησης. Το έτος 2024, η αμερικάνικη εταιρεία λογισμικού προσλήψεων Paradox, εξυπηρετούσε περισσότερους από 1.000 πελάτες παγκοσμίως, όπως η Pfizer, η General Motors, η Amazon και η Mcdonald's. Σύμφωνα με την Wu, συνεργάτιδα του αμερικάνικου καναλιού επιχειρηματικών ειδήσεων CNBC, οι διαδραστικές συνεντεύξεις με την τεχνητή νοημοσύνη θα είναι ο νέος κανόνας το 2025 (Wu, 2024).

Κατά το παρελθόν υπήρξαν προβλήματα και περιορισμοί στην έρευνα και στα πειράματα για την αυτόματη εξαγωγή συμπεράσματος πρόσληψης, ή μη, υποψηφίων, ενώ και τα ποσοστά πρόβλεψης δεν ήταν αρκετά υψηλά. Ο αριθμός δειγμάτων των συνεντεύξεων ήταν χαμηλός, με αποτέλεσμα η εκπαίδευση των αλγορίθμων Μηχανικής Μάθησης να παρουσιάζει κενά. Οι εφαρμογές για την εξαγωγή στοιχείων δεν ήταν ακριβείς. Με συνεχείς όμως πειραματισμούς και έρευνα, η πρόοδος υπήρξε αναμενόμενη. Ένα πεδίο ακόμα που συνεχώς εξελίσσεται, είναι η ποικιλομορφία του δείγματος. Η σωστή έρευνα οφείλει να βασίζεται σε ισορροπημένα δεδομένα με βάση όλα τα απαραίτητα χαρακτηριστικά των υποψηφίων, όπως είναι το φύλο, η εθνικότητα, ο πολιτισμός και η κουλτούρα, η ηλικία και άνθρωποι από διαφορετικό κοινωνικό, οικονομικό και φυλετικό υπόβαθρο (Chen et al, 2017). Η ακρίβεια πρόβλεψης έχει βελτιωθεί σημαντικά, αλλά όχι τόσο ώστε να επέλθει πλήρης αυτοματοποίηση. Ακόμα υπάρχει ο κίνδυνος πρόσληψης ενός ατόμου που δεν πληροί τα πρόσθετα προσόντα που απαιτούνται για την εκάστοτε θέση. Η συνήθης διαδικασία που ακολουθείται στα άρθρα για την επίτευξη του αυτόματου συμπεράσματος του καταλληλότερου υποψηφίου για μία θέση, είναι η ομαδοποίηση (clustering) των χαρακτηριστικών και των στοιχείων, η εξαγωγή τους (extraction) και η ταξινόμηση τους (classification).

Η παρούσα διπλωματική εργασία εξετάζει ποικίλες προτεινόμενες μεθόδους βασιζόμενες στη Μηχανική Μάθηση για την αυτόματη εξαγωγή συμπεράσματος όσον αφορά την καταλληλότητα των υποψηφίων. Η συνεισφορά της έγκειται στην ανασκόπηση και την αξιολόγηση των σύγχρονων μεθόδων Μηχανικής Μάθησης για την αυτόματη επιλογή υποψηφίων υπαλλήλων. Επιπλέον, εξετάζονται ποικίλοι τρόποι και προσεγγίσεις για την εξασφάλιση αμερόληπτων, ορθών και ισορροπημένων αποτελεσμάτων. Η εργασία προσδοκά στο να συμβάλλει στη βελτίωση της αυτοματοποίησης των διαδικασιών πρόσληψης, μειώνοντας τον χρόνο και το κόστος της διαδικασίας επιλογής κατάλληλου προσωπικού, καθώς και στο να παρουσιάσει μεθόδους και προσεγγίσεις με αξιόπιστα και αντικειμενικά αποτελέσματα. Ακόμη, η χρήση δεδομένων με ποικίλες και διαφορετικές παραμέτρους, π.χ. κοινωνικό, πολιτισμικό, οικονομικό υπόβαθρο, φύλο, ηλικία κ.λπ., τονίζεται ώστε να εξαλειφθεί η όποια προκατάληψη που μπορεί να οδηγήσει σε εσφαλμένα συμπεράσματα.

Η παρούσα διπλωματική εργασία ενδέχεται να αποτελέσει τη βάση για περαιτέρω έρευνα στον τομέα της Μηχανικής μάθησης, με ιδιαίτερη έμφαση στον δημόσιο τομέα. Στόχος είναι η υιοθέτηση εξελιγμένων διαδικασιών, ώστε οι προσλήψεις να πραγματοποιούνται με ελάχιστα περιθώρια σφάλματος, χωρίς προκαταλήψεις, και με ελάχιστο κόστος και χρόνο. Οι αναλύσεις των μεθόδων και των προσεγγίσεων που περιγράφονται, μπορούν να λειτουργήσουν ως κατευθυντήριες γραμμές για τη σχεδίαση και την ανάπτυξη βελτιωμένων συστημάτων πρόσληψης στον ιδιωτικό αλλά και στον δημόσιο τομέα. Η συνδυαστική παρουσίαση και αξιολόγηση τόσο των ουσιαστικών, όσο και των τυπικών προσόντων των υποψηφίων, συγκεντρωμένων σε μία εργασία, προσφέρει άμεση και αποδοτική σύγκριση των μεθόδων και των προσεγγίσεων.

Η τρέχουσα αναφορά δομείται ως προς το υπόλοιπό της ως εξής. Στο δεύτερο κεφάλαιο αναλύονται η έννοια της Μηχανικής Μάθησης, τα είδη της, τα πλεονεκτήματα αλλά και οι αδυναμίες της καθώς και οι εφαρμογές της. Ακολουθεί το τρίτο κεφάλαιο που περιγράφει τις διαδικασίες πρόσληψης στον ιδιωτικό και τον δημόσιο τομέα, καθώς και τις αναμενόμενες βελτιώσεις αυτών με τη χρήση της Μηχανικής Μάθησης. Στο τέταρτο κεφάλαιο εξετάζεται η μεθοδολογία της βιβλιογραφικής ανασκόπησης και πραγματοποιείται μια πρώτη ποσοτική ανάλυση, ιδίως στατιστικής φύσεως, της επιλεγμένης βιβλιογραφίας. Η βασική και εκτεταμένη

ανάλυση της επιλεγμένης βιβλιογραφίας πραγματοποιείται στο πέμπτο κεφάλαιο όπου παρουσιάζεται μια ιεραρχική κατηγοριοποίησή της καθώς και συγκριτική ανάλυσή της, μέσω ενός συνόλου από επιλεγμένα κριτήρια. Στο έκτο κεφάλαιο παρουσιάζονται τα βασικά συμπεράσματα από την βιβλιογραφική επισκόπηση που πραγματοποιήθηκε, η συνολική συνεισφορά της διπλωματικής εργασίας, τα ερευνητικά κενά και οι κατευθύνσεις, όπως και η προσωπική εμπειρία που αποκτήθηκε.

2. ΜΗΧΑΝΙΚΗ ΜΑΘΗΣΗ

2.1 Εισαγωγή & Ορισμός

Η Μηχανική Μάθηση (Machine Learning) αποτελεί κατηγορία της τεχνητής νοημοσύνης (AI) στο πεδίο της επιστήμης των υπολογιστών που ασχολείται με τη δημιουργία συστημάτων, τα οποία μπορούν να μαθαίνουν και να βελτιώνονται αυτόματα από εισερχόμενα δεδομένα χωρίς ανθρώπινη παρέμβαση. Κορμός της Μηχανικής Μάθησης (MM) είναι η χρήση αλγορίθμων, οι οποίοι αναλύουν δεδομένα, εντοπίζουν συσχετίσεις και μοτίβα που μπορεί να υπάρχουν, αυτο-βελτιώνονται και χρησιμοποιούν αυτή τη γνώση ώστε να κάνουν προβλέψεις και να αποφασίζουν.

Σκοπός της MM είναι μία μηχανή να έχει τη δυνατότητα να μιμηθεί την ευφυή ανθρώπινη συμπεριφορά. Τα συστήματα τεχνητής νοημοσύνης εκτελούν σύνθετες εργασίες με τρόπο, που προσομοιάζει εκείνον με τον οποίο οι άνθρωποι επιλύουν προβλήματα. Την δεκαετία του 1950, ο αμερικανός επιστήμονας Άρθουρ Σάμιουελ, πρωτοπόρος στην Τεχνητή Νοημοσύνη, όρισε ως Μηχανική Μάθηση την ικανότητα των υπολογιστών να μαθαίνουν χωρίς να είναι ρητά προγραμματισμένοι (Nilsson, 1998).

Η βασική ιδέα της Μηχανικής Μάθησης είναι ότι οι υπολογιστές, οι μηχανές και τα συστήματα εκπαιδεύονται μέσα από τα δεδομένα που συλλέγουν. Τα δεδομένα αυτά μπορεί να είναι φωτογραφίες, βίντεο, ήχοι, κείμενα, νούμερα κ.λπ. Αυτά τα δεδομένα χρησιμοποιούνται είτε ως παραδείγματα, είτε ως δεδομένα εκπαίδευσης (training sets). Όσο περισσότερα δεδομένα χρησιμοποιούνται, τόσο πιο αποδοτικός γίνεται ο αλγόριθμος. Μετά την εισαγωγή των δεδομένων, οι μηχανές αναγνωρίζουν μοτίβα

και σχέσεις που υπάρχουν ανάμεσα στα δεδομένα αυτά και στη συνέχεια λαμβάνουν αποφάσεις και κάνουν προβλέψεις, στηριζόμενες στα παραπάνω ευρήματα. Επομένως, οι αλγόριθμοι μηχανικής μάθησης, εκπαιδεύονται μέσω των δεδομένων, δειγμάτων και των παραδειγμάτων που λαμβάνουν. Ως εκ τούτου, η MM αναφέρεται στην διαδικασία της εκπαίδευσης των μηχανών (Nilsson, 1998; Shai & Shai, 2014).

Αλγόριθμο ονομάζουμε ένα σύνολο από πεπερασμένα βήματα ή οδηγίες που οδηγούν στην εκτέλεση μιας λειτουργίας ή την επίλυση ενός προβλήματος. Λαμβάνει εισόδους, τις επεξεργάζεται μέσω μιας σειράς από λογικά βήματα και έπειτα παράγει την επιθυμητή έξοδο εντός ενός κατάλληλου χρονικού διαστήματος. Ένας αλγόριθμος MM είναι μία δομημένη, υπολογιστική διαδικασία που επιτρέπει σε ένα σύστημα να μαθαίνει μοτίβα, συσχετίσεις ή πληροφορίες από δεδομένα, χωρίς να είναι ρητά προγραμματισμένο. Χρησιμοποιεί μαθηματικές και στατιστικές τεχνικές για την επεξεργασία εισόδου, την προσαρμογή των παραμέτρων και τη βελτίωση της απόδοσης με την πάροδο του χρόνου. Σκοπός του είναι οι προβλέψεις, η ταξινόμηση δεδομένων και ο εντοπισμός κρυφών δομών. Οι αλγόριθμοι MM έχουν τη δυνατότητα να επεξεργάζονται τεράστιους όγκους δεδομένων σε σύντομο χρόνο και να αναγνωρίζουν συγκεκριμένα μοτίβα με τελικό σκοπό την αυτόνομη λειτουργία τους.

Μια κατηγορία αλγόριθμων που αφορούν αυτή την εργασία είναι οι αλγόριθμοι προτάσεων (recommendation algorithms), οι οποίοι κάνουν προτάσεις π.χ. για το ποιος είναι ο καταλληλότερος υποψήφιος για μία θέση πρόσληψης. Μια ακόμα κατηγορία είναι οι αναγνωριστικοί αλγόριθμοι (recognition algorithms), οι οποίοι προσδιορίζουν και επεξεργάζονται εικόνες προσώπου και σώματος. Οι αναγνωριστικοί αλγόριθμοι, είτε εστιάζοντας σε διακριτά σημεία του προσώπου και του σώματος, είτε προσλαμβάνοντάς τα ως ενιαίο σύνολο, αναγνωρίζουν, αναλύουν και ταξινομούν χαρακτηριστικά, όπως το χαμόγελο, η χροιά της φωνής, η κατεύθυνση του βλέμματος, καθώς και συσχετίζουν συναισθήματα και χαρακτηριστικά της προσωπικότητας με τις τιμές που εξάγουν. Η διαδικασία μπορεί να πραγματοποιηθεί σε βίντεο και σε εικόνες, σε πραγματικό χρόνο. Αυτό βοηθά στην άμεση αποτύπωση των χαρακτηριστικών της προσωπικότητας (π.χ. της ομαδικότητας, της ηγεσίας, κ.α.), με βάση τα οποία ο υπεύθυνος θα αποφασίσει αν ένας υποψήφιος είναι ο καταλληλότερος για πρόσληψη.

2.2 Είδη Μηχανικής Μάθησης

Η Μηχανική Μάθηση (MM) χωρίζεται σε τρεις κατηγορίες. Η πρώτη κατηγορία είναι η Επιβλεπόμενη Μάθηση (Supervised Learning), η δεύτερη είναι η Μάθηση χωρίς επίβλεψη (Unsupervised Learning) και η τρίτη είναι η Ενισχυτική Μάθηση (Reinforcement Learning).

2.2.1 Επιβλεπόμενη Μάθηση

Στην Επιβλεπόμενη Μάθηση, η οποία είναι και η πιο διαδεδομένη κατηγορία MM, έχουμε προκαθορισμένα σύνολα δεδομένων εισόδου και εξόδου με ετικέτες. Οι μηχανές εκπαιδεύονται με βάση τα εισερχόμενα δεδομένα ώστε να προβλέπουν το σωστό αποτέλεσμα, αφού αναγνωρίσουν και κατηγοριοποιήσουν τα δεδομένα με τη σωστή ετικέτα. Δύο είναι οι κυρίαρχοι τύποι αλγορίθμων επιβλεπόμενης μάθησης, της ταξινόμησης (classification) και της παλινδρόμησης (regression). Οι αλγόριθμοι ταξινόμησης κάνουν προβλέψεις για δεδομένα με ετικέτα. Ένα απλό παράδειγμα είναι εάν ένα μήνυμα ηλεκτρονικού ταχυδρομείου είναι “spam” ή όχι. Οι πιο διαδεδομένοι αλγόριθμοι ταξινόμησης στην Επιβλεπόμενη Μάθηση είναι οι: Random Forest¹, Decision Tree² και Support Vector Machine (SVM)³.

Οι αλγόριθμοι παλινδρόμησης χρησιμοποιούν και αυτοί δεδομένα με ετικέτα, όμως προβλέπουν συνεχώς μεταβαλλόμενες τιμές. Η πρόβλεψη των τιμών στο Χρηματιστήριο, των τιμών των ακινήτων και η πρόβλεψη του καιρού, αποτελούν χαρακτηριστικά παραδείγματα. Ορισμένοι γνωστοί αλγόριθμοι παλινδρόμησης είναι οι εξής: Ridge Regression⁴, Linear Regression⁵, Logistic Regression⁶, Least Absolute Shrinkage and Selection Operator (LASSO)⁷ και Elastic Net Regression (Dhumne, 2023).

¹ [What Is Random Forest? | IBM](#)

² [What is a Decision Tree? | IBM](#)

³ [What Is Support Vector Machine? | IBM](#)

⁴ [What Is Ridge Regression? | IBM](#)

⁵ [Linear regression in R - IBM Developer](#)

⁶ [What Is Logistic Regression? | IBM](#)

⁷ [What is lasso regression? | IBM](#)

Ορισμένοι αλγόριθμοι ταξινόμησης Επιβλεπόμενης Μάθησης που αναφέρονται στην παρούσα διπλωματική εργασία χρησιμοποιούνται για την ανάλυση συναισθήματος και την αναγνώριση εικόνων και ομιλίας. Αναγνωρίζουν τις γραπτές λέξεις και φράσεις, τις κινήσεις του σώματος και του προσώπου, την τονικότητα και τον τόνο της φωνής, τις εκφράσεις του προσώπου και τις ταξινομούν σε συναισθήματα που εκφράζουν θετικότητα, αρνητισμό και ουδετερότητα. Οι αλγόριθμοι μπορούν να αποκωδικοποιήσουν το συναίσθημα που κρύβεται πίσω από τις κινήσεις, τις λέξεις και γενικά τις μη λεκτικές εκφράσεις.

2.2.2 Μη Επιβλεπόμενη Μάθηση

Στην Μη Επιβλεπόμενη Μάθηση, τα εισερχόμενα σύνολα δεδομένων δεν έχουν ετικέτες. Ένας αλγόριθμος MM που εμπίπτει σε αυτή την κατηγορία αναζητά τα μοτίβα που προκύπτουν, ομαδοποιεί και συσχετίζει τα δεδομένα με βάση τα μοτίβα αυτά, και τέλος εξάγει συμπεράσματα. Παραδείγματα, αποτελούν οι αλγόριθμοι που ταξινομούν άρθρα στις ορθές ενότητες των ειδήσεων, που μεταφράζουν και ταξινομούν κείμενα και που αναγνωρίζουν τις λέξεις και τα κείμενα μετά από μετατροπή τους από ομιλία. Η επεξεργασία της φυσικής γλώσσας (Natural Language Processing - NLP⁸) σίγουρα αποτελεί μεγάλο κομμάτι στη Μη Επιβλεπόμενη Μάθηση. Επίσης, η Μη Επιβλεπόμενη Μάθηση χρησιμοποιείται ευρέως την τελευταία δεκαετία στην ιατρική (Bajwa et al, 2021). Ειδικότερα, σε αυτή την περίπτωση, οι αλγόριθμοι καλούνται να ανιχνεύσουν τις ανωμαλίες που εμφανίζονται σε εικόνες, όπως οι ακτινογραφίες.

Ορισμένοι τύποι αλγορίθμων μη Επιβλεπόμενης Μάθησης είναι:

- οι αλγόριθμοι ομαδοποίησης (clustering algorithms⁹), όπως ο K-means Clustering¹⁰ και ο DBSCAN¹¹ (Density-Based Spatial Clustering of Applications with Noise),

⁸ [What Is NLP \(Natural Language Processing\)? | IBM](#)

⁹ [10 Clustering Algorithms With Python - MachineLearningMastery.com](#). (Brownlee1, 2020).

¹⁰ [What is k-means clustering? | IBM](#) (Kavlakoglu & Winland, 2024).

¹¹ [Clustering Like a Pro: A Beginner's Guide to DBSCAN | by Sachinsoni | Medium](#). (Sachinsoni, 2023).

- οι αλγόριθμοι μείωσης διαστάσεων (Dimensionality Reduction Algorithms¹²), όπως ο PCA¹³,
- οι αλγόριθμοι συσταδοποίησης (Manifold Learning¹⁴), όπως ο Locally Linear Embedding¹⁵ (LLE)
- και τέλος οι αλγόριθμοι ανίχνευσης ανωμαλιών (Anomaly Detection¹⁶), όπως ο Isolation Forest¹⁷.

Οι αλγόριθμοι προτάσεων (recommendation algorithms), οι οποίοι αναφέρθηκαν προηγουμένως ως σχετικοί με το τρέχων αντικείμενο της εργασίας, αποτελούν χαρακτηριστικό παράδειγμα αυτής της κατηγορίας. Εταιρείες και εφαρμογές, όπως το Netflix, το Youtube και οι ηλεκτρονικές ιστοσελίδες προϊόντων, τους χρησιμοποιούν για προτάσεις προϊόντων, βίντεο και ταινιών.

2.2.3 Ενισχυτική Μάθηση

Στην Ενισχυτική Μάθηση, η μηχανή (πράκτορας) εκπαιδεύεται σε συγκεκριμένο, αλλά δυναμικό περιβάλλον έτσι ώστε να εξετάζει διάφορες εναλλακτικές δράσεις προς εκτέλεση και να επιλέγει τη καταλληλότερη, βασιζόμενη στις προηγούμενες σωστές και λανθασμένες επιλογές της (ανατροφοδότηση - feedback). Ο τρόπος εκμάθησης του συστήματος βασίζεται στην επιβράβευση της σωστής απόφασης και στην αποδοκιμασία της λανθασμένης (reward & punishment), δηλαδή εκπαιδεύεται τόσο από την θετική ανατροφοδότηση που λαμβάνει, όσο και από την αρνητική. Συνεπώς, η εκπαίδευση δεν βασίζεται σε δεδομένα με ετικέτες. Σκοπός της μηχανής είναι πάντοτε η μεγαλύτερη επιβράβευση. Στο πλαίσιο αυτού του είδους μάθησης, οι γνώσεις συνεχώς επικαιροποιούνται.

Δύο είναι οι μέθοδοι που ακολουθεί αυτή η κατηγορία Μηχανικής Μάθησης, η εκμετάλλευση (exploitation) και η εξερεύνηση (exploration). Στην πρώτη μέθοδο, η μηχανή εκμεταλλεύεται την ήδη υπάρχουσα γνώση. Το αποτέλεσμα εκτέλεσης της

¹² [6 Dimensionality Reduction Algorithms With Python - MachineLearningMastery.com](#) (Brownlee2, 2020).

¹³ [Step-By-Step Guide to Principal Component Analysis With Example](#). (Dharani, 2025).

¹⁴ [2.2. Manifold learning — scikit-learn 1.5.2 documentation](#)

¹⁵ <https://medium.com/biased-algorithms/an-introduction-to-locally-linear-embedding-lle-f40b2b703040> (Yadav, 2024).

¹⁶ [8 Anomaly Detection Algorithms to Know | Built In](#). (Kumar, 2024).

¹⁷ <https://medium.com/@corymaklin/isolation-forest-799fcea44a4>. (Maklin, 2022).

καταλληλότερης δράσης (με βάση την υπάρχουσα γνώση) αφού είναι δοκιμασμένο είναι σίγουρα ικανοποιητικό, αλλά μπορεί να υπάρχει καλύτερο. Στη δεύτερη μέθοδο, η μηχανή εξερευνά κάτι νέο. Εμφιλοχωρεί όμως ο κίνδυνος να μην εντοπιστεί δράση που να οδηγεί σε καλύτερο αποτέλεσμα από το υπάρχον. Όμως, αν η μηχανή δεν προβεί σε εξερευνήσεις τέτοιους είδους, δεν θα βελτιωθεί ποτέ. Αυτό είναι το δίλημμα της εκμετάλλευσης και της εξερεύνησης (exploitation vs exploration dilemma) (Brown, 2021; Shai & Shai, 2014).

Παραδείγματα τύπων αλγορίθμων Ενισχυτικής Μάθησης είναι:

- οι αλγόριθμοι βασισμένοι στη συνάρτηση αξίας (Value-Based Algorithms¹⁸), όπως ο Q-Learning¹⁹,
- οι αλγόριθμοι βασισμένοι στην πολιτική (Policy-Based Algorithms²⁰), όπως ο Reinforce²¹,
- οι αλγόριθμοι βασισμένοι σε μοντέλο (Model-Based RL²²), όπως ο Monte Carlo Tree Search²³ (MCTS)
- και οι υβριδικοί²⁴, όπως ο AlphaZero²⁵.

Παραδείγματα αυτής της κατηγορίας της Μηχανικής Μάθησης είναι η αυτόνομη οδήγηση, ο έλεγχος φωτεινών σηματοδοτών για την βελτιστοποίηση της κυκλοφοριακής κίνησης κ.λπ. Αξιοσημείωτη είναι η μελέτη και ο αρχικός πειραματισμός δημιουργίας ενός αυτόνομου εργαστηρίου χημείας, του AlphaFlow. Σε αυτό το εργαστήριο θα διεξάγονται περίπλοκα και πολλαπλά πειράματα, καθοδηγούμενα πλήρως από συστήματα ενισχυτικής μάθησης (Volk et al, 2023).

¹⁸ [Value-Based vs Policy-Based Reinforcement Learning | by Papers in 100 Lines of Code | Nov. 2024 | Medium](#)

¹⁹ <https://ieeexplore.ieee.org/document/8836506> (Jang et al, 2019).

²⁰ [Value-Based vs Policy-Based Reinforcement Learning | by Papers in 100 Lines of Code | Nov. 2024 | Medium](#)

²¹ [REINFORCE algorithm — Reinforcement Learning from scratch in PyTorch | by Konstantin Sofeikov | Medium](#). (Sofeikov, 2023).

²² [Understanding Model-Based Reinforcement Learning | by Rakshit Kalra | Medium](#) (Kalra, 2024).

²³ [Monte Carlo Tree Search: a review of recent modifications and applications | Artificial Intelligence Review](#) (Swiechowski, 2022).

²⁴ [What is Hybrid Machine Learning?](#) (Polymer, 2025).

²⁵ [AlphaZero for dummies!](#). (Michelangelo, 2022).

2.3 Πλεονεκτήματα & Μειονεκτήματα

Η εκπαίδευση μιας μηχανής προκειμένου να αναγνωρίζει φωτογραφίες, κινήσεις, οπτικές σκηνές, να κατανοεί ένα κείμενο γραμμένο σε φυσική γλώσσα ή ακόμα και ένα κείμενο εξαγόμενο από ένα βίντεο, εκτός από χρονοβόρα, μπορεί να καταλήξει και αδύνατη. Ωστόσο, η MM παρέχει σημαντικά πλεονεκτήματα, αφού οι αλγόριθμοι, χρησιμοποιώντας δεδομένα εκπαίδευσης και δοκιμών, μπορούν να εκπαιδεύονται συνεχώς και να βελτιώνονται, επιτυγχάνοντας την μεγαλύτερη δυνατή ακρίβεια στα αποτελέσματά τους. Στην MM, το σύστημα εκπαιδεύεται μέσω της εμπειρίας, όπως ένας ευφυής άνθρωπος μαθαίνει και προσαρμόζεται στα νέα δεδομένα και συνθήκες. Με αυτόν τον τρόπο, οι αλγόριθμοι αυτοματοποιούν επαναλαμβανόμενες διαδικασίες, μειώνοντας τον χρόνο και το κόστος που απαιτεί η ανθρώπινη παρέμβαση. Με κατάλληλη εκπαίδευση των αλγορίθμων, η ανάλυση και η ανάγνωση των βιογραφικών σημειωμάτων όπως και ο εντοπισμός και η ταξινόμηση των κατάλληλων υποψηφίων (σε θέσεις εργασίας), καθιστά όλη τη διαδικασία γρηγορότερη και με μειωμένο το κόστος.

Οι αλγόριθμοι έχουν τη δυνατότητα να αναγνωρίζουν πολύπλοκα πρότυπα, μοτίβα και συσχετίσεις στα δεδομένα, που είναι δύσκολο να αναγνωριστούν από ανθρώπους, και να δίνουν εξατομικευμένα αποτελέσματα, πιο ακριβείς προβλέψεις και λύσεις. Σε περιπτώσεις, όπως η αυτόματη εξαγωγή ικανοτήτων και καταλληλότητας των υποψηφίων, μπορούν να διαχειριστούν τη διαδικασία ανάλυσης μεγάλων όγκων δεδομένων με ακρίβεια και ταχύτητα που υπερβαίνει τις ανθρώπινες δυνατότητες. Η μηχανική μάθηση προσφέρει σημαντικές ευκαιρίες, ιδιαίτερα στον αυτοματισμό σύνθετων διαδικασιών, όπως η ανάλυση βιογραφικών και χαρακτηριστικών προσωπικότητας. Ειδικότερα, εξάγουν πληροφορίες με μεγάλη ακρίβεια, εντοπίζουν τις συσχετίσεις μεταξύ δεξιοτήτων, εμπειρίας, χαρακτηριστικών προσωπικότητας και λοιπών προσόντων και τα συνδέουν με τις απαιτήσεις κάθε θέσης εργασίας.

Παρόλα αυτά, η μηχανική μάθηση αντιμετωπίζει και σοβαρές προκλήσεις. Ένα από τα μεγαλύτερα προβλήματα αποτελεί η προκατάληψη στα δεδομένα εκπαίδευσης. Αν οι εισερχόμενες πληροφορίες είναι υποκειμενικές ή μη αντιπροσωπευτικές, τότε και το αποτέλεσμα θα είναι αναξιόπιστο και λανθασμένο. Τα δεδομένα ενδέχεται να διαφέρουν εφόσον εισαχθούν από άτομα διαφορετικού υπόβαθρου, εμπειρίας,

τρόπου ζωής, κουλτούρας και φύλου. Οπότε, τα προκατειλημμένα ή ακόμα και τα περιορισμένα δεδομένα εκπαίδευσης μπορεί να αποκλείσουν άδικα υποψήφιους (Shai & Shai, 2014). Επιπλέον, ένα σύστημα που έχει εκπαιδευτεί και λειτουργεί ικανοποιητικά σε ένα σύνολο δεδομένων, ενδέχεται να μην αποδίδει το ίδιο καλά με άλλα σύνολα δεδομένων. Άλλα ζητήματα όπως η ανασφάλεια και το απόρρητο των δεδομένων που περιλαμβάνονται στα βιογραφικά σημειώματα, εγείρουν ανησυχίες για την προστασία των ευαίσθητων προσωπικών δεδομένων. Πρέπει να διασφαλίζεται η προστασία των δεδομένων από αναρμόδιους και κακόβουλους χρήστες/άτομα. Επίσης, η εξάρτηση από την ανθρώπινη εποπτεία για τη δημιουργία, εκπαίδευση και διαχείριση των συστημάτων MM αποτελεί ένα μειονέκτημα καθώς εισάγει προκαταλήψεις και μπορεί να υποπέσει σε λάθη. Η Μηχανική Μάθηση ήδη συνεισφέρει στην ικανότητα αυτοματοποίησης της διαδικασίας πρόσληψης, όμως προκειμένου να αξιοποιηθεί στο έπακρο, απαιτείται μεγάλη προσοχή από την αρχή της διαδικασίας έως το τέλος της.

2.4 Εφαρμογές Μηχανικής Μάθησης

2.4.1 Γενικές Εφαρμογές

Οι λειτουργίες της Μηχανικής Μάθησης είναι κυρίως τρεις. Η πρώτη είναι η περιγραφική (Descriptive). Η μηχανή χρησιμοποιεί τα δεδομένα για να εντοπίσει συσχετίσεις και μοτίβα ώστε να εξηγήσει και να μας βοηθήσει να κατανοήσουμε τί συμβαίνει. Ένα παράδειγμα αποτελεί η ανάλυση των τάσεων του κοινού π.χ. σε πλατφόρμες πωλήσεων και ταινιών. Κατά κανόνα αυτή η λειτουργία ανήκει στην κατηγορία της Μη Επιβλεπόμενης Μάθησης και η χαρακτηριστικότερη μέθοδός της είναι η ομαδοποίηση (clustering).

Η δεύτερη λειτουργία είναι η ρυθμιστική (Prescriptive) που χρησιμοποιεί τα δεδομένα και προγνωστικά μοντέλα, για να προτείνει τις βέλτιστες ενέργειες που δύνανται να πραγματοποιηθούν. Οι αλγόριθμοι πρότασης (recommendation) είναι παραδείγματα σχετικών αλγορίθμων που χρησιμοποιούνται για εξατομικευμένες προτάσεις π.χ. σε πλατφόρμες ψυχαγωγίας και πωλήσεων. Επίσης, ο επιχειρησιακός τομέας κάνει χρήση της συγκεκριμένης λειτουργίας της MM για τη διαχείριση της εφοδιαστικής αλυσίδας. Οι προτάσεις βελτιστοποιούν και διευκολύνουν διεργασίες, όπως η αγορά,

διαχείριση, αποθήκευση, μεταφορά και διανομή προϊόντων, συνυπολογίζοντας την προσφορά, τη ζήτηση αλλά και το κόστος όλων αυτών. Ο χρηματοοικονομικός τομέας επίσης επωφελείται από τις προτάσεις για επενδύσεις κ.ο.κ. Συνήθως η λειτουργία αυτή εμπίπτει στην Ενισχυτική Μάθηση.

Η τελευταία λειτουργία της Μηχανικής Μάθησης είναι η προγνωστική (Predictive), κατά την οποία η μηχανή χρησιμοποιεί τα δεδομένα για να προβλέψει τι πρόκειται να συμβεί. Πάλι χρησιμοποιώντας μοτίβα, αλλά και βασιζόμενη σε τεκταινόμενα, προβλέπει αποτελέσματα. Η πρόβλεψη πωλήσεων και ασφαλιστικών κινδύνων των ασφαλιστικών εταιρειών, είναι παραδείγματα αυτών των αλγορίθμων. Χαρακτηριστικοί αλγόριθμοι αυτής της λειτουργίας αποτελούν αυτοί της ταξινόμησης και της παλινδρόμησης, και κατατάσσονται στην Επιβλεπόμενη Μάθηση. Μια προγνωστική εφαρμογή που αναφέρεται στην παρούσα εργασία είναι η ανάλυση συναισθημάτων ή αλλιώς εξόρυξη γνώμης (sentiment analysis/opinion mining²⁶). Ανάλυση συναισθημάτων ορίζεται ως η τεχνική διαδικασία ανάλυσης μεγάλου όγκου κειμένου στην επεξεργασία φυσικής γλώσσας (NLP), ώστε να εξαχθούν και να ποσοτικοποιηθούν οι πληροφορίες από τα δεδομένα κειμένου και να καθοριστεί εάν εκφράζουν θετικό, αρνητικό ή ουδέτερο συναίσθημα. Ορισμένοι αλγόριθμοι ανάλυσης συναισθημάτων είναι οι: Support Vector Machine, Naive Bayes²⁷ και Logistic Regression. Οι αλγόριθμοι προβλέπουν το συναίσθημα, συνήθως ως θετικό, αρνητικό ή ουδέτερο, από βίντεο, εικόνες και κείμενο.

Τονίζεται πως υπάρχουν προβλήματα που απαιτούν τον συνδυασμό δύο λειτουργιών. Σε αυτές τις περιπτώσεις χρησιμοποιούνται αλγόριθμοι από διαφορετικές κατηγορίες, όπως η χρήση μεθόδων clustering (Μη Επιβλεπόμενη Μάθηση) για την εξαγωγή χαρακτηριστικών που έπειτα χρησιμοποιούνται στην Επιβλεπόμενη Μάθηση για την διενέργεια προβλεπτικών λειτουργιών.

2.4.2 Εφαρμογές στη Διαδικασία Πρόσληψης Υποψηφίων

Με την έλευση της Μηχανικής Μάθησης, μια μηχανή εκπαιδεύεται ολοένα και με μεγαλύτερη ακρίβεια για να προσδιορίσει και να αναλύσει τον συναισθηματικό τόνο

²⁶ [What Is Sentiment Analysis? | IBM](#)

²⁷ [What Are Naïve Bayes Classifiers? | IBM](#)

των λέξεων. Οι άνθρωποι είναι ικανοί να αναγνωρίζουν μέσω των κινήσεων του σώματος, του τόνου της φωνής, των εκφράσεων προσώπου κ.ο.κ., όπως και διαβάζοντας ένα κείμενο, το συναίσθημα που κρύβεται πίσω από τις μη λεκτικές εκφράσεις. Ως εκ τούτου, η MM μπορεί να χρησιμοποιηθεί από ιδιωτικές επιχειρήσεις, αλλά και φορείς του Δημοσίου για την πρόσληψη υποψηφίων. Η χροιά της φωνής, του νοήματος ενός κειμένου και των κινήσεων ενός ατόμου μέσα από ένα βίντεο, εικόνα ή γραπτό κείμενο επιτρέπει με τη βοήθεια της MM να αποκωδικοποιήσει και να ταξινομήσει τα συναισθήματα ώστε να εξαχθούν τα χαρακτηριστικά της προσωπικότητας του υποψηφίου (Brown, 2021; Shai & Shai, 2014).

Ένα ακόμα πεδίο εφαρμογής της MM που αναφέρεται στην παρούσα εργασία είναι η Επεξεργασία φυσικής γλώσσας (Natural Language Processing - NLP). Η Επεξεργασία φυσικής γλώσσας συνδέεται άμεσα με την ανάλυση συναισθημάτων καθώς είναι μία από τις εφαρμογές της. Η Επεξεργασία φυσικής γλώσσας συνδέεται με την Επιβλεπόμενη και την Μη Επιβλεπόμενη Μάθηση, ανάλογα με την εφαρμογή της. Οι μηχανές εκπαιδεύονται ώστε να κατανοούν και να ανταποκρίνονται στην φυσική γλώσσα, όπως ακριβώς γράφεται και ομιλείται από τους ανθρώπους. Ένα συχνό παράδειγμα είναι το chatbox και το (Chat)GPT (Brown, 2021; Shai & Shai, 2014).

Αξίζει επίσης να αναφερθούν τα νευρωνικά δίκτυα (Neural Networks), τα οποία είναι αλγόριθμοι μηχανικής μάθησης που μαθαίνουν στους υπολογιστές πώς να επεξεργάζονται τα δεδομένα σύμφωνα με την αρχιτεκτονική και την λειτουργία ενός ανθρώπινου εγκεφάλου. Χρησιμοποιούνται ευρέως στην Επεξεργασία φυσικής γλώσσας και είναι ιδιαίτερα χρήσιμα στην ανάλυση συναισθήματος. Παραδείγματα νευρωνικών δικτύων για τη μετάφραση κειμένων είναι τα Recurrent Neural Networks (RNN) και τα Transformers (μετασχηματιστές), όπως το BERT²⁸ (Bidirectional Encoder Representations from Transformers), το οποίο είναι ένα μοντέλο για την κατανόηση της φυσικής γλώσσας. Ακόμα ένα νευρωνικό δίκτυο που αναφέρεται στην σχετική βιβλιογραφία με το αντικείμενο της παρούσας εργασίας είναι το Named Entity Recognition²⁹ (NER), το οποίο χρησιμοποιείται στην αναγνώριση οντοτήτων.

²⁸ [BERT — Bidirectional Encoder Representations from Transformer | by Gayathri siva | Medium](#) (Siva, 2021).

²⁹ [What Is Named Entity Recognition \(NER\)? | Definition from TechTarget](#)

Το NER αναγνωρίζει και κατηγοριοποιεί τις οντότητες που βρίσκονται στα κείμενα, ώστε να παραμείνουν οι κρίσιμες πληροφορίες για την εκπαίδευση των αλγορίθμων. Αναγνωρίζει τα άτομα, τις τοποθεσίες, τα προϊόντα, τους οργανισμούς, τα ουσιαστικά, τα επίθετα, τα ρήματα κ.λπ. (Brown, 2021; Shai & Shai, 2014).

Στις περιπτώσεις που εξετάζονται στο κεφάλαιο της ανάλυσης της επιλεγμένης βιβλιογραφίας, αλγόριθμοι MM εκπαιδεύονται να διαχειρίζονται τον τεράστιο και πολύπλοκο όγκο δεδομένων αυτόματα. Οι αρμόδιοι, χωρίς τη βοήθεια της MM, θα σπαταλούσαν πολλές εργατοώρες, με κίνδυνο να αγνοήσουν χρήσιμες πληροφορίες, για να διαβάσουν, να αναλύσουν και να κατηγοριοποιήσουν εκατοντάδες, αν όχι χιλιάδες σε μεγάλες εταιρείες, βιογραφικά. Με τη βοήθεια της MM, η κατηγοριοποίηση και ταξινόμηση των βιογραφικών σημειωμάτων και ο συσχετισμός των καταλληλότερων βιογραφικών με συγκεκριμένη θέση, πραγματοποιείται αυτόματα. Το αποτέλεσμα είναι οι εταιρείες να απασχολούν μειωμένους ανθρώπινους πόρους για τον εντοπισμό των κατάλληλων υπαλλήλων και να εξοικονομούν χρήματα με την αυτοματοποίηση της επεξεργασίας του τεράστιου όγκου δεδομένων. Επίσης, με την ελαχιστοποίηση των λαθών και την μεγαλύτερη ακρίβεια στην αξιολόγηση των βιογραφικών, οι εταιρείες ενδέχεται να αποφύγουν τυχόν απολύσεις και επαναξιολογήσεις, λόγω λανθασμένης επιλογής υποψηφίου (Shai & Shai, 2014).

3. ΔΙΑΔΙΚΑΣΙΑ ΠΡΟΣΛΗΨΗΣ

Η εφαρμογή της Τεχνητής Νοημοσύνης, τα τελευταία χρόνια διαρκώς αυξάνεται και αναπτύσσεται στις υπηρεσίες του Δημόσιου Τομέα, κυρίως όμως στις ιδιωτικές επιχειρήσεις (Langer et al, 2020). Η Μηχανική Μάθηση χρησιμοποιείται ευρέως σε όλους τους επαγγελματικούς κλάδους και πολλοί την ταυτίζουν με την Τεχνητή Νοημοσύνη, καθώς η πλειονότητα των τεχνολογικών εξελίξεων αφορά την τελευταία. Μια έρευνα της αμερικανικής πολυεθνικής εταιρείας Deloitte, εταιρείας παροχής επαγγελματικών υπηρεσιών, που διεξήχθη το 2020, διαπίστωσε ότι το 67% των εταιρειών ήδη χρησιμοποιούσε την Μηχανική Μάθηση και την επόμενη χρονιά θα χρησιμοποιούσε ή σχεδίαζε να την χρησιμοποιήσει το 97% (Brown, 2021). Η

Μηχανική Μάθηση, όχι μόνο είναι πιο γρήγορη στην λειτουργία της, αλλά και πιο οικονομική, όπως προαναφέρθηκε.

Η περίπτωση που εξετάζει η παρούσα διπλωματική εργασία αφορά στην αυτόματη εξαγωγή συμπερασμάτων χαρακτηριστικών και ικανοτήτων των υποψηφίων για την αυτοματοποίηση της διαδικασίας πρόσληψης. Αυτή η βελτιωμένη διαδικασία θα διευκολύνει τους υπαλλήλους που ασχολούνται με την ανίχνευση των κατάλληλων βιογραφικών, καθώς η πλειοψηφία τους θεωρεί ότι η διαδικασία αυτή, είναι το δυσκολότερο κομμάτι (Suhas & Manjunath, 2020). Στη συνέχεια, αναλύεται συντόμως η παραδοσιακή διαδικασία πρόσληψης και οι ιδιαιτερότητές της στο Δημόσιο Τομέα. Έπειτα, προσδιορίζεται ο τρόπος που αυτή η διαδικασία μπορεί να βελτιωθεί και να αυτοματοποιηθεί μέσω της χρήσης της MM.

3.1 Παραδοσιακή Διαδικασία Πρόσληψης

Η Μηχανική Μάθηση χρησιμοποιείται από εταιρείες κολοσσούς, όπως το Netflix και η μηχανή αναζήτησης του Google. Γιατί λοιπόν να μην ενσωματωθεί και στις διαδικασίες πρόσληψης από επιχειρήσεις και τον δημόσιο τομέα; Το τμήμα Ανθρώπινου Δυναμικού (Human Resources, HR) επωμίζεται ως επί των πλείστον τις διαδικασίες πρόσληψης των υποψηφίων. Η αποκωδικοποίηση, ανάλυση και κατηγοριοποίηση των βιογραφικών είναι δραστηριότητες χρονοβόρες και δαπανηρές. Μετά την κατάθεση των αιτήσεων των υποψηφίων, ηλεκτρονικά ή εγγράφως, συνήθως οι αρμόδιοι υπάλληλοι κάνουν τη πρώτη διαλογή σύμφωνα με τα τυπικά τους προσόντα. Έπειτα, οι υπεύθυνοι προσλήψεων επεξεργάζονται τα επιλεγμένα βιογραφικά, είτε ηλεκτρονικά είτε χειρόγραφα. Ύστερα επιλέγουν ποια από αυτά είναι κατάλληλα για να προχωρήσουν σε συνεντεύξεις, είτε διαδικτυακά είτε δια ζώσης.

Οι συνεντεύξεις πραγματοποιούνται κυρίως για να αξιολογήσουν τα χαρακτηριστικά της προσωπικότητας του υποψηφίου, δηλαδή τα ουσιαστικά προσόντα. Κατά κανόνα, ο υπεύθυνος πρόσληψης αξιολογεί, εκτός από τα τυπικά προσόντα, τα χαρακτηριστικά της προσωπικότητας, την ευγλωττία, τις ικανότητες, τη συμπεριφορά και τις δεξιότητες του υποψηφίου προκειμένου να προσλάβει το καταλληλότερο άτομο για τη συγκεκριμένη θέση. Τα τυπικά προσόντα, η λεκτική, αλλά και η μη λεκτική επικοινωνία αξιολογούνται ισοβαρώς διότι όλα αυτά τα χαρακτηριστικά είναι

εξίσου σημαντικά. Ειδικότερα, οι κινήσεις του σώματος, η προσωπικότητα, οι εκφράσεις του προσώπου, ο τρόπος ομιλίας, οι χειρονομίες κ.λπ. παίζουν εξίσου ρόλο στην πρόσληψη ενός εργαζομένου.

Τελευταία, οι συνεντεύξεις δύνανται να πραγματοποιούνται ασύγχρονα. Ο ενδιαφερόμενος συνδέεται σε μία πλατφόρμα στο διαδίκτυο και απαντά σε προκαθορισμένες ερωτήσεις καθώς βιντεοσκοπεί τον εαυτό του με κάμερα, smartphone ή τάμπλετ. Αργότερα, ο υπεύθυνος προσλήψεων συνδέεται στην ίδια πλατφόρμα, αξιολογεί τις απαντήσεις και αποφασίζει αν θα προχωρήσει σε διαζώσης συνέντευξη (Hemamou et al, 2019). Εν κατακλείδι, τα τυπικά και τα ουσιαστικά προσόντα αξιολογούνται ταυτόχρονα από τους αξιολογητές. Αυτό βέβαια συνεχώς αλλάζει καθώς όλο και περισσότερες επιχειρήσεις, στο εξωτερικό κυρίως, βασίζονται στη Μηχανική Μάθηση για τις προσλήψεις προσωπικού (Anju Raj & Farzana, 2023; Shai & Shai, 2014).

3.2 Ιδιαιτερότητες Διαδικασίας Πρόσληψης στο Δημόσιο Τομέα – Α.Σ.Ε.Π.

Η παραπάνω διαδικασία αφορά κυρίως τον ιδιωτικό τομέα, καθώς οι προσλήψεις στον δημόσιο τομέα ακολουθούν διαφορετική πορεία, αυτή του γραπτού διαγωνισμού του Ανώτατου Συμβουλίου Επιλογής Προσωπικού (Α.Σ.Ε.Π.). Η πλειοψηφία των προσλήψεων πραγματοποιείται με γραπτό διαγωνισμό. Υπάρχουν όμως περιπτώσεις που πραγματοποιούνται μέσω προκήρυξης θέσης και συνήθως περιλαμβάνουν τη διενέργεια συνεντεύξεων.

Το Α.Σ.Ε.Π. ιδρύθηκε το 1994 με τον νόμο 2190/1994, ο οποίος επικαιροποιήθηκε με τον νόμο 4765/2021. Ο λόγος που δημιουργήθηκε ο θεσμός αυτός, ήταν η επιτακτική ανάγκη για διαφάνεια στις προσλήψεις προσωπικού στο Δημόσιο. Όλοι οι πολίτες έχουν δικαίωμα για διαφάνεια, αντικειμενικότητα, αμεροληψία και για αξιοκρατικά κριτήρια αξιολόγησης σε ένα σύστημα πρόσληψης και στελέχωσης της δημόσιας διοίκησης. Το Α.Σ.Ε.Π., ως ανεξάρτητη αρχή, είναι ο θεσμικός εγγυητής για τον έλεγχο της εφαρμογής των νόμων και των διατάξεων αυτών. Ως ανεξάρτητη αρχή, το Α.Σ.Ε.Π. δεν εποπτεύεται και δεν ελέγχεται από την κυβέρνηση ή οποιοσδήποτε διοικητικές αρχές, παρά μόνο υποβάλλεται σε κοινοβουλευτικό έλεγχο ενώ οι πράξεις

του υπόκεινται μόνο σε δικαστικό έλεγχο. Επίσης, απολαμβάνει διοικητική και οικονομική αυτοτέλεια. Τα μέλη του Α.Σ.Ε.Π. είναι ανώτατοι κρατικοί λειτουργοί, με τις αρμοδιότητές τους να περιορίζονται σε κανονιστικές, ελεγκτικές και κυρωτικές (ΑΣΕΠ 1, 2024).

Η επιλογή τακτικού, ή μη, προσωπικού του δημόσιου και ευρύτερου δημόσιου τομέα, ο έλεγχος διαδικασιών πρόσληψης, η διεξαγωγή διαδικασιών επιλογής θέσης ευθύνης, όπως και ο έλεγχος τήρησης των διατάξεων περί πρόσληψης, αποτελούν τις κυριότερες αρμοδιότητες του θεσμού αυτού (ΑΣΕΠ 1, 2024). Ανάλογα με τον κλάδο/ειδικότητα και την κατηγορία εκπαίδευσης που ανήκει ο καθένας, μπορεί να επισκεφθεί την ιστοσελίδα του Α.Σ.Ε.Π.³⁰ και να υπολογίσει πώς μοριοδοτούνται τα τυπικά προσόντα του, όπως οι τίτλοι σπουδών, η εμπειρία/προϋπηρεσία, οι ξένες γλώσσες και η εντοπιότητα (ΑΣΕΠ 2, 2024).

Δεν πραγματοποιούνται όμως όλες οι προσλήψεις κατά τον προαναφερόμενο τρόπο στο δημόσιο τομέα. Υπάρχουν περιπτώσεις προσλήψεων με διαγωνισμό προκαθορισμένων κριτηρίων (παραδείγμα η προκήρυξη πρόσληψης ΑΔΑ: 6ΜΚΠ46ΜΤΛΚ-Η5Τ, 2022³¹) ή στελέχωσης θέσης ευθύνης. Η στελέχωση θέσεων ευθύνης διεξάγεται είτε με διαγωνισμό προκαθορισμένων κριτηρίων, είτε με κρίσεις, στις οποίες αξιολογείται ο υποψήφιος για τα τυπικά και ουσιαστικά προσόντα. Στους διαγωνισμούς προκαθορισμένων κριτηρίων οι συμμετέχοντες κρίνονται βάσει συγκεκριμένων και σαφώς προκαθορισμένων κριτηρίων. Κάθε προκήρυξη πρόσληψης έχει διαφορετικά κριτήρια πρόσληψης, ανάλογα με τη θέση στελέχωσης. Τα τυπικά προσόντα βρίσκονται στην αίτηση που καταθέτει ο υπάλληλος. Τα ουσιαστικά προσόντα αξιολογούνται κατά τη διάρκεια της συνέντευξης. Συνήθως, οι διαδικασίες και τα μόρια στις κρίσεις Διευθυντών και Προϊσταμένων του δημόσιου τομέα καθορίζονται από συγκεκριμένο νόμο και τις διατάξεις του, οι δε οδηγίες εφαρμογής παρέχονται με εγκυκλίους. Παραδείγματα των παραπάνω δύο περιπτώσεων αποτελούν κυρίως οι ειδικές θέσεις, όπως αυτές του ειδικού επιστημονικού προσωπικού και οι θέσεις με θητεία. Εξαίρεση του γραπτού διαγωνισμού επίσης προβλέπεται για το ιατρικό προσωπικό, τους δικαστικούς

³⁰ [Πώς μοριοδοτούνται τα προσόντα | ΑΣΕΠ Ενημερωτική Πύλη](#)

³¹

<https://ypergasias.gov.gr/wp-content/uploads/2022/07/6%CE%9C%CE%9A%CE%A046%CE%9C%CE%A4%CE%9B%CE%9A-%CE%975%CE%A4-%CE%A0%CE%A1%CE%9F%CE%9A%CE%97%CE%A1%CE%A5%CE%9E%CE%97-%CE%93%CE%94-%CE%94%CE%A5%CE%A0%CE%91.pdf>

λειτουργούς, το ερευνητικό προσωπικό, το εκπαιδευτικό και διδακτικό προσωπικό, τους στρατιωτικούς, τους υπαλλήλους του διπλωματικού κλάδου, τους υπαλλήλους της Ειδικής Νομικής Υπηρεσίας, τους ειδικούς συμβούλους και συνεργάτες, τους υπαλλήλους απευθείας διευθυντικού βαθμού κ.ο.κ. (Ν.3614/2022; ΑΣΕΠ 1, 2024).

Όπως γίνεται αντιληπτό, η διαδικασία για τις παραπάνω (ειδικές) περιπτώσεις θα μπορούσε να πραγματοποιείται με τη βοήθεια της μηχανικής μάθησης. Στις περιπτώσεις εξαιρέσεων, προβλέπεται ατομική συνέντευξη μετά την θετική αξιολόγηση του βιογραφικού σημειώματος. Τη δεδομένη στιγμή, υπάρχει στο Α.Σ.Ε.Π. ένα Πληροφοριακό Σύστημα για τη διαχείριση των ηλεκτρονικών αιτήσεων και των δικαιολογητικών. Όμως η εισαγωγή της μηχανικής μάθησης σε αυτές τις διαδικασίες, θα το έκανε γρηγορότερο και αποτελεσματικότερο, όχι μόνο στην παραλαβή των αιτήσεων, αλλά και στις διαδικασίες πρόσληψης (Ν.3614/2022; ΑΣΕΠ1, 2024).

Όταν οι αιτήσεις είναι ελάχιστες, είναι εύκολο να επεξεργαστούν, αν όμως ο αριθμός των αιτήσεων είναι τεράστιος είναι σχεδόν βέβαιο ότι θα υπάρξουν λάθη. Κάποια λάθη ενδέχεται να γίνονται στην ταξινόμηση και αξιολόγηση των προσόντων. Για παράδειγμα, στον ανθρώπινο παράγοντα είναι πιθανό να διαφύγει κάποιο σημαντικό στοιχείο. Η ανθρώπινη προκατάληψη είναι επίσης βασικός αρνητικός παράγοντας λανθασμένης αξιολόγησης. Οι προσωπικές εντυπώσεις και οι υποσυνείδητες προκαταλήψεις είναι κάτι που ελαχιστοποιείται, αν όχι αποφεύγεται, από την αντικειμενικότητα των συστημάτων. Περαιτέρω, στα συστήματα δεν επέρχεται κόπωση όπως στους ανθρώπους, ούτε μειώνεται η απόδοσή τους εξαιτίας αυτής. Γενικά, τα συστήματα είναι αποτελεσματικά, αποδοτικά, γρήγορα και όσο το δυνατόν με λιγότερες προκαταλήψεις, όμως αυτό δυστυχώς εξαρτάται από τα εισερχόμενα δεδομένα εκπαίδευσης.

Έγινε αναφορά νωρίτερα στο ότι συνεκτιμώνται τα τυπικά και τα ουσιαστικά προσόντα. Τα τυπικά προσόντα είναι οι τίτλοι σπουδών, οι ξένες γλώσσες, η προϋπηρεσία κ.λπ. Τα ουσιαστικά προσόντα είναι πιο περίπλοκα και διαχωρίζονται σε τέσσερις κατηγορίες, στα σωματικά, τα πνευματικά, τα διοικητικά και στα επαγγελματικά. Η πρώτη κατηγορία αφορά την υγεία, την αρτιμέλεια και την αντοχή στην καταβολή προσπαθειών. Η δεύτερη αφορά τα πνευματικά προσόντα, όπως η νοημοσύνη, η κρίση, η αντίληψη, η προσαρμοστικότητα, η δημιουργικότητα, η

πρωτοτυπία, και η δύναμη της έκφρασης. Στην τρίτη (διοικητικά) περιλαμβάνονται, το κύρος, η ομαδικότητα, η ηγετικότητα, η ενεργητικότητα, η προβλεπτικότητα και η ικανότητα οργάνωσης και συντονισμού. Στην τέταρτη κατηγορία (επαγγελματικά) εντάσσονται, η ικανότητα αφομοίωσης νέων εξελίξεων, η κατάρτιση για περαιτέρω επαγγελματική εξέλιξη, η μεθοδικότητα, η συναδελφικότητα και η αλληλεγγύη (Π.Δ. 318/1992 (Α' 161) και τροποποίηση αυτού (ΑΔΑ: 71Ψ6Χ – Μ5Λ, 2014).

Τα ουσιαστικά προσόντα είναι ποιοτικά και όχι ποσοτικά. Οι υποψήφιοι του δημόσιου τομέα που συμμετέχουν σε γραπτό διαγωνισμό, δεν εξετάζονται για τα ουσιαστικά τους προσόντα. Ο επιτυχής γραπτός διαγωνισμός αρκεί για την πρόσληψη. Στις ειδικές θέσεις του δημόσιου τομέα, αξιολογούνται τα τυπικά αλλά και τα ουσιαστικά προσόντα. Στον ιδιωτικό τομέα, αφού αξιολογηθεί ένας υποψήφιος για τα τυπικά προσόντα του, σε δεύτερη φάση περνάει από συνέντευξη για να αξιολογηθεί και για τα ουσιαστικά. Επομένως, η διαδικασία πρόσληψης μπορεί να θεωρηθεί παρόμοια στον ιδιωτικό τομέα και στο δημόσιο μόνο όσον αφορά τις ειδικές θέσεις. Στον ιδιωτικό τομέα υπάρχει ευελιξία για την αξιολόγηση των υποψήφιων εργαζόμενων ενώ στον δημόσιο τομέα όλα καθορίζονται από το νομικό πλαίσιο που διέπει τη χώρα.

3.3 Βελτιώσεις Διαδικασίας Πρόσληψης μέσω Μηχανικής Μάθησης

Με βάση όλα τα παραπάνω, η χρήση της ΜΜ μπορεί αυτόματα να επεξεργαστεί, ταξινομήσει, κατηγοριοποιήσει και να αντιστοιχίσει το κατάλληλο άτομο με συγκεκριμένη θέση εργασίας. Η ανίχνευση των προσόντων, τυπικών ή ουσιαστικών, πραγματοποιείται αυτόματα με τους κατάλληλους αλγορίθμους, από τα βιογραφικά σημειώματα και τα βίντεο, μετά την συλλογή των απαραίτητων στοιχείων από τις αντίστοιχες εφαρμογές εξαγωγής. Αποτέλεσμα αυτών είναι ένας υποψήφιος να θεωρείται κατάλληλος ή μη για μία θέση εργασίας. Ταυτόχρονα αποτιμάται εάν είναι και ο καλύτερος υποψήφιος. Οι ικανότητες πρόβλεψης, η ταχύτητα λήψης αποφάσεων των αλγορίθμων και η μείωση της απαιτούμενης χειρωνακτικής εργασίας από τους αρμόδιους εργαζόμενους, μετατρέπουν την Τεχνητή Νοημοσύνη και την Μηχανική Μάθηση ειδικότερα σε ένα απαραίτητο εργαλείο για τη σημερινή κοινωνία σε διάφορους τομείς, όπως αυτός που εξετάζεται εδώ.

Απώτερος σκοπός είναι, ο άνθρωπος να μην συμμετέχει καθόλου στη διαδικασία πρόσληψης, παρά μόνο στην επιβεβαίωση της τελικής επιλογής. Με τον τρόπο αυτό, μειώνεται ο χρόνος της υποβολής και συλλογής των αιτήσεων και των δικαιολογητικών των υποψηφίων, της καταγραφής των συνεντεύξεων, της αρχικής αξιολόγησης/βαθμολόγησης των υποψηφίων, με συνέπεια το τμήμα ανθρώπινου δυναμικού να ασχολείται μόνο κατά την τελευταία και πιο κρίσιμη φάση της επιλογής εργαζόμενου. Η MM αυτοματοποιεί ακόμα και το βήμα της συμπλήρωσης των τιμών, στις περιπτώσεις που τα κριτήρια συμπληρώνονται χειρωνακτικά από τους αξιολογητές. Επιπλέον, μειώνονται οι προκαταλήψεις και ο υποκειμενισμός στο ελάχιστο, αν όχι μηδενίζονται, αυτό όμως εξαρτάται από την ποιότητα και το περιεχόμενο των δεδομένων που χρησιμοποιούνται για την εκπαίδευση των σχετικών αλγορίθμων MM. Είναι σχεδόν αδύνατο, ακόμα και για τους πλέον έμπειρους αξιολογητές, να μην επηρεαστούν από προκαταλήψεις και να μην υπάρχουν στοιχεία από βιογραφικά σημειώματα που μπορεί να παραβλεφθούν.

Η σύνδεση με πλατφόρμες ευρέσεως εργασίας, όπως το LinkedIn, θα μπορούσε να βοηθήσει στην σκιαγράφηση των ουσιαστικών και των τυπικών προσόντων των υποψηφίων καθώς μία προηγούμενη θέση εργασίας είναι ενδεικτική για την καταλληλότητα ενός ατόμου για μια νέα παρόμοια θέση. Με τον ίδιο τρόπο θα μπορούσαν να συσχετιστούν τα αναγκαία πιστοποιητικά για μια θέση. Ο εντοπισμός τους θα πραγματοποιείται βάσει παρόμοιων παρελθοντικών θέσεων και των απαιτήσεών τους. Με παρόμοιο τρόπο θα μπορούν να χρησιμοποιηθούν και οι συστατικές επιστολές, με τιμές που θα πρέπει να καλύπτουν για μια νέα θέση εργασίας σε σχέση με παρόμοιες παλαιότερες. Επίσης, οι τελευταίες θα μπορούσαν να σκιαγραφούνται με βάση τις επιλογές των ατόμων που πραγματοποιήθηκαν σε προηγούμενες θέσεις εργασίας. Σύμφωνα με τους (Sowjanya et al, 2023), πάνω από 50.000 ιστοσελίδες εύρεσης εργασίας υπάρχουν παγκοσμίως. Αυτό σημαίνει ότι ο εργοδότης όχι μόνο θα μπορεί να ερευνήσει τους υποψηφίους εάν αυτές οι ιστοσελίδες είναι ανοιχτές, αλλά και ότι η MM σε αυτόν τον τομέα θα είναι εξαιρετικά χρήσιμη για τις ιστοσελίδες που αυτός θα μπορεί να επισκεφθεί και να χρησιμοποιήσει. Επίσης, όχι μόνο θα μπορούσε να εξάγεται πληροφορία για τους αιτούντες, αλλά επιπλέον αυτοί θα μπορούσαν να ειδοποιούνται προληπτικά για μία θέση που ταιριάζει με τα προσόντα τους.

4. ΜΕΘΟΔΟΛΟΓΙΑ ΒΙΒΛΙΟΓΡΑΦΙΚΗΣ ΑΝΑΣΚΟΠΗΣΗΣ

Σε αυτή την ενότητα περιγράφεται η μεθοδολογία που ακολουθήθηκε για την βιβλιογραφική ανασκόπηση, ειδικότερα σχετικά με τη χρήση της Μηχανικής Μάθησης για την αυτόματη εξαγωγή δεδομένων από βιογραφικά σημειώματα και τα χαρακτηριστικά προσωπικότητας από συνεντεύξεις, για την τελική κατάταξη των υποψηφίων, καθώς και τα ερευνητικά ερωτήματα που προέκυψαν και απαντήθηκαν από τη μελέτη της εργασίας. Στη συνέχεια θα ακολουθήσει μια πρώτη ποσοτική ανάλυση των αποτελεσμάτων της βιβλιογραφικής ανασκόπησης, κυρίως στατιστικής φύσεως, συνοδευόμενη από συγκεκριμένα γραφήματα.

4.1 Στόχος Ανασκόπησης

Στόχος της ανασκόπησης είναι να διερευνηθούν, να συγκεντρωθούν και να αξιολογηθούν οι υφιστάμενες επιστημονικές προσεγγίσεις και μέθοδοι που αφορούν στο θέμα της παρούσας διπλωματικής εργασίας. Θα εξεταστούν οι ελλείψεις, τα πλεονεκτήματα, οι περιορισμοί, οι προτάσεις για το μέλλον και οι προσεγγίσεις των υπάρχουσών ερευνών. Πιο συγκεκριμένα, θα εξεταστούν οι εφαρμογές και οι αλγόριθμοι Μηχανικής Μάθησης που χρησιμοποιούνται για την αυτόματη ανάλυση των βιογραφικών και των χαρακτηριστικών προσωπικότητας των υποψηφίων εργαζόμενων. Η ανασκόπηση θα επικεντρωθεί στις πιο αποτελεσματικές μεθόδους και εφαρμογές, όπως και στις προκλήσεις των ερευνών και σε μελλοντικές κατευθύνσεις αντιμετώπισής τους.

4.2 Ερευνητικά Ερωτήματα

Τα ερευνητικά ερωτήματα που προέκυψαν και έχουν άμεση σχέση με τον στόχο της ανασκόπησης είναι τα ακόλουθα:

- Ε1: Ποιές προσεγγίσεις προτείνονται για την αυτόματη εξαγωγή χαρακτηριστικών, είτε προσωπικότητας, είτε τυπικών προσόντων των υποψηφίων;

- E2: Ποιές από αυτές τις προσεγγίσεις είναι οι καλύτερες με βάση ποιό κριτήριο;
- E3: Ποιοί αλγόριθμοι Μηχανικής Μάθησης χρησιμοποιούνται και ποιοί επιτυγχάνουν καλύτερα ποσοστά ακρίβειας ως προς την ταξινόμηση των υποψηφίων σε συγκεκριμένη θέση εργασίας;
- E4: Ποιοί αλγόριθμοι Μηχανικής Μάθησης χρησιμοποιούνται και ποιοί φέρουν καλύτερα ποσοστά ακρίβειας ως προς την εξαγωγή των σωστών χαρακτηριστικών/στοιχείων των υποψηφίων;
- E5: Ποιά χαρακτηριστικά της προσωπικότητας θεωρείται ότι συνδέονται άμεσα με την ικανότητα και τις δεξιότητες για την αξιολόγηση ενός υποψηφίου για μια συγκεκριμένη θέση;
- E6: Ποιοί είναι οι περιορισμοί και οι προκλήσεις της τρέχουσας έρευνας;
- E7: Ποιες είναι οι ερευνητικές προτάσεις για τη βελτίωση των υπάρχουσών εργασιών και την αντιμετώπιση των προκλήσεων;

4.3 Τρόπος Αναζήτησης Άρθρων

Η αρχική αναζήτηση άρθρων πραγματοποιήθηκε στο Google Scholar. Ορισμένες λέξεις-κλειδιά που χρησιμοποιήθηκαν υπήρξαν το “automation inference hiring”, “machine learning algorithms in hiring”, “automatic prediction of CV”, “recruitment prediction algorithms”, “hiring prediction by machine learning from CV”, και “machine learning algorithms in resume ranking”. Το μοτίβο που προέκυψε για τον αποδοτικότερο εντοπισμό άρθρων ήταν το “(((Supervised || Unsupervised || Reinforcement) && Learning) || "Neural Network") && ((automatic || automated || automation) && (ranking || prediction || sorting || inference) && (recruitment || CV || resume || hiring)”. Το Google Scholar υπήρξε βασικό μέσο αναζήτησης άρθρων διότι χρησιμοποιείται ευρέως σε βιβλιογραφικές ανασκοπήσεις. Επίσης, χρησιμοποιήθηκαν σε μεγάλο βαθμό τα: IEEE, Semantic Scholar, Research Gate, Science Direct, Springer και Wiley. Οι παραπάνω βάσεις και εκδότες πληρούσαν τα κριτήρια που είχαν τεθεί, δηλαδή είχαν το κύρος, την επιστημονική αξιοπιστία και άρθρα με τις τελευταίες εξελίξεις. Ειδικότερα, είναι γνωστό ότι οι συγκεκριμένοι εκδότες ακολουθούν αυστηρές διαδικασίες αξιολόγησης ώστε να διασφαλίζεται ότι οι δημοσιεύσεις είναι επιστημονικά ακριβείς και αξιόπιστες.

Πρέπει να αναφερθεί ότι εκτός από την αναζήτηση μέσω των προαναφερόμενων βιβλιογραφικών βάσεων, εφαρμόστηκε και η μέθοδος «snowball effect». Το «snowball effect» είναι η διαδικασία αναζήτησης άρθρων, και παραπομπών αυτών, που παραπέμπουν σε νέα άρθρα. Με αυτόν τον τρόπο, το ένα άρθρο οδηγεί σε επόμενα με έναν επαναληπτικό τρόπο (Cranfield University, 2024). Με βάση αυτή τη μέθοδο, καθώς γινόταν η μελέτη ενός επιλεγμένου άρθρου (ως σχετικό με τον στόχο της βιβλιογραφικής ανασκόπησης), εντοπίζονταν ταυτόχρονα μέσω παραπομπών επιπλέον άρθρα, τα οποία αναζητούνταν και εκφορτώνονταν προς έλεγχο και μελέτη.

4.4 Φιλτράρισμα Άρθρων

Υπήρξαν αρκετά άρθρα τα οποία αποκλείστηκαν από την παρούσα διπλωματική εργασία και οι λόγοι ήταν κυρίως έλλειψης συνάφειας, αξιοπιστίας, ποιότητας και η παλαιότητα. Η θεματική συνάφεια με το ερευνητικό αντικείμενο της εργασίας υπήρξε κριτήριο ένταξης ή αποκλεισμού ενός άρθρου. Συνοπτικά, τα κριτήρια που χρησιμοποιήθηκαν για το φιλτράρισμα των αρχικά εντοπισμένων άρθρων διακρίνονται ως εξής:

Κριτήρια αποκλεισμού / αποδοχής

- *Θεματική συνάφεια*: Τα άρθρα έπρεπε να είναι σχετικά με το θέμα της διπλωματικής εργασίας, δηλαδή με την αυτόματη εξαγωγή συμπερασμάτων για την επιλογή και ταξινόμηση των υποψηφίων με τη βοήθεια της MM. Αποκλείστηκαν άρθρα με παραπλήσια θέματα που δεν εξυπηρετούσαν επακριβώς το αντικείμενο της έρευνας.
- *Γλώσσα δημοσίευσης*: Τα άρθρα που αναζητήθηκαν ήταν στην αγγλική γλώσσα, καθώς η γλώσσα αυτή κυριαρχεί σε ένα πολύ μεγάλο ποσοστό των επιστημονικών δημοσιεύσεων, τόσο σε περιοδικά όσο και σε συνέδρια υψηλού κύρους. Στην ελληνική γλώσσα δεν βρέθηκε κάποιο άρθρο που να πληροί τα κριτήρια της διπλωματικής εργασίας.
- *Χρονολογία δημοσίευσης*: Ο τομέας της Μηχανικής Μάθησης είναι ταχέως εξελισσόμενος. Για τον λόγο αυτό, αποκλείστηκαν παλαιότερα άρθρα αφού θεωρούνται ξεπερασμένα. Ειδικότερα, τα άρθρα έπρεπε να έχουν δημοσιευθεί στο χρονικό διάστημα από το 2020 έως και το 2024. Εξαίρεση αποτελούν τα

άρθρα που αφορούν στα ουσιαστικά προσόντα καθώς η χρονολογία αναγκαστικά επεκτάθηκε από το 2016 μέχρι σήμερα. Ειδικότερα, πριν το 2020 βρέθηκαν μοναδικές προσεγγίσεις που έπρεπε να καλυφθούν.

- *Πηγή δημοσίευσης*: Εξετάστηκαν και επιλέχθηκαν άρθρα μόνο από αξιόπιστες πηγές, ειδικότερα από διεθνώς αναγνωρισμένα συνέδρια και περιοδικά έγκριτου κύρους.
- *Έκταση άρθρων*: Άρθρα μικρότερα από 5 δίστηλες σελίδες αποκλείστηκαν λόγω της μικρής τους έκτασης και προφανώς της μη επαρκούς ανάλυσης που παρείχαν.

Κριτήρια ποιότητας

- *Κατάλληλη χρήση της αγγλικής γλώσσας*: Επιλέχθηκαν άρθρα με σωστή χρήση της αγγλικής γλώσσας.
- *Επιστημονική μεθοδολογία*: Προτιμήθηκαν άρθρα με ορθή επιστημονική μεθοδολογία.
- *Επίπεδο κατανόησης*: Επιλέχθηκαν άρθρα με ικανοποιητικό επίπεδο κατανόησης και αφαιρέθηκαν τα δυσνόητα και όπως αυτά που χρειάζονται μεγάλη προσπάθεια για την κατανόησή τους.
- *Συμβολή και καινοτομία*: Επιλέχθηκαν άρθρα με καινοτόμες ιδέες και προσεγγίσεις.

4.5 Ποσοτική Ανάλυση

Ο αρχικός όγκος των άρθρων ήταν περί τα τριπλάσια από την τελική διαλογή, δηλαδή περίπου στα ογδόντα. Ο τελικός αριθμός άρθρων που επιλέχθηκαν είναι δεκαεννέα. Από αυτά τα άρθρα που παρουσιάζονται στην διπλωματική εργασία, τα έντεκα είναι από επιστημονικά περιοδικά και τα οκτώ από συνέδρια. Οι σελίδες των άρθρων κυμαίνονται από πέντε έως και είκοσι στην πλειοψηφία τους δίστηλες.

Παρακάτω παρατίθενται, στην υποενότητα 4.5.1, δύο διαγράμματα πίτας που επιδεικνύουν το ποσοστό των άρθρων ανά είδος και εκδότη, αντίστοιχα. Έπειτα, στην υπο-ενότητα 4.5.2 επιδεικνύεται ένα γράφημα μπάρας που προσδιορίζει τον αριθμό των επιλεγμένων άρθρων ανά έτος δημοσίευσης. Με αυτό τον τρόπο, παρέχουμε

στον αναγνώστη μια πρώτη εικόνα από την ποσοτική ανάλυση των άρθρων που επιλέξαμε από τη βιβλιογραφία.

4.5.1 Διαγράμματα Πίτας

Παρακάτω παρατίθενται διαγράμματα πίτας. Το πρώτο (δείτε Εικόνα 1) δείχνει το ποσοστό των άρθρων ανά είδος: συγκεκριμένα το 40% των άρθρων προέρχονται από συνέδρια και το 60% από επιστημονικά περιοδικά. Συνεπώς, παρατηρούμε πως υπάρχει λίγο-πολύ μια ισορροπία ως προς το είδος των επιλεγμένων άρθρων.

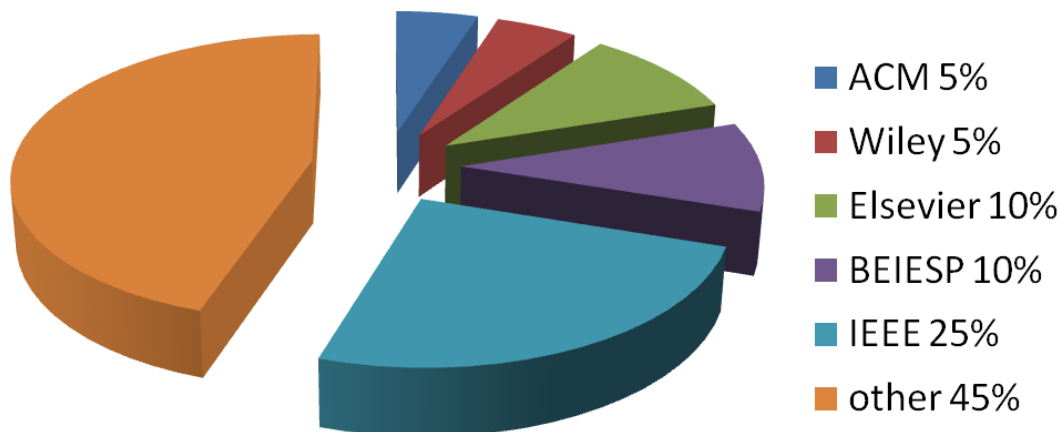


Εικόνα 1 - Γράφημα πίτας για τα ποσοστά άρθρων ανά είδος

Το δεύτερο διάγραμμα πίτας (δείτε Εικόνα 2) δείχνει τα ποσοστά σε σχέση με το από πού αντλήθηκαν τα άρθρα, δηλαδή ποιός οργανισμός τα έχει εκδώσει. Από ότι φαίνεται, το μεγαλύτερο ποσοστό άρθρων σε μεμονωμένο εκδότη απολαμβάνει ο πολύ γνωστός οργανισμός έγκριτου κύρους IEEE και έπειτα ακολουθεί ο Elsevier. Αλλά το μεγαλύτερο ποσοστό συνολικά αντιστοιχεί στην περίπτωση 'other' που καλύπτει πολλούς, όχι τόσο ευρέως γνωστούς εκδότες. Αυτό επιδεικνύει πως τα άρθρα επιλέχθηκαν από μια μεγάλη ποικιλία από διαφορετικούς εκδότες. Όμως, να τονιστεί σε αυτό το σημείο πως η άντληση των άρθρων έγινε υπό αυστηρά κριτήρια

ώστε να διατηρηθούν μόνο τα άρθρα που δημοσιεύτηκαν έπειτα από μια διαδικασία blind-based peer review (αξιολόγησης από ομότιμους).

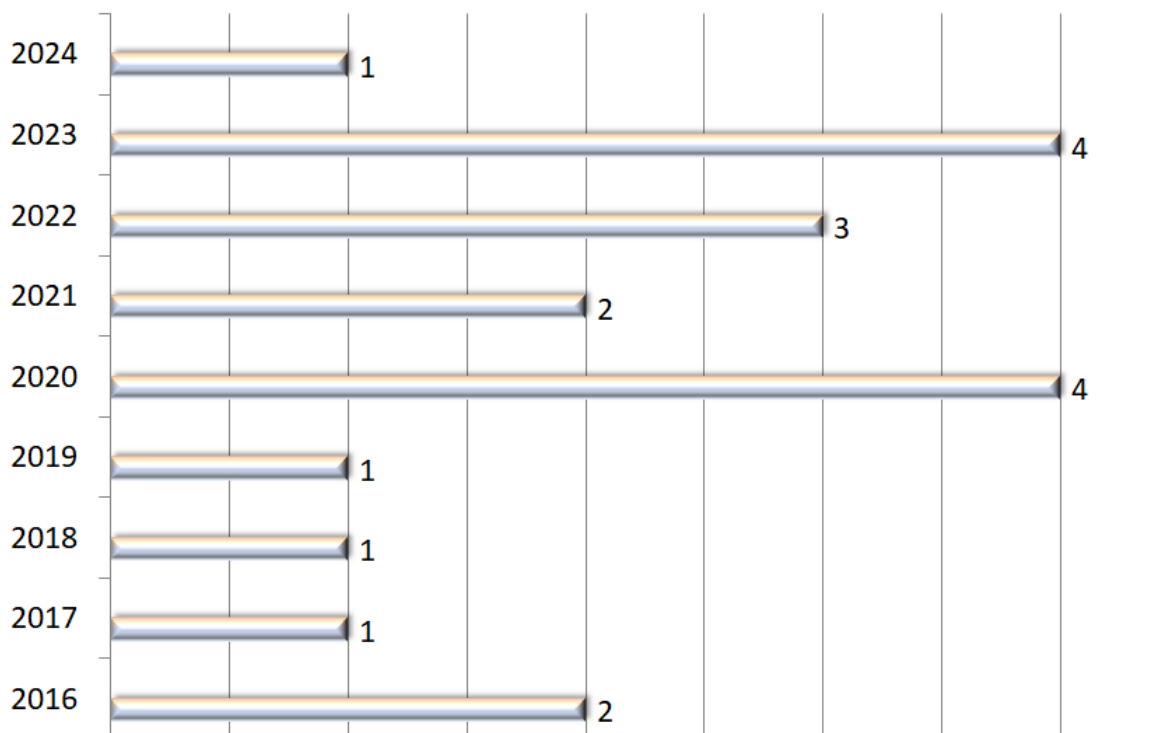
Ποσοστά άρθρων ανά εκδότη



Εικόνα 2 - Γράφημα πίτας για τα ποσοστά άρθρων ανά εκδότη

4.5.2 Γραφήματα

Στην Εικόνα 3 παρουσιάζεται ο αριθμός των άρθρων που εξετάστηκαν ανά έτος δημοσίευσης. Όπως μπορεί να παρατηρηθεί, διακρίνεται μια αυξητική τάση μετά το 2019. Ειδικότερα, πριν το 2020, ο αριθμός των άρθρων είναι ελάχιστος. Ενώ αυξάνεται στο χρονικό διάστημα 2020-2023. Θα πρέπει να ληφθεί υπόψη πως θα πρέπει να εξαιρεθεί από αυτό το συμπέρασμα ο αριθμός των άρθρων για το 2024 διότι η ανασκόπηση πραγματοποιήθηκε εντός του 2024 και επομένως δεν ήταν δυνατό να καλυφθούν τελικώς δημοσιευμένα άρθρα που μπορεί να ήταν ακόμη υπό την διαδικασία αξιολόγησης. Με βάση την παραπάνω δικαιολόγηση, μπορούμε να συμπεράνουμε πως το αντικείμενο της μελέτης μας χαίρει αυξανόμενης ερευνητικής έντασης με όλο και περισσότερα άρθρα να σχετίζονται με αυτό τα τελευταία χρόνια.



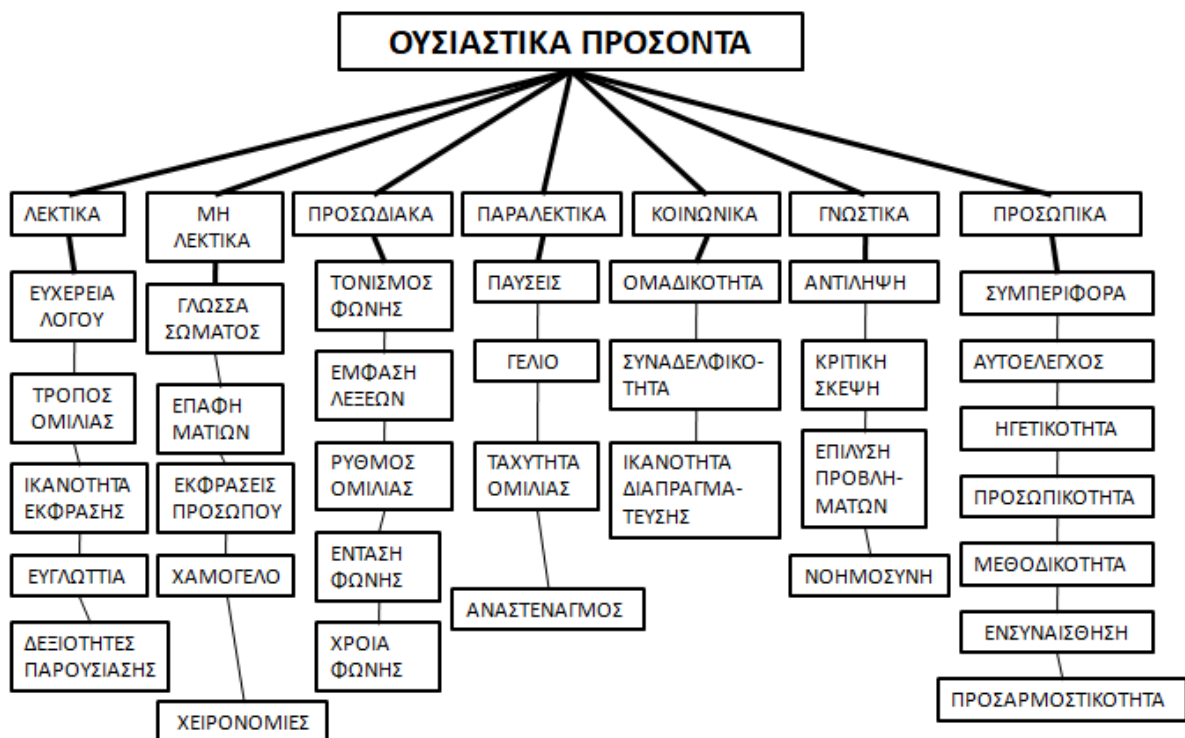
Εικόνα 3 - Αριθμός άρθρων ανά έτος δημοσίευσης

5. Ανάλυση Βιβλιογραφίας

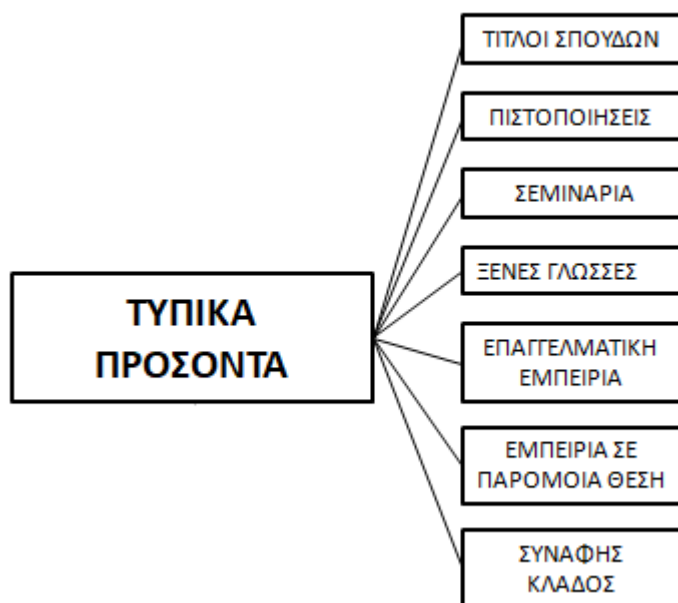
5.1 Ιεραρχική Κατηγοριοποίηση Άρθρων

Τα επιλεγμένα άρθρα κατηγοριοποιήθηκαν σύμφωνα με διάφορες διαστάσεις. Η πρώτη διάσταση αφορά το είδος των προσόντων στα οποία εστιάζουν τα άρθρα. Οι δύο βασικές κατηγορίες προσόντων είναι τα ουσιαστικά και τα τυπικά προσόντα. Τα ουσιαστικά προσόντα συνδέονται με τις επτά υπο-κατηγορίες τους, δηλαδή με τις κατηγορίες: λεκτικά, μη λεκτικά, προσωδιακά, παραλεκτικά, κοινωνικά, γνωστικά και προσωπικά. Τα ουσιαστικά προσόντα συνδέονται με τις υπο-κατηγορίες τους με έντονη γραμμή. Κάτω από αυτές τις κατηγορίες βρίσκονται τα γνωρίσματά τους. Τα τυπικά προσόντα συνδέονται με όλα τα παραδείγματά τους, πλάι από αυτά. Στην Εικόνα 4 παρατίθεται η κατηγοριοποίηση των ουσιαστικών προσόντων, ενώ στην

Εικόνα 5, η κατηγορία των ουσιαστικών προσόντων με τα παραδείγματά της. Σε όλες τις εικόνες που ακολουθούν, οι υποκατηγορίες ή παραδείγματα/γνωρίσματα κατηγοριών συνδέονται με τις πατρικές τους κατηγορίες μέσω έντονων γραμμών (πχ. ΛΕΚΤΙΚΑ με ΟΥΣΙΑΣΤΙΚΑ ΠΡΟΣΟΝΤΑ), ενώ όλοι οι υπόλοιποι αδελφικοί κόμβοι που βρίσκονται στο ίδιο επίπεδο (και έχουν την ίδια πατρική κατηγορία) συνδέονται με απλές γραμμές (π.χ. ΕΥΧΕΡΕΙΑ ΛΟΓΟΥ & ΤΡΟΠΟΣ ΟΜΙΛΙΑΣ είναι ΛΕΚΤΙΚΑ ΠΡΟΣΟΝΤΑ).

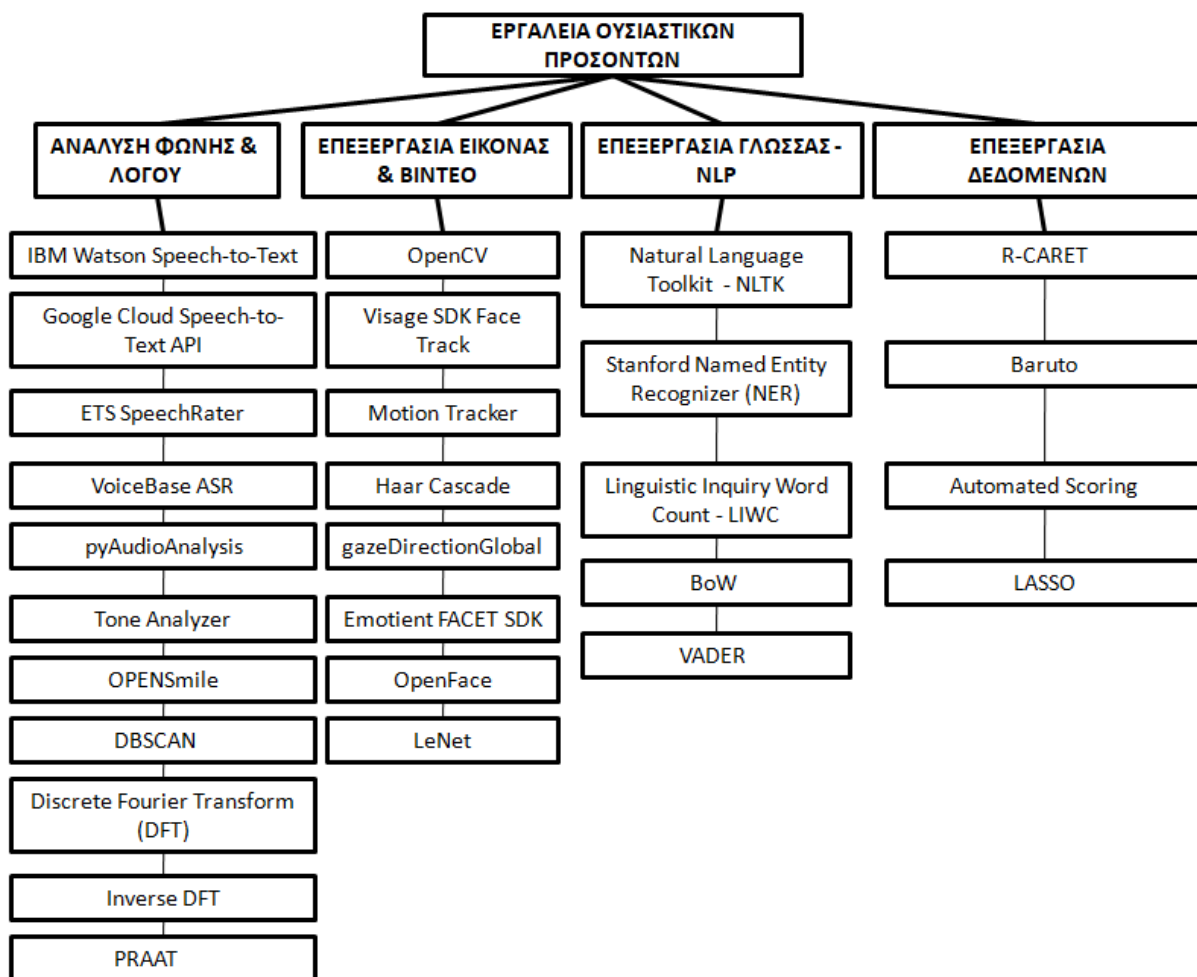


Εικόνα 4 - Οι κατηγορίες των ουσιαστικών προσόντων και οι υποκατηγορίες τους.



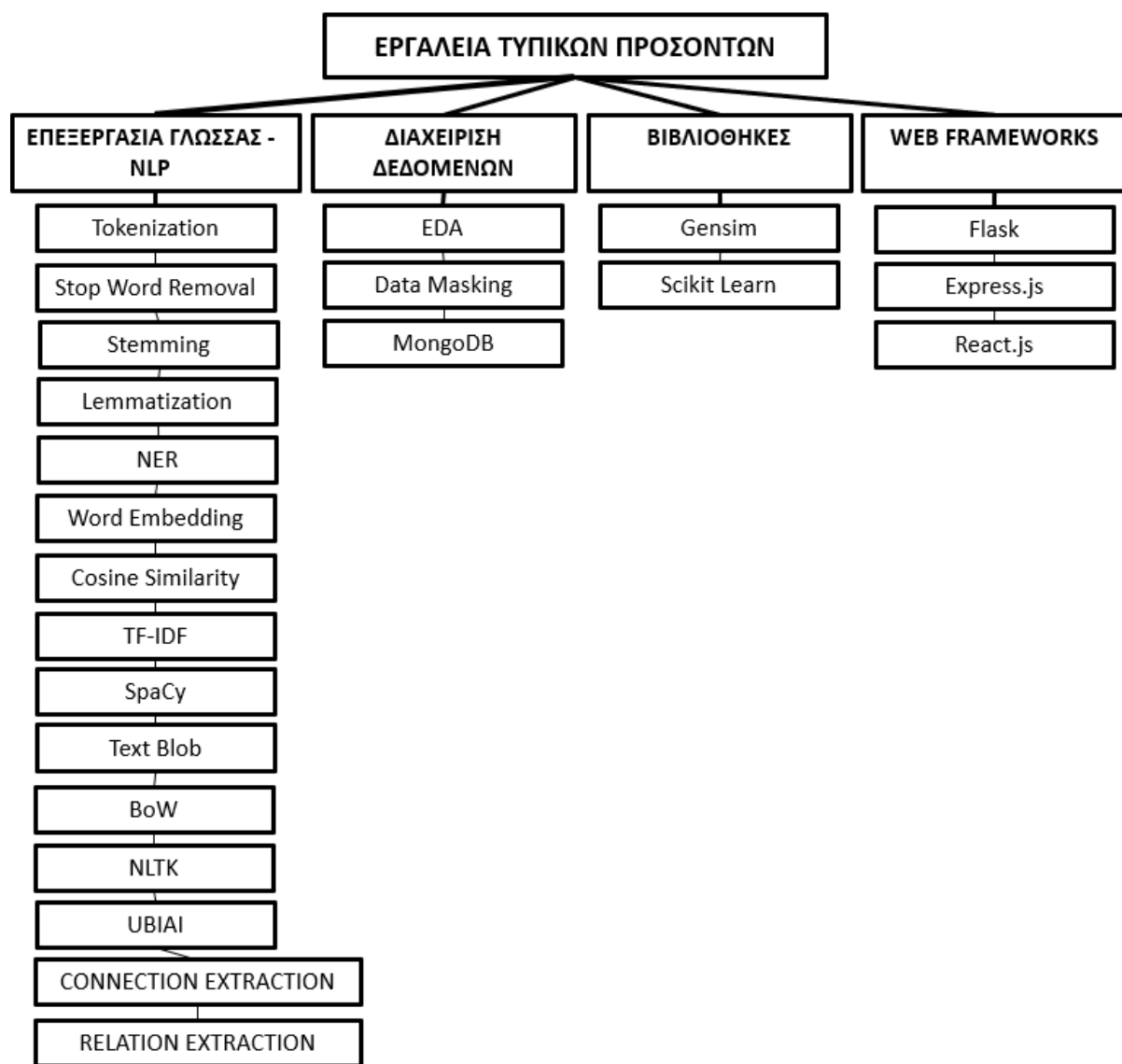
Εικόνα 5 - Οι κατηγορίες των τυπικών προσόντων.

Η δεύτερη διάσταση αφορά την επεξεργασία των δεδομένων. Στην Εικόνα 6 απεικονίζονται τα εργαλεία και οι εφαρμογές που χρησιμοποιήθηκαν για την εξαγωγή, την ανάλυση, ή την επεξεργασία των χαρακτηριστικών και των στοιχείων ουσιαστικών προσόντων για να εκπαιδευτούν στη συνέχεια οι αλγόριθμοι Μηχανικής Μάθησης. Η ανάλυση της φωνής και του λόγου, η εικόνα και το βίντεο, η ανάλυση και επεξεργασία του κειμένου και των δεδομένων, είναι οι βασικές κατηγορίες. Κάτω από αυτές τις τέσσερις κατηγορίες βρίσκονται τα εργαλεία τους.



Εικόνα 6 - Παράθεση των εργαλείων και εφαρμογών που χρησιμοποιήθηκαν στα άρθρα με τα ουσιαστικά προσόντα.

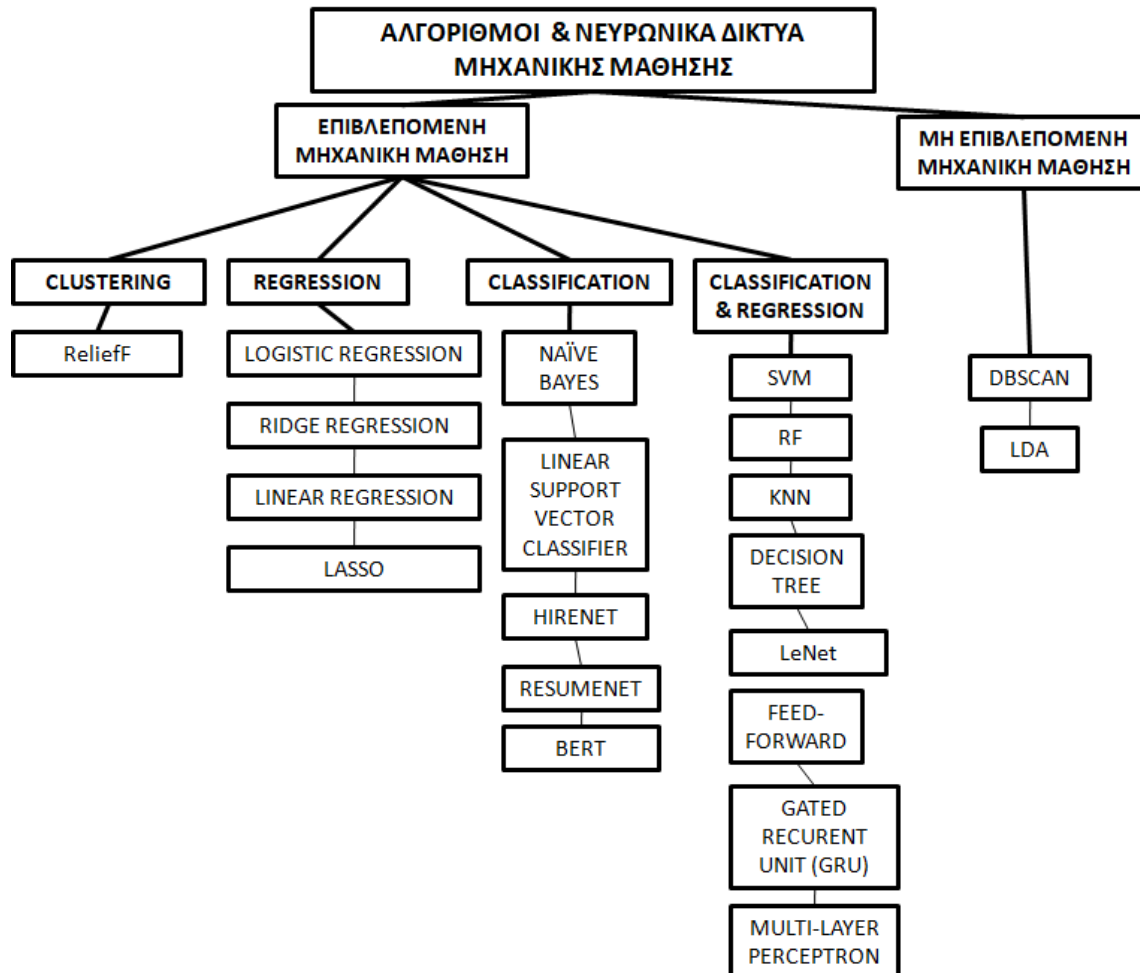
Συνεχίζοντας με τη δεύτερη διάσταση, στην Εικόνα 7 απεικονίζονται τα εργαλεία και οι εφαρμογές που χρησιμοποιήθηκαν για την εξαγωγή, την ανάλυση, ή την επεξεργασία των χαρακτηριστικών τυπικών προσόντων για να εκπαιδευτούν στη συνέχεια με βάση αυτά οι αλγόριθμοι Μηχανικής Μάθησης. Οι βασικές κατηγορίες είναι η επεξεργασία φυσικής γλώσσας (NLP), η διαχείριση των δεδομένων, οι βιβλιοθήκες και τέλος τα Web Frameworks (δηλ. τα πλαίσια ανάπτυξης εφαρμογών ιστού).



Εικόνα 7 - Παράθεση των εργαλείων και εφαρμογών που χρησιμοποιήθηκαν στα άρθρα με τα τυπικά προσόντα.

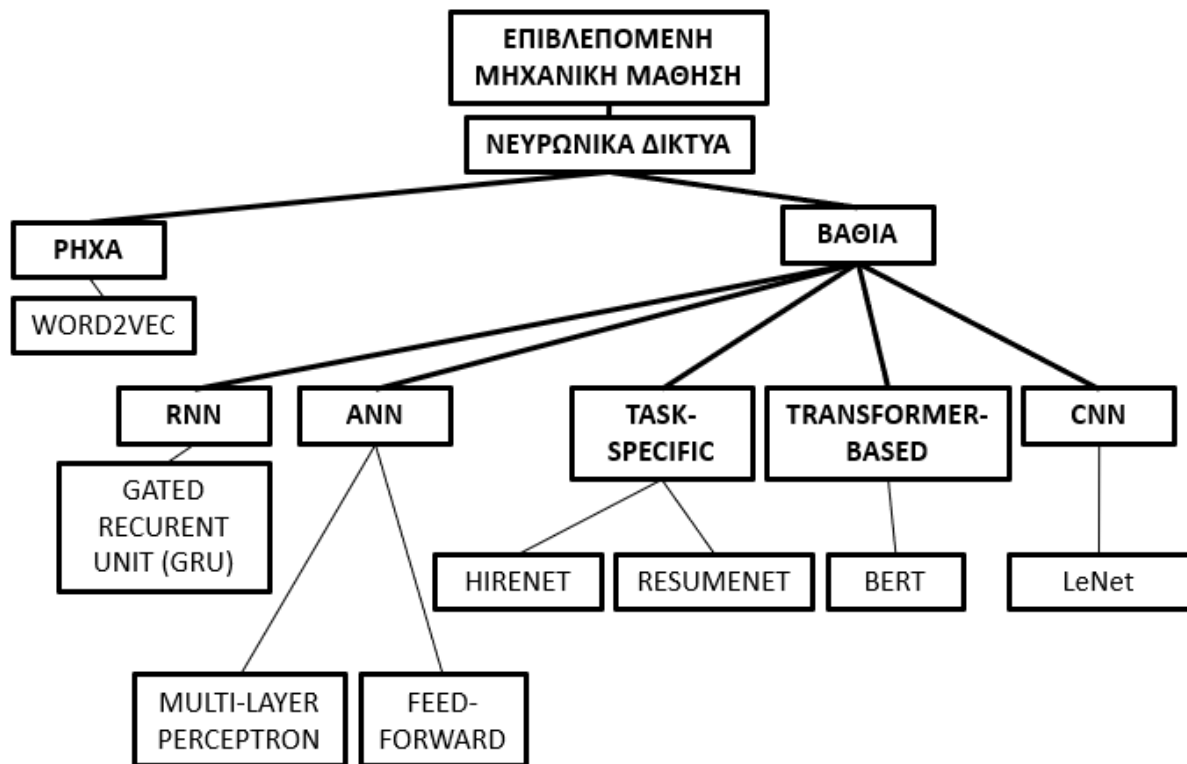
Η τρίτη διάσταση αφορά τις κατηγορίες αλγορίθμων και νευρωνικών δικτύων (Neural Networks - NNs) που χρησιμοποιήθηκαν για εκπαίδευση στα άρθρα που εξετάστηκαν. Οι αλγόριθμοι που ανήκουν στην κατηγορία της επιβλεπόμενης Μηχανικής Μάθησης χωρίζονται σε: clustering (ομαδοποίηση), regression (παλινδρόμηση), και classification (ταξινόμηση). Υπάρχουν όμως και αλγόριθμοι και NN που ανήκουν και στις δύο κατηγορίες των classification και regression. Οι αλγόριθμοι MM που χρησιμοποιούνται περισσότερο και στις δύο κατηγορίες άρθρων είναι οι: Random Forest, Support Vector Machine, ενώ ο KNN, με δεύτερο τον Naive

Bayes είναι αυτοί που χρησιμοποιούνται περισσότερο στα άρθρα των τυπικών προσόντων. Στην Εικόνα 8 απεικονίζονται οι κατηγορίες αλγορίθμων MM και NN.



Εικόνα 8 - Οι κατηγορίες των αλγορίθμων MM και NN.

Συνεχίζοντας με τη τέταρτη διάσταση, εστιάζουμε στα NN που χρησιμοποιήθηκαν στα άρθρα που μελετήθηκαν. Ειδικότερα, η Εικόνα 9 αναφέρεται στα ρηχά και βαθιά νευρωνικά δίκτυα και τις υπο-κατηγορίες αυτών. Σημειώνεται πως τα Νευρωνικά Δίκτυα ανήκουν στην κατηγορία της Επιβλεπόμενης Μηχανικής Μάθησης.



Εικόνα 9 - Τα νευρωνικά δίκτυα Μηχανικής Μάθησης που χρησιμοποιήθηκαν στις έρευνες.

Τέλος, έχουμε τη διάσταση των συνεντεύξεων, η οποία διαχωρίζεται στις υπο-διαστάσεις του τύπου της συνέντευξης και του μέσου με το οποίο έγινε η συλλογή και μετέπειτα η αξιολόγηση των δεδομένων της συνέντευξης. Οι συνεντεύξεις κατηγοριοποιήθηκαν σε πραγματικές και εικονικές, και οι υποψήφιοι αξιολογήθηκαν είτε από άτομα, είτε από σύστημα. Εντοπίστηκαν τα χαρακτηριστικά και τα στοιχεία εκείνα, ώστε να κριθεί και να βαθμολογηθεί κάθε υποψήφιος για τα ουσιαστικά του προσόντα. Η αξιολόγηση στο πλαίσιο των συνεντεύξεων αναφέρεται στον εντοπισμό των σημασιολογικών δεδομένων, σύμφωνα με τα οποία βαθμολογούν οι αξιολογητές την καταλληλότητα του υποψήφιου για την κενή θέση εργασίας. Το σύστημα επίσης εντοπίζει και συλλέγει δεδομένα από τις συνεντεύξεις για να εκπαιδευτούν στη συνέχεια οι αλγόριθμοι MM με αυτά. Οι συνεντεύξεις λειτουργούν ως πηγή δεδομένων για τα ουσιαστικά προσόντα και τα χαρακτηριστικά των υποψηφίων. Τα συλλεγόμενα δεδομένα, είτε από τους αρμόδιους υπαλλήλους είτε από το σύστημα, χρησιμοποιούνται για την απόφαση πρόσληψης, ή μη, του

υποψήφιου. Στην Εικόνα 10 απεικονίζεται η κατηγοριοποίηση με βάση αυτή τη διάσταση.



Εικόνα 10 - Οι κατηγορίες των συνεντεύξεων και από ποιους αξιολογούνται οι υποψήφιοι.

5.2 Ανάλυση Άρθρων

Ακολουθεί η ανάλυση των εξεταζόμενων άρθρων. Αρχικά θα παρουσιαστούν τα άρθρα των ουσιαστικών προσόντων, δηλαδή αυτών που ασχολούνται με την εξαγωγή και ανάλυση των χαρακτηριστικών προσωπικότητας. Στη συνέχεια θα παρουσιαστούν τα άρθρα που βασίζονται στα τυπικά προσόντα, όπως το πτυχίο, η προϋπηρεσία, ο κλάδος, οι ξένες γλώσσες κ.ο.κ. Κατά κανόνα, στα αναλυόμενα άρθρα ακολουθείται η ίδια διαδικασία. Σε πρώτη φάση συγκεντρώνονται οι υποβαλλόμενες αιτήσεις, ενώ οι συνεντεύξεις μπορούν να πραγματοποιούνται ασύγχρονα. Η εξαγωγή χαρακτηριστικών και γνωρισμάτων των υποψηφίων γίνεται με εφαρμογές εξαγωγής στοιχείων. Στη συνέχεια οι αλγόριθμοι Μηχανικής Μάθησης εκπαιδεύονται για την καλύτερη δυνατή πρόβλεψη επιλογής υπαλλήλου.

5.2.1 Ανάλυση Άρθρων Εστιαζόμενων στα Ουσιαστικά Προσόντα

5.2.1.1 Leveraging Multimodal Behavioral Analytics for Automated Job Interview (Αξιοποίηση Πολυτροπικής Ανάλυσης Συμπεριφοράς για Αυτοματοποιημένη Συνέντευξη Εργασίας)

Στο άρθρο των Agrawal και συν. (2020) αναλύεται ποικιλοτρόπως η συμπεριφορά των υποψηφίων, με αυτόματη αξιολόγηση απόδοσης και ανατροφοδότηση των συνεντεύξεών τους. Το δείγμα των συνεντεύξεων, όπου όλες είναι εικονικές, ανακτήθηκε από τη βάση δεδομένων του MIT. Οι συμμετέχοντες στο διαμορφωμένο δείγμα είναι 69 και οι συνεντεύξεις 138, οι οποίες αξιολογήθηκαν από τους Amazon Mechanical Turk Workers. Οι Amazon Mechanical Turk Workers³² είναι εργαζόμενοι στην πλατφόρμα Amazon Mechanical Turk, η οποία είναι μία online υπηρεσία. Αυτοί οι εργαζόμενοι προσφέρουν υπηρεσίες (όπως αξιολόγησης) για ένα μικρό αντίτιμο. Είναι ανώνυμοι, ανεξάρτητοι και κατοικούν οπουδήποτε στον κόσμο. Οι αξιολογητές επικεντρώθηκαν στις εξής παραμέτρους: την επαφή με τα μάτια, τον ρυθμό ομιλίας, την αφοσίωση, την ηρεμία, τις παύσεις, την ψυχραιμία, την επικέντρωση, την αυθεντικότητα και την επιδεξιότητα. Ακολούθως οι Amazon Mechanical Turk Workers εξήγαγαν συμπεράσματα σχετικά με την καταλληλότητα για πρόσληψη, σύμφωνα με τα ακόλουθα κριτήρια: όσον αφορά τον ήχο, από την ισχύ, την ένταση, τη διάρκεια, το ύψος της φωνής και την ενέργεια, όσον αφορά δε τα βίντεο, τα χαρακτηριστικά του προσώπου και τον θετικό/αρνητικό αντίκτυπο που έχει σε αυτούς, όπως το χαμόγελο. Τέλος, αξιολόγησαν τα λεκτικά ερεθίσματα, όπως τον ρυθμό ομιλίας, την επάρκεια, την ευχέρεια, το σύνολο των μοναδικών λέξεων, ουσιαστικών, ρημάτων και επιθέτων που χρησιμοποιήθηκαν, όπως επίσης και το συναίσθημα.

Στη συνέχεια, χρησιμοποιήθηκαν οι εξής μέθοδοι επεξεργασίας δεδομένων και εξαγωγής χαρακτηριστικών, για να εξαχθούν τα επιθυμητά στοιχεία, ώστε να γίνει η ταξινόμηση και η πρόβλεψη από τους αλγόριθμους Μηχανικής Μάθησης. Το Discrete Fourier Transform (DFT) (Delgutte & Greenberg, 1999) υπολογίζει την συχνότητα και το μέγεθος των κυρίαρχων χαρακτηριστικών, το Inverse DFT (Delgutte & Greenberg, 1999) υπολογίζει τα χρονικά δεδομένα από τις συχνότητες, ενώ το PRAAT³³ αναλύει και επεξεργάζεται την φωνή, όπως το τονικό ύψος (pitch) και την ένταση. Το pyAudioAnalysis (Giannakopoulos, 2015), αφού αναλύσει και επεξεργαστεί τον ήχο,

³² [Amazon Mechanical Turk](#)

³³ [Praat - Download](#)

εξάγει ακουστικά χαρακτηριστικά, όπως η ένταση και η μελωδία. Το OpenCV³⁴ εξάγει χαρακτηριστικά του προσώπου μετά την επεξεργασία μιας εικόνας, ενώ το LeNet³⁵ είναι ένα νευρωνικό δίκτυο που αναγνωρίζει χειρόγραφα ψηφία και εικόνες, στην περίπτωση της τρέχουσας αναλυόμενης εργασίας ανιχνεύει το χαμόγελο. Το Google Cloud Speech-to-Text API³⁶, μετατρέπει την ομιλία σε κείμενο, ενώ το Natural Language Toolkit (NLTK)³⁷ χωρίζει το κείμενο σε μικρότερες μονάδες, αφού επεξεργαστεί και αναλύσει την γλώσσα. Το τελευταίο, δύναται να αναγνωρίσει ονόματα, να αναλύσει το συντακτικό και το συναίσθημα. Το Tone Analyzer αναλύει τα συναισθήματα, εντοπίζοντάς τα σε ένα κείμενο, όπως τη λύπη, το θυμό, τη χαρά και τον ενθουσιασμό. Τέλος, το Stanford Named Entity Recognizer (NER) αναγνωρίζει και κατηγοριοποιεί τις οντότητες των κειμένων, όπως τα άτομα, τις τοποθεσίες, τα προϊόντα και τους οργανισμούς. Στη συγκεκριμένη περίπτωση αναγνωρίζει κυρίως τα ουσιαστικά, τα επίθετα, τα ρήματα κλπ.

Έπειτα, εκπαιδεύτηκαν οι εξής αλγόριθμοι ταξινόμησης μηχανικής μάθησης για να προβλέψουν την βαθμολογία των αξιολογητών: Random Forest (RF), Support Vector Machine (SVM), Multi-Layer Perceptron (MLP) και Least Absolute Shrinkage and Selection Operator (LASSO). Οι αλγόριθμοι δοκίμασαν όλους τους συνδυασμούς των εξαγόμενων χαρακτηριστικών του ήχου, των βίντεο και των λεκτικών ερεθισμάτων. Η βαθμολογία είναι σε κλίμακα 1-7, βασισμένη στις προαναφερθείσες παραμέτρους. Ο RF είχε την μεγαλύτερη πιστότητα (accuracy) χρησιμοποιώντας ταυτόχρονα όλες τις παραμέτρους (ήχο, λεκτικά στοιχεία, βίντεο), που έφτανε το 96,43%. Άλλες μετρικές δεν εξετάστηκαν.

Θετικό σε αυτή τη μελέτη είναι ότι προσφέρει ανατροφοδότηση προς τους συνεντευξιαζόμενους για τα ισχυρά και αδύναμα χαρακτηριστικά τους (ενδείξεις συμπεριφοράς, γλώσσα του σώματος), ώστε να είναι εφικτή η βελτίωσή τους. Είναι κάτι που στο μέλλον θα πραγματοποιείται σε πραγματικό χρόνο. Αν και ενσωματώνονται αρκετά χαρακτηριστικά, θα μπορούσε να συμπεριληφθεί η κίνηση των χεριών και η στάση του σώματος, κάτι που είναι σημαντικό στη μη λεκτική

³⁴ [OpenCV Face Recognition powered by Seventh Sense - OpenCV](#)

³⁵ [LeNet: The Deep Learning Model That Revolutionized Image Recognition | by Neha Purohit | AI monks.io | Medium](#)

³⁶ [Speech-to-Text documentation | Cloud Speech-to-Text Documentation | Google Cloud](#)

³⁷ Το Natural Language Toolkit (NLTK) είναι μία βιβλιοθήκη Python και χρησιμοποιείται για την επεξεργασία φυσικής γλώσσας (NLP) με σκοπό την ανάλυση και την κατανόηση ενός κειμένου. Κάποιες χρήσεις του είναι: tokenization, stemming, lemmatization, Named Entity Recognition (NER), ανάλυση συναισθημάτων. ([NLTK](#))

επικοινωνία. Επιπλέον πλεονέκτημα της έρευνας είναι η χρήση πολλών εργαλείων για την εξαγωγή μεγάλου αριθμού ουσιαστικών χαρακτηριστικών από διαφορετικά είδη δεδομένων. Αυτό συμβάλλει στην ακρίβεια και την αξιοπιστία των μοντέλων της Μηχανικής Μάθησης και αποδεικνύεται στο υψηλό ποσοστό απόδοσης του Random Forest. Ωστόσο, ένα αρνητικό της έρευνας είναι ότι οι συνεντεύξεις ήταν εικονικές, καθώς οι υποψήφιοι ενδέχεται να μην έχουν το ίδιο επίπεδο άγχους όπως στις πραγματικές και αυτό επηρεάζει τα αποτελέσματα. Επίσης, ο αριθμός των συνεντευζομένων ήταν μικρός.

5.2.1.2 Automatic Prediction of Fluency in Interface-Based Interviews (Αυτόματη Πρόβλεψη της Ευχέρειας Λόγου σε Συνεντεύξεις που Βασίζονται σε Διεπαφή)

Στο άρθρο των Rasipuram και συν. (2016) προτείνεται ένα σύστημα που προβλέπει την ευχέρεια λόγου (fluency) μέσα από εικονικές βιντεοσκοπημένες συνεντεύξεις. Οι συνεντεύξεις ήταν 106, αποτελούμενες από 58 άντρες και 48 γυναίκες, όλοι φοιτητές του πανεπιστημίου International Institute of Information Technology, στο Bangalore της Ινδίας. Οι αξιολογητές ήταν τρεις, οι οποίοι αξιολόγησαν την ευχέρεια λόγου με κλίμακα από το 1 έως το 5: {1 (ακατάλληλος), 2 (αρκετός), 3 (επαρκής), 4 (καλός) και 5 (εξαιρετικός)}. Οι δύο από τους τρεις αξιολογητές είχαν εμπειρία στην αξιολόγηση συνεντεύξεων συμπεριφοράς για περισσότερα από οκτώ χρόνια.

Οι συνεντευζοόμενοι έπαιρναν ανατροφοδότηση από το σύστημα διεπαφής κατά τη διάρκεια της συνέντευξης. Η ευχέρεια λόγου προσδιορίζεται ως η συνολική επάρκεια στην ομιλία, η ικανότητα να μιλάμε με λίγες παύσεις, με συνεκτικότητα, αιτιολογημένα, πυκνά, κατάλληλα, δημιουργικά και ευφάνταστα. Υπάρχουν λεκτικά, και μη, στοιχεία που σχετίζονται με την ευχέρεια. Αυτά είναι πρώτον, τα προσωδιακά χαρακτηριστικά (prosodic) που αποτελούνται από τον τονισμό, το ρυθμό, την ένταση και την παύση κατά τη διάρκεια της ομιλίας και δεύτερον, τα γλωσσικά χαρακτηριστικά (linguistic) στα οποία εμπεριέχονται οι ήχοι, η προφορά, η δομή των προτάσεων, η γραμματική, το λεξιλόγιο και τα σημασιολογικά. Τρίτον, τα λεξιλογικά χαρακτηριστικά (lexical) αναφέρονται στο μήκος και τη συχνότητα κάθε λέξης. Αυτά λοιπόν τα χαρακτηριστικά εξήχθησαν και αξιολογήθηκαν από αλγόριθμους Μηχανικής Μάθησης.

Τα λεξιλογικά χαρακτηριστικά εξήχθησαν από το VoiceBase ASR (Automatic Speech Recognition³⁸), το οποίο αναγνωρίζει αυτόματα και μετατρέπει τα ηχητικά δεδομένα σε κείμενο. Το OPENSmile³⁹ είναι εφαρμογή που αναλύει τα φωνητικά χαρακτηριστικά και ανιχνεύει τα συναισθήματα, οπότε χρησιμοποιείται για την εξαγωγή των προσωδιακών χαρακτηριστικών, όπως και το PRAAT. Αρχικά, και πριν από την χρήση των άλλων αλγορίθμων MM, χρησιμοποιήθηκε ο αλγόριθμος reliefF που διευκολύνει την κατανόηση και αξιολογεί την βαρύτητα των εξαγόμενων χαρακτηριστικών, ώστε οι επόμενοι αλγόριθμοι MM να εκπαιδεύονται με τα πιο σημαντικά χαρακτηριστικά. Οι αλγόριθμοι MM που χρησιμοποιήθηκαν είναι οι: Linear Regression, Logistic Regression, και Support Vector Machine (SVM) με RBF Kernel (πυρήνας για την ταξινόμηση διανυσματικών μηχανών). Το ποσοστό της ακρίβειας (mean average precision) για τα προσωδιακά χαρακτηριστικά έφτασε το 83% και για τα γλωσσικά το 78%.

Αυτό που δεν ανταποκρίθηκε στις προσδοκίες των ερευνητών του άρθρου ήταν ότι το ASR δεν αποκωδικοποιούσε πάντα σωστά τις λέξεις. Επίσης, επειδή ο αριθμός των συνεντεύξεων και των υποψηφίων ήταν μικρός καθώς και το ότι οι συνεντεύξεις ήταν εικονικές, δεν μπορούν να εξασφαλιστούν εμπειριστατωμένα αποτελέσματα. Ακόμη, αν και οι δύο αξιολογητές είχαν πολυετή εμπειρία, δεν αναφέρεται η εμπειρία του τρίτου αξιολογητή.

Από την άλλη μεριά, μεγάλο πλεονέκτημα είναι ότι η διαδικασία της κατηγοριοποίησης μπορεί να πραγματοποιηθεί οπουδήποτε και οποτεδήποτε χωρίς την ανάγκη έμψυχου υλικού. Θετικό επίσης είναι ότι χρησιμοποιήθηκε ποικιλία εργαλείων για την εξαγωγή των χαρακτηριστικών και ο reliefF για την ταξινόμηση και επιλογή των χαρακτηριστικών με βάση τη βαρύτητά τους. Τέλος, η χρήση πολυάριθμων λεκτικών και μη χαρακτηριστικών για την αξιολόγηση της ευχέρειας λόγου, ήταν παραπάνω από ικανοποιητική. Στο άρθρο όμως δεν αναφέρεται εάν οι εικονικές συνεντεύξεις πραγματοποιούνται για την πλήρωση μίας ή περισσότερων θέσεων (Rasipuram et al., 2016).

³⁸ [How VoiceBase Automatic Speech Recognition \(ASR\) Works - VoiceBase](#)
³⁹ [openSMILE: Audio Fingerprinting and Feature Extraction](#)

5.2.1.3 Automatic Scoring of Monologue Video Interviews Using Multimodal Cues (Αυτόματη Βαθμολόγηση Βίντεο από Μονολόγες Συνεντεύξεις με Χρήση Πολλαπλών Χαρακτηριστικών)

Σκοπός των Chen και συν. (2016) είναι η δημιουργία ενός συστήματος αυτόματης αξιολόγησης για βίντεο-συνεντεύξεις με μονόλογο. Με την εξέλιξη της τεχνολογίας, πλέον είναι εφικτή η ποσοτικοποίηση μιας συνέντευξης. Ο ερευνητικός τομέας της Επεξεργασίας του Κοινωνικού Σήματος, Social Signal Processing (SSP) (Vinciarelli, 2017), επιτρέπει στους υπολογιστές να ανιχνεύουν, να αναγνωρίζουν, να αναλύουν, να ερμηνεύουν και να κατανοούν τα ανθρώπινα κοινωνικά σήματα, όπως οι εκφράσεις του προσώπου, ο τόνος της φωνής, οι χειρονομίες και η γλώσσα του σώματος, μέσω βίντεο. Η Επεξεργασία του Κοινωνικού Σήματος γενικότερα, αναλύει την ανθρώπινη επικοινωνία παρατηρώντας τη μη λεκτική συμπεριφορά.

Οι συμμετέχοντες στις συνεντεύξεις ήταν 36. Η πλειοψηφία συμμετείχε δωρεάν αφού προήλθαν από την Υπηρεσία Εκπαιδευτικών Δοκιμών των Η.Π.Α. (Educational Testing Service - ETS⁴⁰), έναν ιδιωτικό μη κερδοσκοπικό εκπαιδευτικό οργανισμό διαγωνισμών και αξιολόγησης του κόσμου, στον οποίο εργάζονται οι συγγραφείς. Οι υπόλοιποι πληρώθηκαν 60 USD καθώς δεν ανήκαν στον οργανισμό. Οι εικονικές συνεντεύξεις ήταν 419 στο σύνολο και οι αξιολογητές 9. Οι ερωτήσεις ήταν 12 και προκαθορισμένες. Οι αξιολογητές βαθμολόγησαν την επίδοση, την προσωπικότητα και συνολικά τις κοινωνικές δεξιότητες κάθε συμμετέχοντα σύμφωνα με το βίντεο-μονολόγό του. Εστίασαν στα “Πέντε Μεγάλα χαρακτηριστικά προσωπικότητας” (Big Five personality traits⁴¹) και τα μετέφεραν στην κλίμακα 1-7 του Likert⁴². Τα πέντε μεγάλα χαρακτηριστικά είναι η εξωστρέφεια (extraversion), η συγκαταβατικότητα (agreeableness), η διαφάνεια (openness), η ευσυνειδησία (conscientiousness) και ο νευρωτισμός/συναισθηματική σταθερότητα (neurotism/emotional stability). Με μια κάμερα Kinect κατέγραφαν τις κινήσεις του σώματος και με μια web κάμερα βιντεοσκοπούσαν τις εικονικές συνεντεύξεις.

Χρησιμοποιήθηκαν για την εξαγωγή των χαρακτηριστικών: το Linguistic Inquiry Word Count (LIWC⁴³) που εξάγει τα λεξιλογικά χαρακτηριστικά (λέξεις που περιγράφουν

⁴⁰ [ETS](#) - Educational Testing Service

⁴¹ [Big 5 Personality Traits: The 5-Factor Model of Personality](#) (Lim, 2023).

⁴² [Likert scale | Social Science Surveys & Applications | Britannica](#) Britannica.

⁴³ [Welcome to LIWC-22](#)

συναισθήματα), το Automated Scoring (AS), μία τεχνολογία που μετρά την προσωδία, την ακρίβεια και τη χρήση της γλώσσας, το ETS SpeechRater⁴⁴, ένα σύστημα που παράγει χαρακτηριστικά δεξιοτήτων ομιλίας, το Visage SDK Face Track που μετρά την στάση του κεφαλιού, το gazeDirectionGlobal⁴⁵, το οποίο είναι ένα εργαλείο για την εξαγωγή της κατεύθυνσης του βλέμματος και το Emotient FACET SDK⁴⁶, το οποίο είναι επίσης εργαλείο εξαγωγής για τα αισθήματα και τις εκφράσεις του προσώπου, όπως το χαμόγελο. Οι αλγόριθμοι που χρησιμοποιήθηκαν είναι οι: SVM, LASSO και Ridge Regression. Για να μετρηθούν τα ποσοστά απόδοσης, χρησιμοποιήθηκε το R-CARET⁴⁷ (Classification And Regression Training), ένα πακέτο παλινδρομικής προσέγγισης στην γλώσσα προγραμματισμού R. Χρησιμοποιήθηκαν τα δεδομένα των 35 συνεντευξιαζόμενων, έναντι ενός, για εκπαίδευση και αυτό επαναλήφθηκε 36 φορές και για διαφορετική κατηγορία χαρακτηριστικών (οπτικά, λεκτικά και λεξιλογικά, συνολικά) και συνδυασμούς τους κάθε φορά. Η καλύτερη απόδοση πρόβλεψης προέκυψε όταν οι αλγόριθμοι δοκίμασαν συνολικά τα εξαγόμενα χαρακτηριστικά, δηλαδή τις λεξιλογικές, λεκτικές και οπτικές πτυχές ταυτόχρονα. Τα αποτελέσματα για κάθε αλγόριθμο είναι: SVM 44.6%, LASSO 45,2% και RR 45,8%. Τα ποσοστά αυτά δείχνουν το κατά πόσο οι αλγόριθμοι ταυτίστηκαν με τα αποτελέσματα των αξιολογητών.

Οι τιμές αυτές είναι απογοητευτικές καθώς η απόδοση των αλγορίθμων είναι χαμηλή. Το άρθρο αναφέρει ότι χάθηκαν δεδομένα συνεντεύξεων λόγω προβλημάτων στο υλικό και το λογισμικό, όμως δεν αναφέρει τον ακριβή αριθμό. Οι ερευνητές δεν εξηγούν το λόγο της απώλειας δεδομένων, θα μπορούσαν όμως να προετοιμαστούν καλύτερα, έχοντας προβλέψει εφεδρικά μηχανήματα και λογισμικό. Σίγουρα αυτή η απώλεια επηρεάζει την ποιότητα της έρευνας. Αν και οι εναπομείνανες συνεντεύξεις ήταν 419, μπορεί να θεωρηθεί ικανός αριθμός για συλλογή δεδομένων. Από την άλλη, αυτές οι 419 συνεντεύξεις δεν προέρχονται από 419 διαφορετικά άτομα, αλλά μόνο από 36. Θετική ήταν η χρήση πολυτροπικών δεδομένων (λεξιλογικών, λεκτικών, οπτικών) για την αξιολόγηση κοινωνικών δεξιοτήτων αν και η ανάγκη βελτίωσης στον τρόπο εξαγωγής και συνδυασμού τους είναι φανερή (Chen et al, 2016).

⁴⁴ [SpeechRater Service: Advanced Spoken-response Scoring Application](#) (ETS, 2025).

⁴⁵ [VisageTracker-Configuration-Manual.pdf](#) - Visage Technologies

⁴⁶ [Emotient and iMotions partner for integrated facial expression recognition, bio sensor and eye-tracking solution | Biometric Update](#) (Vrankulj, 2013).

⁴⁷ [Caret package in R](#) (Hoeting, 2020).

5.2.1.4 Automated Video Interview Judgment on a Large-Sized Corpus Collected Online (Αυτοματοποιημένη Αξιολόγηση Συνέντευξης Βίντεο σε Μεγάλου Μεγέθους Σύνολο Δεδομένων Συλλεγόμενων από το Διαδίκτυο)

Στην έρευνα των Chen και συν. (2017) συνελέγησαν 1891 μονόλογοι από εικονικές βίντεο-συνεντεύξεις σε πραγματικό χρόνο από 260 άτομα. Η έρευνα εστιάζει στην πρόβλεψη για σύσταση πρόσληψης βασιζόμενη στα χαρακτηριστικά προσωπικότητας. Σημειώνεται πως εφόσον είναι εικονικές οι συνεντεύξεις, δεν αφορούν συγκεκριμένη θέση εργασίας. Τα 260 συμμετέχοντα άτομα ήταν από το Amazon Mechanical Turk Workers, εγκατεστημένα στις Η.Π.Α. και πληρώθηκαν 10 δολάρια για τη συμμετοχή τους. Οι αξιολογητές ήταν πέντε, έμπειροι στη βαθμολόγηση συγγραφής δοκιμίων και την αξιολόγηση απόδοσης ατόμων σε βίντεο καθώς εργάζονταν στην Υπηρεσία Εκπαιδευτικών Δοκιμών των Η.Π.Α. (Educational Testing Service - ETS). Οι ερωτήσεις ήταν οκτώ, προκαθορισμένες και αποσκοπούσαν στο να εντοπίσουν τέσσερις τύπους κοινωνικών δεξιοτήτων, φιλτράροντας τα χαρακτηριστικά προσωπικότητας των συνεντευξιαζόμενων. Οι τέσσερις τύποι κοινωνικών δεξιοτήτων είναι: οι δεξιότητες επικοινωνίας στο εργασιακό περιβάλλον, οι διαπροσωπικές δεξιότητες, η ηγετικότητα και η πειθώ/διαπραγμάτευση. Το αποτέλεσμα ήταν να επιλεχθούν οι κατάλληλοι υπάλληλοι. Αυτό πραγματοποιήθηκε με την αυτόματη ανάλυση του περιεχομένου ομιλίας, τις εκφράσεις του προσώπου και τα προσωδιακά χαρακτηριστικά. Το RecordRTC⁴⁸, ένα πρόγραμμα επεξεργασίας κειμένου που επιτρέπει στον χρήστη να προσθέσει σχόλια ήχου και βίντεο, χρησιμοποιήθηκε για την εγγραφή των βίντεο. Όλα βασίστηκαν στην τεχνολογία του Social Signal Processing, όπως και στην προηγούμενη έρευνα των συγγραφέων αλλά με χρήση μεγαλύτερου δείγματος και με διαφοροποιήσεις στα εργαλεία και στους αλγόριθμους Μηχανικής Μάθησης που χρησιμοποίησαν.

Η κλίμακα που χρησιμοποιήθηκε είναι η Likert, με κλίμακα από 1 έως 7, για να αξιολογηθεί ο βαθμός στον οποίο εκδηλώνονται τα “Πέντε Μεγάλα χαρακτηριστικά προσωπικότητας”. Οι δεξιότητες επικοινωνίας με το εργασιακό περιβάλλον συνδέονται με την εξωστρέφεια, την διαφάνεια και τον νευρωτισμό/συναισθηματική

⁴⁸ [RecordRTC: Home](#)

σταθερότητα. Οι διαπροσωπικές δεξιότητες συνδέονται με την συγκαταβατικότητα, την εξωστρέφεια και τον νευρωτισμό/συναισθηματική σταθερότητα. Η ηγετικότητα συνδέεται με την εξωστρέφεια, ευσυνειδησία, τον νευρωτισμό/συναισθηματική σταθερότητα και την διαφάνεια. Η πειθώ/διαπραγμάτευση συνδέεται με την εξωστρέφεια, την συγκαταβατικότητα, τη διαφάνεια και τον νευρωτισμό/συναισθηματική σταθερότητα.

Οι εφαρμογές εξαγωγής που χρησιμοποιήθηκαν ήταν: η VoiceBase ASR, η pyAudioAnalysis, η OpenFace (Baltrusaitis et al, 2016), η οποία αναλύει τις μικροεκφράσεις του προσώπου, το βλέμμα, τη στάση του κεφαλιού και τα συναισθήματα καθώς και η BoW (Bag of Words⁴⁹), η οποία εξάγει τα χαρακτηριστικά των κειμένων. Οι αλγόριθμοι MM που χρησιμοποιήθηκαν είναι ο Random Forest και ο SVM. Τα δεδομένα εκπαίδευσης (train set) ήταν 1591 και το δείγμα δοκιμής (test set) 372. Η πρόβλεψη και των πέντε χαρακτηριστικών της προσωπικότητας με τις μετρικές F1 measure, Precision και Recall κυμαίνονταν από 78% έως 86% για τα χαρακτηριστικά που εξήχθησαν από κείμενο και από 54% έως 56% για τα χαρακτηριστικά που εξήχθησαν από βίντεο. Το γενικό ποσοστό πρόβλεψης για την σύσταση πρόσληψης κυμαίνεται από 63% έως 67%.

Θετικό σε αυτή την έρευνα είναι πως το δείγμα των συνεντεύξεων είναι μεγάλο και υπάρχει ποικιλομορφία, αφού αποτελείται από άτομα από διαφορετικές περιοχές των Η.Π.Α., φύλο, εθνικότητες και ηλικίες. Έτσι, παρέχει πλούσιο υλικό για ανάλυση καθώς και ενισχύεται η αξιοπιστία των αποτελεσμάτων. Η αξιολόγηση διαφορετικών παραμέτρων (ομιλία, εκφράσεις προσώπου, προσωδιακά χαρακτηριστικά) επιτρέπει την πολυδιάστατη αξιολόγηση των συνεντευξιαζόμενων.

Δυστυχώς, οι εικονικές συνεντεύξεις σε καμία περίπτωση δεν αποτυπώνουν τις πραγματικές συνθήκες μιας πραγματικής συνέντευξης. Αρνητικό θα μπορούσε επίσης να είναι ότι πρόβλεψη δεν συνδέεται με συγκεκριμένη θέση εργασίας, αλλά με ακαθόριστη. Όμως, δεν απαιτούν όλες οι θέσεις εργασίας τα ίδια προσόντα. Η απόκλιση των ποσοστών προβλέψεων από τα δεδομένα των βίντεο και των κειμένων, υποδηλώνει ότι υπάρχουν πολλά περιθώρια βελτίωσης στην πρόβλεψη από τα δεδομένα των βίντεο, όπως και στο ποσοστό πρόβλεψης για σύσταση πρόσληψης. Τα χαμηλά ποσοστά πρόβλεψης από βίντεο θα μπορούσαν να

⁴⁹ [What is bag of words? | IBM](#) (Murel & Kavlakoglu, 2024).

οφείλονται στην χαμηλή ανάλυση/ποιότητα του βίντεο (κάποιοι συνεντευξιαζόμενοι χρησιμοποίησαν το κινητό τους τηλέφωνο). Η βελτίωση της ανάλυσης, της ποιότητας, του φωτισμού της εικόνας θα επέφερε καλύτερα αποτελέσματα καθώς θα βοηθούσε στην ακριβέστερη εξαγωγή χαρακτηριστικών.

5.2.1.5 Computational Inference of Candidates' Oratory Performance in Employment Interviews Based on Candidates' Vocal Analysis (Υπολογιστικό Συμπέρασμα της Ρητορικής Απόδοσης των Υποψηφίων σε Συνεντεύξεις Εργασίας με Βάση τη Φωνητική Ανάλυσή τους)

Ο στόχος των Sawla και συν. (2018) στην έρευνά τους είναι η δημιουργία ενός συστήματος για την αυτόματη πρόβλεψη της προφορικής επίδοσης των υποψηφίων, βασισμένη στη φωνητική τους ανάλυση. Το σύνολο δεδομένων που χρησιμοποίησαν αντλήθηκε από το MIT. Τα άτομα που έλαβαν μέρος ήταν 69, φοιτητές του πανεπιστημίου, και οι εικονικές συνεντεύξεις ήταν 138. Η αξιολόγηση πραγματοποιήθηκε σε επτά χαρακτηριστικά της προσωπικότητάς τους και στις κοινωνικές τους δεξιότητες. Τα επτά χαρακτηριστικά ήταν τα εξής: νοητική ικανότητα (mental capability), γνώση και δεξιότητες (knowledge and skills), βασικές τάσεις προσωπικότητας (basic personality tendencies), εφαρμοσμένες κοινωνικές δεξιότητες (applied social skills), ενδιαφέροντα και προτιμήσεις (interests and preferences), οργανωτική προσαρμογή (organizational fit) και φυσικές ιδιότητες (physical attributes). Για το λόγο αυτό, το δείγμα που χρησιμοποιήθηκε είναι ακουστικό. Κατέληξαν ότι τα τέσσερα προσωδιακά στοιχεία που πιστεύεται ότι βοηθούν το σύστημα για καλύτερη πρόβλεψη είναι η ενεργητικότητα (energetic), η σαφήνεια (clarity), η ευχέρεια του λόγου (fluency) και η αυτοπεποίθηση (confidence). Οι συγγραφείς βαθμολόγησαν κάθε προσωδιακό στοιχείο ξεχωριστά από το 1 έως το 10 για το πόσο αδύναμο (1) ή κυρίαρχο (10) ήταν σε κάθε άτομο.

Το PRAAT και το Baruto package στην R είναι οι δύο εφαρμογές εξαγωγής των τεσσάρων προσωδιακών γνωρισμάτων που χρησιμοποιήθηκαν στην αναλυόμενη έρευνα. Επίσης, χρησιμοποιήθηκαν οι αλγόριθμοι Random Forest, SVM και το Τεχνητό Νευρωνικό Δίκτυο ANN (Artificial Neural Network⁵⁰), το οποίο γίνεται

⁵⁰ [Simply understanding Artificial Neural Networks \(ANN\): Deep Learning | by Jeslur Rahman | Medium](#). (Rahman & Kavlakoglu, 2024).

εκμεταλλεύσιμο για την ταξινόμηση, την πρόβλεψη και την αναγνώριση προτύπων, όπως η επεξεργασία εικόνας (αναγνώριση προσώπου) και η επεξεργασία φυσικής γλώσσας (NLP). Το σύνολο δεδομένων χωρίστηκε σε 75% δεδομένα εκπαίδευσης και 25% δεδομένα δοκιμής. Με την πιστότητα (accuracy) μετρήθηκε κατά πόσο τα αποτελέσματα των αλγορίθμων συμφωνούν με τα αποτελέσματα αξιολογητών. Για το σύνολο των τεσσάρων προσωδιακών γνωρισμάτων, ο Random Forest έχει τα μεγαλύτερα ποσοστά / τιμές πιστότητας, από 70.52% μέχρι 79.76%. Ο SVM έχει κοντινά αλλά χειρότερα ποσοστά, με τον ANN να έρχεται τρίτος και κατά πολύ χειρότερος (55.49% με 67.63%).

Η δυσκολία που αντιμετώπισαν οι ερευνητές ήταν ότι δεν κατέσται δυνατή η αφαίρεση του “θορύβου” και των χαμηλών συχνοτήτων από τα βίντεο. Επομένως, σημαντική παράλειψη των ερευνητών είναι ότι θα μπορούσαν να χρησιμοποιήσουν κάποιο εργαλείο για την αφαίρεση αυτών. Επίσης, οι συνεντεύξεις θα πρέπει να πραγματοποιούνται σε κλειστά και ελεγχόμενα περιβάλλοντα για να μην επηρεάζεται αρνητικά η ποιότητα των δεδομένων. Στο άρθρο, οι συγγραφείς αναφέρονται στην βελτίωση της διαδικασίας που ακολουθήθηκε προς αύξηση της ποιότητας των δεδομένων εκπαίδευσης. Ειδικότερα, για την απόσταση του μικροφώνου από τον συνεντευξαζόμενο, την μη ικανοποιητική απόδοση του υλικού (hardware) της συσκευής εγγραφής του ήχου, τους ήχους του περιβάλλοντος και τις ενδιάμεσες συνομιλίες του αξιολογητή μεταξύ των συνεντεύξεων. Επίσης, το δείγμα των συνεντεύξεων ήταν περιορισμένο και σίγουρα δεν επαρκεί για ασφαλή συμπεράσματα. Όλα τα παραπάνω ενδέχεται να επηρέασαν στο μέγιστο τα αποτελέσματα.

Κάτι που θα μπορούσε να εξελιχθεί στο μέλλον είναι η προσθήκη διαφορετικών διαλέκτων, γλωσσών και παραλλαγών της ανθρώπινης φωνής στο δείγμα, ώστε να μην υπάρχουν δείγματα μόνο με την ίδια διάλεκτο και γλώσσα και να διευρυνθεί η έρευνα. Είναι σχεδόν σίγουρο ότι τα αποτελέσματα θα ήταν διαφορετικά (δηλ. χειρότερα) καθώς κάποιες λέξεις πιθανώς να μην ήταν κατανοητές και να μην εξαγόταν ορθά.

5.2.1.6 Slices of Attention in Asynchronous Video Job Interviews (Κομμάτια Άξια Προσοχής σε Ασύγχρονες Βίντεο-Συνεντεύξεις Εργασίας)

Η μελέτη των Hemamou και συν. (2019) εστιάζει στα μη λεκτικά κοινωνικά σήματα (non-verbal social signals) από ασύγχρονες συνεντεύξεις, με τα οποία η Μηχανική Μάθηση θα συμπεράνει εάν ο υποψήφιος είναι κατάλληλος για πρόσληψη. Αρχικά εντόπισαν και εξήγαγαν τα κομμάτια της συνέντευξης που ξεχώριζαν (attention slices) και τα σύγκριναν με τυχαία κομμάτια (sample slices). Κατέληξαν στο συμπέρασμα ότι τα attention slices όντως φέρουν ουσιώδεις πληροφορίες για πρόσληψη και επικεντρώνονται στις οπτικές ενδείξεις που σχετίζονται με το άγχος και την εναλλαγή συνομιλίας, ενώ τα τυχαία κομμάτια περιέχουν αδιάφορες πληροφορίες. Ειδικά στην αρχή και το τέλος της απάντησης κάθε ερωτώμενου, τα σημεία αυτά είναι πιο έντονα. Τα γνωρίσματα του άγχους σχετίζονται με το ανοιγοκλείσιμο των ματιών, με το σφίξιμο του στόματος κ.α. Τα αποτελέσματα (από την εφαρμογή της έρευνας) δόθηκαν στους υποψηφίους για να βελτιώσουν τη μη λεκτική στρατηγική τους. Τα δεδομένα αποτελούνταν από 7.938 αληθινές συνεντεύξεις Γάλλων σε ασύγχρονα βίντεο. Οι αρμόδιοι αξιολόγησαν τις συνεντεύξεις και κατέληξαν σε δύο βασικές τιμές τελικής αξιολόγησης: “προσλαμβάνεται” και “δεν προσλαμβάνεται”.

Χρησιμοποιήθηκε το νευρωνικό δίκτυο HireNet, το οποίο αναλύει αυτόματα ασύγχρονες βιντεοσκοπημένες συνεντεύξεις. Το HireNet εντοπίζει και αξιολογεί υποψηφίους κάνοντας χρήση τεχνολογιών ανάλυσης δεδομένων με τη Μηχανική Μάθηση. Ειδικότερα, εξετάζει μία συνέντευξη ως ακολουθία ερωταπαντήσεων, εστιάζοντας στα σημαντικότερα κοινωνικά σήματα. Κάποιες από τις λειτουργίες του είναι να αναλύει βιογραφικά εξάγωντας τις σημαντικές πληροφορίες, να συνδυάζει λέξεις που αναφέρονται στα βιογραφικά με λέξεις που υπάρχουν στην περιγραφή θέσης εργασίας, να αντιστοιχίζει τους υποψηφίους με τις κατάλληλες θέσεις εργασίας, καθώς και να τους αξιολογεί και να τους κατατάσσει. Βασικά, αυτοματοποιεί και επιταχύνει τις φάσεις της διαδικασίας πρόσληψης, αποτρέπει τα ανθρώπινα λάθη και αντιστοιχεί με ακρίβεια υποψήφιος με θέσεις εργασίας. Το HireNet χρησιμοποιεί ήχο, βίντεο και λεκτικό περιεχόμενο για να εφαρμόσει όλα τα παραπάνω. Οι συγγραφείς εξήγαγαν το λεκτικό περιεχόμενο χρησιμοποιώντας το Google Speech-to-Text API. Οι ερευνητές αφαίρεσαν τον “θόρυβο” με το Density-Based Spatial Clustering of Applications with Noise (DBSCAN) που ομαδοποιεί δεδομένα και ανιχνεύει τον “θόρυβο”, ώστε να μετρήσουν τα μη λεκτικά γνωρίσματα. Προκειμένου να εντοπίσουν και να εξάγουν τα attention slices των

συνεντεύξεων χρησιμοποίησαν το Gated Recurrent Unit (GRU) που αποκωδικοποιεί τις πληροφορίες. Με το OpenFace εξήγαγαν τα οπτικά χαρακτηριστικά.

Όπως αναφέρθηκε παραπάνω, το HireNet χρησιμοποιεί όλα τα είδη δεδομένων εισόδου (οπτικά, ακουστικά και λεκτικά). Ο μέσος όρος του ποσοστού της F1, για το εάν πρέπει να προσληφθεί κάποιος με τη χρήση του HireNet ανήλθε στο 61.8%. Για τα αποτελέσματα των αλγορίθμων Μηχανικής Μάθησης που αφορούν τα Attention Slices, επέλεξαν την μετρική αξιολόγησης Area Under the Curve (AUC). Τα ποσοστά των αποτελεσμάτων των αλγορίθμων LASSO, Random Forest και SVM σε σύγκριση με τα αποτελέσματα του HireNet, ήταν κατώτερα. Συγκεκριμένα, ο RF πέτυχε 55.4%, ο LASSO 55% και ο SVM 54.3%.

Το άρθρο εξετάζει τη χρήση μη λεκτικών κοινωνικών σημάτων και συγκεκριμένα το κατά πόσο επηρεάζουν τα Attention Slices την επιλογή ενός υποψηφίου για πρόσληψη. Επομένως, αυτό είναι αρκετά καινοτόμο. Ωστόσο, ακόμα και αν με τη χρήση των attention slices η πρόβλεψη είναι καλύτερη από τη χρήση των τυχαίων κομματιών, το τελικό αποτέλεσμα δεν είναι ξεκάθαρο, δηλαδή αν κάποιος “προσλαμβάνεται” ή “δεν προσλαμβάνεται” σύμφωνα με τα attention slices. Όλα τα ποσοστά των αλγορίθμων Μηχανικής Μάθησης βρίσκονται κοντά στο 50% και αυτό δεν μπορεί να καθορίσει στο ελάχιστο αν κάποιος πρέπει να προσληφθεί ή όχι. Μικρό ποσοστό θεωρείται και το 61.8% με τη χρήση του HireNet για να εφαρμοστεί από πραγματικούς εργοδότες. Η χαμηλή ακρίβεια πρόβλεψης περιορίζει τη χρησιμότητα του συστήματος και πρέπει να βελτιωθεί προκειμένου να λειτουργήσει με μεγαλύτερα ποσοστά. Το δείγμα δεδομένων από 7.938 πραγματικές ασύγχρονες συνεντεύξεις ήταν άκρως ικανοποιητικό και παρέχει ένα αξιόπιστο και αντιπροσωπευτικό δείγμα προς ανάλυση. Βέβαια, θα ήταν αντιπροσωπευτικότερο εάν συμμετείχαν άτομα και από άλλες χώρες.

5.2.1.7 Developing and Evaluating Language-Based Machine Learning Algorithms for Inferring Applicant Personality in Video Interviews (Ανάπτυξη και Αξιολόγηση Αλγορίθμων Μηχανικής Μάθησης Βασισμένων στη Γλώσσα ως προς την Εξαγωγή της Προσωπικότητας του Αιτούντος σε Βίντεο Συνεντεύξεις)

Οι Hickman και συν. (2021) ανέπτυξαν αλγόριθμους Μηχανικής Μάθησης για εξαγωγή συμπερασμάτων της προσωπικότητας του αιτούντος. Με αυτόν τον τρόπο, οι αλγόριθμοι ήταν σε θέση να συμπεράνουν αυτόματα τις γνώσεις, τις ικανότητες και άλλα χαρακτηριστικά προσωπικότητας των υποψηφίων. Το δείγμα των εικονικών βίντεο συνεντεύξεων είναι 441 και πάρθηκε από φοιτητές ψυχολογίας. Αυτοί που ορίστηκαν αξιολογητές, ήταν προπτυχιακοί βοηθοί ερευνητές. Και σε αυτή την έρευνα οι αλγόριθμοι προβλέπουν, κατηγοριοποιούν και ταξινομούν τα “Πέντε Μεγάλα χαρακτηριστικά προσωπικότητας” (Big Five personality traits)⁵¹. Αυτό που επιδιώκουν οι συγγραφείς είναι να μπορούν να αποκωδικοποιούν οι αλγόριθμοι και οι εφαρμογές (MM) όχι μόνο το τί λέει κάποιος, αλλά και το πώς το λέει.

Με το IBM Watson Speech-to-Text οι συγγραφείς μετέτρεψαν την ομιλία σε κείμενο. Με το Linguistic Inquiry Word Count (LIWC), κατηγοριοποίησαν τα λεξιλογικά χαρακτηριστικά που περιγράφουν τα συναισθήματα συσχετίζοντάς τα με τα Πέντε Μεγάλα χαρακτηριστικά προσωπικότητας. Με το πακέτο R-CARET εκπαιδύσαν και δοκίμασαν μοντέλα Μηχανικής Μάθησης με προσεγγίσεις ελαστικής παλινδρόμησης (elastic net regression) για να υπολογίσουν την ακρίβεια της πρόβλεψης. Οι αλγόριθμοι αυτοί είναι οι: Ridge Regression και LASSO. Το αποτέλεσμα ήταν, ο μέσος όρος πρόβλεψης των χαρακτηριστικών προσωπικότητας από το σύστημα σε σχέση με τους αξιολογητές ($r=0,39$). Αυτό δείχνει μία μέτρια πρόβλεψη των αλγορίθμων σε σχέση με την πρόβλεψη των αξιολογητών. Το καλύτερο μεμονωμένο αποτέλεσμα ήταν $r=0.49$ και αφορούσε την εξωστρέφεια. Οι πιο ισχυρές προβλέψεις ($r=0,45$) ήταν οι συσχετίσεις των Πέντε Μεγάλων χαρακτηριστικών με τις προβλέψεις του LIWC σε δύο σημεία, στην εξωστρέφεια και στον αριθμό λέξεων που χρησιμοποιούν οι συνεντευξιαζόμενοι. Η μετρική r είναι ο συντελεστής συσχέτισης του Pearson και χρησιμοποιείται για την αξιολόγηση συσχέτισης μεταξύ μεταβλητών. Το αποτέλεσμα κυμαίνεται από -1 έως και +1, όπου το -1 είναι η τέλεια αρνητική συσχέτιση, το 0 καμία συσχέτιση και το +1 η τέλεια θετική συσχέτιση. Οι ερευνητές δεν αναφέρονται σε κάποια μετρική ακρίβειας, παρά μόνο στον συντελεστή συσχέτισης του Pearson.

⁵¹ Εξωστρέφεια (extraversion), συγκαταβατικότητα (agreeableness), διαφάνεια (openness), ευσυνειδησία (conscientiousness) και νευρωτισμός/συναισθηματική σταθερότητα (neurotism/emotional stability).

Το αποτέλεσμα της έρευνας δεν υπήρξε ικανοποιητικό, καθώς οι προβλέψεις συσχετίσεων των αλγορίθμων με τις προβλέψεις των αξιολογητών υπήρξαν μέτριες. Η μικρή διάρκεια των συνεντεύξεων (6 λεπτά) επηρέασε επίσης το μέτριο αποτέλεσμα των αλγορίθμων καθώς δεν επαρκεί για την σωστή εκπαίδευση και δοκιμή των αλγορίθμων. Έρευνες έδειξαν ότι για σωστή εκπαίδευση και δοκιμή των αλγορίθμων χρειάζονται περίπου είκοσι λεπτά συνέντευξης. Οπότε με την σωστή διάρκεια συνεντεύξεων, το ποσοστό της ακρίβειας της πρόβλεψης αναμένεται να αυξηθεί. Οι αλγόριθμοι πλέον εκπαιδεύονται όχι μόνο για να κατανοούν το περιεχόμενο ενός κειμένου ή μίας ομιλίας (τι λέει κάποιος), αλλά και για τον τρόπο που εκφράζεται κάποιος (πώς το λέει), κάτι που ενισχύει την κατανόηση της προσωπικότητας. Αυτό που παρατηρήθηκε είναι ότι οι αλγόριθμοι προέβλεψαν καλύτερα τα εξαγόμενα χαρακτηριστικά των ανδρών έναντι των γυναικών. Ακόμα ένα θέμα που επηρέασε αρνητικά την ακρίβεια πρόβλεψης αφορά το γεγονός πως η γλώσσα συνέντευξης δεν ήταν η μητρική για ορισμένους συνεντευξιάζοντες. Η έρευνα περιορίζεται στη χρήση του συντελεστή συσχέτισης Pearson και δεν παρέχει άλλες μετρικές που ενδεχομένως να ήταν απαραίτητες για καλύτερη και πιο ολοκληρωμένη αξιολόγηση και αποτελεσματικότητα του συστήματος.

5.2.1.8 Bias and Fairness in Multimodal Machine Learning: A Case Study of Automated Video Interviews (Μεροληψία και Δικαιοσύνη στην Πολυτροπική Μηχανική Μάθηση: Μια Μελέτη Περίπτωσης Αυτοματοποιημένων Βίντεο-Συνεντεύξεων)

Οι Booth και συν. (2021) εξερευνούν πόσο μπορεί να είναι δίκαιη και χωρίς προκαταλήψεις η αξιολόγηση για την πρόσληψη εργαζόμενων. Η πρωτοτυπία εδώ έγκειται στο ότι αφαιρούν τα χαρακτηριστικά που προσδιορίζουν το φύλο για αμεροληψία. Η εξαγωγή των λεκτικών, παραλεκτικών και οπτικών χαρακτηριστικών, η αφαίρεση των στοιχείων του φύλου και στη συνέχεια η μέτρησή τους για πρόβλεψη πρόσληψης είναι αυτό που δοκιμάζεται στην παρούσα έρευνα. Παραδείγματα παραλεκτικών χαρακτηριστικών είναι η φωνή και πώς αυτή εκφράζεται μέσα από τον τόνο της, τον ήχο της, τον τονισμό των λέξεων, την ένταση και τον ρυθμό ομιλίας. Οι αρχικές βιντεοσκοπημένες εικονικές συνεντεύξεις πραγματοποιήθηκαν από 733 φοιτητές και ο αριθμός τους ήταν 4255. Ο τελικός αριθμός των συνεντεύξεων που

χρησιμοποιήθηκαν ήταν 524, διότι υπήρχαν 465 γυναίκες, 262 άντρες και 6 non-binary. Οπότε, προκειμένου να υπάρχει ισορροπία στα φύλα, οι συγγραφείς αφαίρεσαν τα non-binary και επέλεξαν 262 γυναίκες και 262 άντρες. Όλες οι συνεντεύξεις αξιολογήθηκαν από εκπαιδευμένους βαθμολογητές.

Η εξαγωγή των λεκτικών στοιχείων έγινε από το Linguistic Inquiry Word Count (LIWC), των παραλεκτικών ή προσωδιακών από το OpenSmile και των οπτικών από το λογισμικό Motion Tracker για τα γνωρίσματα του προσώπου και του άνω μέρους του σώματος, ενώ το Emotient FACET SDK χρησιμοποιήθηκε για τα στοιχεία συναισθήματος και τις εκφράσεις του προσώπου. Εντοπίστηκαν 5.653 λεκτικά γνωρίσματα, 125 παραλεκτικά και 250 οπτικά. Δοκιμάστηκε ο Random Forest αρχικά στα γνωρίσματα που δεν χαρακτηρίζουν το φύλο και ύστερα σε όλα. Επιπλέον, δοκιμάστηκε ξεχωριστά σε κάθε ένα από τα τρία χαρακτηριστικά (λεκτικά, παραλεκτικά και οπτικά) και σε όλα μαζί. Η μετρική που χρησιμοποιήθηκε είναι η πιστότητα συσχετισμού (correlational accuracy) για να υπολογιστεί η προκατάληψη γενικά, αλλά και του φύλου. Αυτή μετρήθηκε από τον Συντελεστή Συσχέτισης Spearman (Spearman's rank correlation coefficient), ο οποίος αξιολογεί τη σχέση μεταξύ δύο μεταβλητών και χρησιμοποιείται για να προσδιοριστεί εάν δύο μεταβλητές σχετίζονται μεταξύ τους και με ποιον τρόπο (θετικά ή αρνητικά). Και τα δύο καλύτερα αποτελέσματα ήταν .46 με τον συνδυασμό των τριών χαρακτηριστικών, το πρώτο για την πρόβλεψη της ακρίβειας των γνωρισμάτων με βάση αυτά των αξιολογητών και το δεύτερο για το κατά πόσο υπάρχει προκατάληψη λόγω φύλου.

Αυτό που εξετάστηκε στην έρευνα, είναι η προσπάθεια για δίκαιη και χωρίς προκατάληψη πρόσληψη του εργαζομένου, με την αφαίρεση των χαρακτηριστικών του φύλου. Αυτό, διασφαλίζει την αποφυγή μηνύσεων κατά των εργοδοτών και παράλληλα ενισχύει τη διαφάνεια και τη δικαιοσύνη. Η χρήση πολυτροπικών δεδομένων προσφέρει μία ολιστική προσέγγιση. Η εξισορρόπηση των φύλων, με την αφαίρεση αρκετών συνεντεύξεων, συνέβαλε στη μείωση πιθανών και πάλι προκαταλήψεων στα αποτελέσματα. Από την άλλη, η χρήση μόνο 524 συνεντεύξεων από τις 4.255, είναι απογοητευτική. Θα μπορούσαν από την αρχή να εξασφαλίσουν πιο ισορροπημένες συνεντεύξεις ώστε να μην χρειαστεί να αφαιρεθεί τόσο μεγάλος αριθμός δεδομένων. Η χαμηλή πιστότητα συσχετισμού, έδειξε ότι υπάρχει μεγάλο

περιθώριο βελτίωσης στη διαδικασία. Ενδιαφέρον θα είχε να συμπεριλάμβαναν στην έρευνα χαρακτηριστικά προκατάληψης, όπως η ηλικία, η εθνικότητα και η θρησκεία.

5.2.1.9 Interview Based Human Behavior Analysis with Machine Learning (Ανάλυση Ανθρώπινης Συμπεριφοράς Βασιζόμενη σε Συνέντευξη μέσω Μηχανικής Μάθησης)

Οι Anju Raj και Farzana (2023) παρουσιάζουν ένα σύστημα που αναλύει τις εκφράσεις του προσώπου, τα γλωσσικά και τα προσωδιακά στοιχεία των υποψηφίων από βιντεοσκοπημένες συνεντεύξεις. Το δοθέν ερωτηματολόγιο περιλαμβάνει επτά ερωτήσεις και η κάθε μία σχετίζεται με κάποιο χαρακτηριστικό προσωπικότητας. Τα χαρακτηριστικά αφορούσαν την ηρεμία, τον ενθουσιασμό, την συγκέντρωση, την φιλικότητα και το άγχος. Αυτά τα χαρακτηριστικά κρίθηκαν από έξι έμπειρους επαγγελματίες αξιολογητές και ο μέσος όρος των αξιολογήσεών τους θεωρείται η σωστή κρίση. Το σύστημα προβλέπει την απόδοση του υποψήφιου στα συμπεριφορικά γνωρίσματά του, την προσωπικότητά του και αν είναι κατάλληλος για πρόσληψη. Σημειώνεται πως η μη λεκτική επικοινωνία είναι υποσυνείδητη και είναι δύσκολο να την διαχειριστεί κάποιος. Σημαντικό ρόλο παίζουν όχι μόνο οι λέξεις και οι κινήσεις, αλλά και το παρουσιαστικό και το ντύσιμο.

Στο πείραμα που επιτελέστηκε, οι συγγραφείς διεξήγαγαν εικονικές συνεντεύξεις για μία πραγματική θέση εργασίας. Μία κάμερα συνέλεγε τα οπτικά στοιχεία και ένα μικρόφωνο τα λεκτικά. Μέσω μιας διεπαφής χρήστη (user interface), τα άτομα που έπαιρναν τη συνέντευξη είχαν τη δυνατότητα να βλέπουν σε πραγματικό χρόνο το smile score variation graph. Το τελευταίο είναι ένα γράφημα που οι τιμές του μεταβάλλονται ανάλογα με την ένταση του χαμόγελου του υποψηφίου.

Η εξαγωγή των προσωδιακών χαρακτηριστικών και η αφαίρεση του “θορύβου” έγινε με το PRAAT με χρήση της pyaudio βιβλιοθήκης. Οι εκφράσεις και τα χαρακτηριστικά του προσώπου ανιχνεύθηκαν και ταξινομήθηκαν από το Haar Cascade, μία μέθοδο που αναγνωρίζει αντικείμενα και εξειδικεύεται στην αναγνώριση προσώπων και των χαρακτηριστικών τους. Το μοντέλο Markov χρησιμοποιήθηκε για την εξαγωγή του χρόνου που κοιτούσε ο υποψήφιος τον αξιολογητή. Το Valence Aware Dictionary and sentiment Reasoner (VADER) είναι ένα εργαλείο ανάλυσης συναισθήματος (sentiment analysis) που έγινε εκμεταλλεύσιμο από την έρευνα. Χρησιμοποιεί

συγκεκριμένο λεξικό (sentiment lexicon), το οποίο περιέχει λέξεις και φράσεις που εκφράζουν το θετικό (π.χ. χαρά, καλοσύνη), αρνητικό (π.χ. θυμός, λύπη) ή ουδέτερο συναίσθημα που προκαλεί η κάθε λέξη.

Τα δεδομένα που προέκυψαν από τις προαναφερόμενες εξαγωγές, συγκρίθηκαν με την αξιολόγηση των βαθμολογητών. Με το πέρας της συνέντευξης και αφού είχαν αξιολογήσει, οι αξιολογητές έβλεπαν στη διεπαφή χρήστη τις προβλέψεις του συστήματος, για να μην επηρεαστούν από αυτές. Σημειώνεται πως με βάση τη προτεινόμενη προσέγγιση, αν το επέλεγαν, οι αξιολογητές μπορούσαν να βλέπουν τη στιγμή της συνέντευξης τον γράφο smile score variation, σε πραγματικό χρόνο. Οι βασικές παράμετροι ήταν η επαφή βλέμματος, η ηρεμία, ο ενθουσιασμός, η συγκέντρωση, η φιλικότητα κ.α.

Οι Anju Raj και Farzana (2023) παρουσιάζουν το σύστημά τους, όμως δεν καταγράφουν τα ποσοστά ακρίβειας της πρόβλεψης. Αυτό που έχει σημασία για αυτούς, είναι η παρουσίαση της πρότασής τους και όχι το αποτέλεσμα αυτής. Η ολιστική προσέγγισή τους, συνδυάζει τα λεκτικά, τα προσωδιακά και τα μη λεκτικά χαρακτηριστικά, κάτι που αναδεικνύει την σφαιρική κατανόηση της ανθρώπινης επικοινωνίας. Η βασική καινοτομία της προσέγγισής τους έγκειται στην προσφερόμενη διεπαφή χρήστη, την οποία δύνανται να βλέπουν οι αξιολογητές τη στιγμή της συνέντευξης για να παρατηρούν άμεσα τις αντιδράσεις των υποψηφίων. Αν και ο γράφος δεν εμφανίζει προβλέψεις, παρόλα αυτά, μπορεί να επηρεαστεί η κρίση τους από αυτόν και να επηρεαστεί η αντικειμενικότητά τους. Τέλος, θα μπορούσε το ερωτηματολόγιο να περιέχει περισσότερες από μία ερωτήσεις για κάθε χαρακτηριστικό προσωπικότητας, ώστε τα αποτελέσματα να αποδίδουν καλύτερα τα χαρακτηριστικά αυτά.

5.2.2 Ανάλυση Άρθρων Εστιαζόμενων στα Τυπικά Προσόντα

5.2.2.1 Resume Parser and Summarizer (Αναλυτής και Συνοψιστής Βιογραφικών Σημειωμάτων)

Οι Farzana και συν. (2023) προτείνουν ένα σύστημα, το οποίο μπορεί να εξάγει τις χρήσιμες πληροφορίες ενός βιογραφικού σημειώματος και να τις αναλύει, να τις αποθηκεύει, να τις οργανώνει και να τις ταξινομεί. Το σύστημα αυτό ονομάζεται «αναλυτής βιογραφικών» (resume parser) και μπορεί να μετατρέπει τα αποδομημένα δεδομένα (όπως ένα βιογραφικό σημείωμα), σε δομημένα. Έπειτα, μπορεί να κατηγοριοποιεί αυτόματα στοιχεία, όπως την εργασιακή εμπειρία, τις επαγγελματικές πιστοποιήσεις, τις δεξιότητες, τις γνώσεις και τα στοιχεία επαφής. Η πιστότητα (accuracy) τέτοιων συστημάτων είχε φτάσει έως και το 87% για την ορθή εισαγωγή των δεδομένων και την ορθή κατηγοριοποίησή τους κατά την ημερομηνία δημοσίευσης του άρθρου σε προηγούμενες έρευνες από άλλους ερευνητές, όπως αναφέρεται στο άρθρο. Το ποσοστό είναι αρκετά υψηλό αν αναλογιστεί κανείς ότι η πιστότητα (accuracy) στους ανθρώπους δεν ξεπερνά το 96%.

Το προτεινόμενο σύστημα των Farzana και συν. (2023), το οποίο θα βρίσκεται σε δική τους ιστοσελίδα, θα δημιουργεί δοκιμαστικά μία βάση δεδομένων από βιογραφικά σημειώματα που θα αντλεί από ιστοσελίδες, όπως το Kaggle⁵² και το GitHub⁵³. Μέσω αυτού του συστήματος είναι δυνατόν να εγγράφονται τόσο οι αιτούντες εργασία όσο και οι εταιρείες που αναζητούν υπαλλήλους. Οι αιτούντες θα ανεβάζουν το βιογραφικό τους σε οποιαδήποτε μορφή θέλουν, όπως pdf, text, docx, rift και html. Έπειτα, το σύστημα θα δομεί και θα αποθηκεύει τις πληροφορίες στην βάση δεδομένων σε μορφή json. Επιπλέον, οι υπεύθυνοι προσλήψεων θα μπορούν να αναζητούν βιογραφικά σύμφωνα με τις προτιμήσεις τους (λέξεις κλειδιά), αφού το σύστημα θα τα έχει ήδη εξάγει, κατηγοριοποιήσει και ταξινομήσει.

Τα εργαλεία και η τεχνολογία που χρησιμοποιήθηκαν για την υλοποίηση του αναλυόμενου συστήματος είναι: η γλώσσα προγραμματισμού Python, η MongoDB που είναι ένα σύστημα διαχείρισης NoSQL βάσεων δεδομένων, η HTML (Hyper Text Markup Language) και η CSS (Cascading Style Sheets) για την εμφάνιση της διεπαφής χρήστη του συστήματος, το πλαίσιο Flask για την υλοποίηση του backend του συστήματος, η τεχνολογία NLP (Natural Language Processing) που αποτελεί διεπιστημονικό υποτομέα της γλωσσολογίας, της επιστήμης των υπολογιστών και της τεχνητής νοημοσύνης, ο οποίος ασχολείται με την αλληλεπίδραση των υπολογιστών και της ανθρώπινης γλώσσας και τέλος η SpaCy, η οποία είναι μια

⁵² [Kaggle: Your Machine Learning and Data Science Community](#)

⁵³ [GitHub · Build and ship software on a single, collaborative platform · GitHub](#)

προηγμένη βιβλιοθήκη NLP που χρησιμοποιείται για την ανίχνευση των σημαντικών πληροφοριών των κειμένων.

Το νέο αυτό σύστημα, από την πλευρά των εταιριών, θα βοηθήσει στην γρήγορη ανάλυση και εύρεση των πλέον κατάλληλων υποψηφίων, θα βελτιώσει την αποδοτικότητα της διαδικασίας και θα εξαλείψει κατά το δυνατόν τα ανθρώπινα λάθη. Επίσης, η λίστα κατάταξης των υποψηφίων, η οποία προκύπτει από τα κριτήρια που δίνουν οι εταιρείες, είναι κάτι που θα διευκολύνει τη λήψη απόφασης. Από την πλευρά των αιτούντων, είναι εύχρηστο καθώς θα υποστηρίζει το ανέβασμα όλων των μορφών κειμένου.

Η ασφάλεια της ιστοσελίδας είναι κάτι που πρέπει να προβληματίσει σοβαρά τους κατασκευαστές της επειδή πρέπει να διασφαλίζεται/προστατεύεται η ιδιωτικότητα τόσο των αιτούντων, όσο και των εταιρειών. Το σύστημα που παρουσιάζουν οι συγγραφείς φαίνεται ότι θα είναι χρήσιμο για όλους τους ενδιαφερόμενους, ωστόσο θα ήταν προτιμότερο να εφαρμοστεί και αξιολογηθεί πριν την δημοσίευση της σχετικής έρευνα, καθώς κατά την λειτουργία του συστήματος ενδέχεται να προέκυπταν προβλήματα που δεν είναι σε θέση να γνωρίζουν πριν την εφαρμογή του και μπορεί να υπήρχαν ριζικές αλλαγές στη διαδικασία.

5.2.2.2 Resume Ranking Using Natural Language Processing (Κατάταξη Βιογραφικού με τη Χρήση Επεξεργασίας Φυσικής Γλώσσας)

Οι Naveed και συν. (2024) παρουσιάζουν ένα μοντέλο που κατατάσσει τους υποψήφιους για κάθε κατηγορία θέσης μέσω πλατφόρμας που θα βρίσκεται σε ιστοσελίδα τους και θα διευκολύνει υπηρεσίες που δέχονται μεγάλο ογκο βιογραφικών. Ο υποψήφιος θα μεταφορτώνει το βιογραφικό του στην ιστοσελίδα τους. Με τη χρήση της επεξεργασίας της φυσικής γλώσσας (NLP), που αναλύεται παρακάτω, θα ταξινομούνται τα προσόντα από τα μη δομημένα βιογραφικά. Η πλατφόρμα τους θα είναι προσπελάσιμη από τους υποψήφιους για να βρουν τις εργασίες που τους ενδιαφέρουν και από τις εταιρείες ή τις κυβερνήσεις για να βρίσκουν κατάλληλα βιογραφικά, μετά από εγγραφή τους στην ιστοσελίδα. Οι τελευταίοι θα λαμβάνουν τα αποτελέσματα σε αρχεία κειμένου json. Οι εταιρείες θα δημοσιεύουν τις κενές θέσεις εργασίας (δεν αναφέρεται ρητά σε ποιά μορφή, αλλά

υπονοείται σε μορφή pdf όπως οι υποψήφιοι) και οι υποψήφιοι τα βιογραφικά τους σε μορφή pdf. Το αποτέλεσμα που θα λαμβάνουν οι εταιρείες θα ταξινομείται σύμφωνα με τα προσόντα και τα χρόνια εμπειρίας των υποψηφίων. Η έρευνα πραγματοποιήθηκε με είκοσι βιογραφικά.

Το σύστημα λογισμικού που χρησιμοποίησαν, αναπτύχθηκε μέσω του Visual Studio Code, το οποίο είναι επεξεργαστής κώδικα της Microsoft, και της στοίβας MERN. Η τελευταία χρησιμοποιείται για την ανάπτυξη διαδικτυακών εφαρμογών και περιλαμβάνει τις εξής τεχνολογίες: τη MongoDB NoSQL βάση δεδομένων, το πλαίσιο React JS για την ανάπτυξη του front-end, το Express.js για την ανάπτυξη του backend και το Node JS ως το περιβάλλον εκτέλεσης και των 2 βασικών μερών μιας διαδικτυακής εφαρμογής.

Πριν την ανάλυση, εξαγωγή και ταξινόμηση των χαρακτηριστικών, πρέπει να ακολουθηθούν συγκεκριμένα βήματα. Η διαδικασία διαχωρισμού του κειμένου σε μικρότερα τμήματα (tokenization), το Stop Word Removal που αφαιρεί λέξεις όπως άρθρα και αντωνυμίες και ο ορθογραφικός έλεγχος. Οι ερευνητές δίνουν έμφαση στη μορφολογική, συντακτική και σημασιολογική ανάλυση και χρησιμοποίησαν τα παρακάτω εργαλεία και εφαρμογές για την επεξεργασία τους. Η ανάλυση του κειμένου και η κατηγοριοποίηση των λέξεων πραγματοποιείται από το Natural Language Toolkit (NLTK). Για την ανάλυση των συναισθημάτων και την μετατροπή του κειμένου χρησιμοποιείται το Text Blob, ένα εργαλείο της Python. Στην αναγνώριση ονομαστικών και αριθμητικών οντοτήτων χρησιμοποίησαν τη βιβλιοθήκη SpaCy για εκπαίδευση μοντέλου τύπου Named Entity Recognition (NER), το οποίο προβλέπει τις οντότητες.

Για την εκπαίδευση του NER από το SpaCy, χρησιμοποιήθηκε η online πλατφόρμα Dataturks, με την οποία εντοπίζονται κείμενα, λέξεις, εικόνες και μπορεί να χρησιμοποιηθεί για την δημιουργία καθαρών και δομημένων δεδομένων εκπαίδευσης. Τα δεδομένα που χρησιμοποίησαν για την εκπαίδευση προήλθαν από τα είκοσι βιογραφικά. Σε αυτά εφάρμοσαν τη διαδικασία που περιγράφηκε παραπάνω. Οι μετρικές F1 Score, Recall και Precision σχεδόν άγγιξαν το 100% στην πρόβλεψη των σωστών οντοτήτων, μέσω των οποίων πραγματοποιείται το σωστό ταίριασμα του βιογραφικού με τη θέση εργασίας (train set). Οι σωστή πρόβλεψη των οντοτήτων συνεπάγεται και τη σωστή κατάταξη των υποψηφίων. Οι ερευνητές

επεξεργάστηκαν⁵⁴ 23.499 φράσεις, από τις οποίες οι σωστές προβλέψεις ήταν 23.317. Ενώ, οι ίδιες μετρικές για το δείγμα δοκιμής (test set) κυμάνθηκαν στο 90% με 5.252 σωστές προβλέψεις φράσεων έναντι 5.942.

Το μοντέλο που παρουσιάστηκε, μπορεί να διευκολύνει τη διαδικασία πρόσληψης σε διαγωνισμούς όλου του δημοσίου, αλλά και ανά υπουργείο ή υπηρεσία. Με αυτό τον τρόπο, ο χρόνος επιλογής εργαζόμενων θα μειωθεί κατά πολύ. Αυτό που δεν διευκρινίστηκε είναι το ποσό που θα πληρώνουν είτε οι αιτούντες, είτε οι εταιρείες ή το κράτος. Για να συμμετέχει ο τομέας του δημοσίου στην πλατφόρμα θα πρέπει να θωρακιστεί έτσι ώστε να εξασφαλιστεί η ασφάλεια και ιδιωτικότητα των δεδομένων. Οι ερευνητές δίνουν έμφαση στην παρουσίαση του συστήματός τους και δεν επικεντρώνονται στη σωστή εκπαίδευση του, αφού ο αριθμός των βιογραφικών που χρησιμοποιούν είναι μόνο είκοσι. Επίσης, δεν αναφέρονται συγκεκριμένοι αλγόριθμοι Μηχανικής Μάθησης που χρησιμοποιήθηκαν για τις προβλέψεις. Επίσης, θα μπορούσαν να ορίσουν και άλλα κριτήρια ως προς την κατάταξη των υποψηφίων, όπως και την βαρύτητα του καθενός. Τέλος, θα πρέπει να επιτρέπουν να επιλέγει η κάθε υπηρεσία τα κριτήρια που προτιμά ανάλογα με τη θέση που δημοσιεύει/προκηρύσσει.

5.2.2.3 Smart Resume Analyzer (Έξυπνος Αναλυτής Βιογραφικών)

Οι Sowjanya και συν. (2023) προτείνουν ένα σύστημα, το οποίο περιλαμβάνει έναν αυτοματοποιημένο αλγόριθμο βασισμένο στη Μηχανική Μάθηση, που φιλοδοξούν να υιοθετηθεί από πολυεθνικές εταιρείες, κυβερνητικούς οργανισμούς και γενικά διοικητικές υπηρεσίες, οι οποίες εξετάζουν σχεδόν σε καθημερινή βάση μεγάλο όγκο βιογραφικών. Σκοπός τους είναι το σύστημα αυτό να αναγνωρίζει, επιλέγει και προτείνει στους εργοδότες τους καλύτερους υποψηφίους προς πρόσληψη, αλλά και να ενημερώνει τους υποψηφίους για τα αδύναμα στοιχεία του βιογραφικού τους και τη συνολική τους βαθμολογία. Οι υποψήφιοι υποβάλουν τα βιογραφικά τους σε μορφή PDF, μετά την εγγραφή τους στο σύστημα και για συγκεκριμένες θέσεις

⁵⁴ Χρησιμοποιήθηκαν οι εξής κατηγορίες για την αναγνώριση και ταξινόμηση οντοτήτων (NER): LOC για την εύρεση τοποθεσιών (location), PER για ονόματα (person), ORG για ονομασίες οργανισμών, εταιρειών (organization) και MISC για όλες τις υπόλοιπες κατηγορίες (miscellaneous).

εργασίας. Έπειτα, οι ενδιαφερόμενες υπηρεσίες λαμβάνουν προτάσεις για τους καλύτερους υποψήφιους ως προς την πλήρωση κάθε κενής θέσης εργασίας.

Αρχικά, το νευρωνικό δίκτυο ResumeNet αξιολογεί αυτόματα την ποιότητα των βιογραφικών, ενσωματώνοντας τεχνικές επεξεργασίας κειμένου για την πρόβλεψη της ποιότητας κάθε βιογραφικού. Στην περίπτωση που ο υποψήφιος επιλεγθεί από τον εργοδότη, θα του προτείνονται βίντεο προετοιμασίας συνέντευξης στο YouTube. Στην περίπτωση που το βιογραφικό του δεν είναι ποιοτικά καλό, θα του προτείνονται βίντεο για την προετοιμασία του βιογραφικού. Για την εκπαίδευση του συγκεκριμένου συστήματος συγκεντρώθηκαν πενήντα βιογραφικά από το Kaggle σε μορφή DOC και DOCX, που μετατράπηκαν σε PDF.

Το σύστημα αποθηκεύει τα δεδομένα σε μορφή json για να τα επεξεργαστεί. Το προτεινόμενο σύστημα, χρησιμοποιεί τεχνικές της Επεξεργασίας της φυσικής γλώσσας (NLP) για την εξαγωγή των κατάλληλων στοιχείων από τα βιογραφικά σημειώματα. Ειδικότερα, χρησιμοποιείται και το SpaCy, λογισμικό/βιβλιοθήκη για την ανίχνευση των σημαντικών πληροφοριών των κειμένων. Μετά την εξαγωγή των στοιχείων, αφαιρούνται τα μη χρήσιμα μέσω Stop Word Removal. Το stemming και το lemmatization περικόπτουν γράμματα από τις λέξεις μέχρι αυτές να φθάσουν στη βασική τους μορφή (ρίζα). Οι διαφορές του stemming και του lemmatization είναι ότι το πρώτο είναι πιο γρήγορο και συχνά καταλήγει σε μη φυσικές λέξεις, ενώ το δεύτερο διατηρεί λεξιλογικά τη σωστή μορφή τους και είναι περισσότερο ακριβές σημασιολογικά. Παρατηρείται ότι στα άρθρα που εξετάζονται στην παρούσα διπλωματική εργασία, γίνεται χρήση και των δύο. Το πρώτο για την γρήγορη ανάλυση και το δεύτερο για την ακρίβεια. Ο συνδυασμός των δύο φαίνεται να λειτουργεί πιο αποτελεσματικά. Η μέθοδος “chunking” μειώνει την έκταση των βιογραφικών, δημιουργώντας μία μορφή μικρής περίληψης. Το TF-IDF (Term Frequency-Inverse Document Frequency) είναι μέθοδος που χρησιμοποιήθηκε για την παραγωγή διανυσμάτων λέξεων που αναπαριστούν προτάσεις ή ολόκληρα έγγραφα. Το Cosine Similarity, εντόπισε τις ομοιότητες μεταξύ δύο τέτοιων διανυσμάτων, όπου το ένα μπορεί να αφορά ένα βιογραφικό υποψηφίου και το δεύτερο μία περιγραφή θέσης εργασίας. Η μέθοδος EDA (Exploratory Data Analysis - Εξερευνητική Ανάλυση Δεδομένων) χρησιμοποιήθηκε για να αναλύσει τα δεδομένα, να βρει ανωμαλίες, μοτίβα, συσχετίσεις, να εντοπίσει ελλείψεις ή προβλήματα και τέλος για να τα οπτικοποιήσει με γραφήματα και στατιστικά για την εύκολη

κατανόησή τους. Η χρήση του αλγόριθμου Latent Dirichlet Allocation (LDA), μέσω της βιβλιοθήκης Gensim, εντόπισε τις κρυφές συσχετίσεις και τα θέματα σε μεγάλο όγκο δεδομένων και δημιούργησε περιλήψεις για τα βιογραφικά και για τις περιγραφές θέσεων εργασίας. Χρησιμοποιήθηκε για την ακόμα καλύτερη συσχέτιση βιογραφικού με θέση εργασίας και ταυτόχρονα δημιούργησε μία λίστα με κατάταξη των υποψηφίων από την υψηλότερη μέχρι την χαμηλότερη βαθμολογία, σύμφωνα με τη σειρά συνάφειας με την κενούμενη θέση εργασίας. Οι συστάσεις, βασισμένες στο περιεχόμενο (Content-based Recommendations), διευκόλυναν στην καλύτερη αντιστοίχιση των βιογραφικών με τις θέσεις εργασίας. Για να βρεθούν τα βιογραφικά που βρίσκονται να ταιριάζουν περισσότερο στην περιγραφή θέσης εργασίας, χρησιμοποιήθηκε ο αλγόριθμος MM KNN. Το προτεινόμενο σύστημα έφερε 85% πιστότητα ανάλυσης κειμένου (text parsing accuracy) και πιστότητα ανάλυσης (ranking accuracy) 92%. Το κείμενο δεν αναφέρει τα ποσοστά που χρησιμοποιήθηκαν για δεδομένα εκπαίδευσης (train set) και για το δείγμα δοκιμής (test set).

Η χρήση της μεθόδου Exploratory Data Analysis (EDA) με την προσθήκη του Latent Dirichlet Allocation (LDA) σε συνδυασμό με τον αλγόριθμο KNN όπως περιγράφηκε παραπάνω, φαίνεται ότι λειτούργησε ενισχυτικά για την καλύτερη αντιστοίχιση της θέσης εργασίας με τα βιογραφικά. Η δωρεάν και επί πληρωμή πρόταση σεμιναρίων και πιστοποιήσεων για την εξέλιξη των γνώσεων και της επάρκειας των υποψηφίων είναι κάτι καινοτόμο, όπως και οι συμβουλές για τη δημιουργία του βιογραφικού και η προετοιμασία συνέντευξης μέσω βίντεο στο YouTube. Οι εξατομικευμένες συμβουλές ενισχύουν την μελλοντική ετοιμότητα των υποψηφίων. Επίσης, η οπτικοποίηση των δεδομένων μέσω γραφημάτων και στατιστικών διευκολύνει την κατανόηση των δεδομένων, ώστε οι εργοδότες να λάβουν την απόφασή τους γρηγορότερα. Τέλος, η υψηλή ακρίβεια του συστήματος έδειξε ότι ήταν αποτελεσματικό.

Παρόλα αυτά, ο μικρός αριθμός των βιογραφικών (50) για ακόμα μία φορά είναι ανεπαρκής. Με έναν μεγαλύτερο αριθμό βιογραφικών οι αλγόριθμοι θα εκπαιδεύονταν καλύτερα. Ακολούθως, η απουσία αναφοράς της κατανομής των δεδομένων εκπαίδευσης και δοκιμής είναι κάτι αρνητικό για την έρευνα. Οι συγγραφείς εστιάζουν στην αρχιτεκτονική και τη διαδικασία του συστήματος και αγνοούν τα ποσοστά κατανομής.

5.2.2.4 Resume Parsing using Natural Language Processing (Αναλυτής Βιογραφικών με τη Χρήση της Επεξεργασίας Φυσικής Γλώσσας)

Η ανάλυση του βιογραφικού σημειώματος με τη συμβολή της Επεξεργασίας της Φυσικής Γλώσσας (NLP) είναι το αντικείμενο με το οποίο ασχολούνται οι Chavare και Patil (2023). Προτείνουν ένα μοντέλο Μηχανικής Μάθησης που επιλέγει αυτόματα τα κατάλληλα βιογραφικά βάσει της περιγραφής θέσης εργασίας. Έτσι, αποφεύγεται η απόρριψη, κατά λάθος, των κατάλληλων υποψηφίων και η πρόσληψη ακατάλληλων. Το μοντέλο δέχεται βιογραφικά, δηλαδή κείμενο μη δομημένο (unstructured) από τους υποψηφίους και παράλληλα περιγραφές θέσεων εργασίας, από τους εργοδότες.

Με την χρήση της NLP, το σύστημα εντοπίζει όχι μόνο τις λέξεις, τις οντότητες κ.α., αλλά ακόμα και την αργκό, την ζωντανή γλώσσα, δηλαδή την καθομιλουμένη, και τα ακρωνύμια. Με το μοντέλο νευρωνικών δικτύων BERT, κατανοείται η φυσική γλώσσα. Με το μοντέλο Spacy NER εξάγονται οι απαιτούμενες οντότητες από τα μη δομημένα δεδομένα που προκύπτουν από τα κείμενα ενώ οι αλγόριθμοι ταξινόμησης Joint NER και Relation Extraction (Εξαγωγή Σχέσεων), ταξινομούν τις χρήσιμες πληροφορίες από τα κείμενα. Το Joint NER αναγνωρίζει και ταξινομεί οντότητες (ονόματα, τίτλους θέσεων, εμπειρία, δεξιότητες, δίπλωμα, ειδικότητα κ.λπ.) ενώ το Relation Extraction αναγνωρίζει τις συσχετίσεις μεταξύ των οντοτήτων αυτών. Το τελευταίο είναι εξαιρετικά χρήσιμο καθώς συσχετίζει τις πληροφορίες των βιογραφικών με τις περιγραφές θέσεων εργασίας. Επίσης, χρησιμοποιήθηκε το μοντέλο αλγορίθμου ταξινόμησης “Connection Extraction” (Εξαγωγή Συνδέσεων) που προβλέπει την εκάστοτε σχέση μεταξύ δύο οντοτήτων. Ειδικότερα, το Connection Extraction εστιάζει στην εύρεση συσχετίσεων, μοτίβων και διασυνδέσεων μεταξύ δεδομένων οποιασδήποτε μορφής. Το Relation Extraction αποτελεί μέρος του Connection Extraction καθώς είναι περισσότερο εξειδικευμένο, αφού αναγνωρίζει συσχετίσεις οντοτήτων μόνο σε κείμενα φυσικής γλώσσας, ενώ το Connection Extraction αναγνωρίζει όλα τα δεδομένα εισόδου, όπως πίνακες και γραφήματα. Τέλος, δημιουργείται ένα γράφημα γνώσης (Knowledge Graph) με τιμές: οντοτήτων δεξιοτήτων και χρόνο εμπειρίας. Το γράφημα δημιουργήθηκε από το pyvis⁵⁵ και το

⁵⁵ Το Pyvis είναι ένα διαδραστικό εργαλείο για οπτικοποιήσεις γραφημάτων. [Interactive network visualizations — pyvis 0.1.3.1 documentation](#)

networkx⁵⁶, εργαλεία της python. Αρχικά επιλέγεται συγκεκριμένη περιγραφή εργασίας, έπειτα εξάγονται οι κατάλληλες δεξιότητες και δημιουργείται το γράφημα.

Για όλα τα παραπάνω, χρησιμοποιήθηκε η UBAI⁵⁷ (Universal Business Intelligence and Artificial Intelligence), μία πλατφόρμα για τη διαχείριση και κατηγοριοποίηση των κειμένων που προσφέρει εργαλεία, όπως το NER και το Relation Extraction. Η ισχυρή μονάδα επεξεργασίας γραφικών (Graphics Processing Unit - GPU) ήταν απαραίτητη καθώς απαιτείται υψηλή υπολογιστική ισχύς για τις πολλαπλές διεργασίες που εκτελούνται ταυτόχρονα.

Καινοτομία στο μοντέλο αυτό αποτελεί η ταυτόχρονη χρήση του Joint NER και του Relation/Connection Extraction. Ο συνδυασμός αυτός ενισχύει ακόμα περισσότερο την κατανόηση αδόμητου κειμένου. Δεύτερη καινοτομία είναι η δημιουργία του γραφήματος γνώσης. Οι ερευνητές επικεντρώνονται στην διαδικασία, τα εργαλεία, τις εφαρμογές και την αρχιτεκτονική του συστήματος που προτείνουν και δεν αναφέρουν ποιά δεδομένα χρησιμοποιήθηκαν. Κάνουν αναφορά ότι μπορούν να χρησιμοποιηθούν δεδομένα από μη δομημένα κείμενα, από συμβόλαια, από οικονομικής φύσεως έγγραφα, από αρχεία υγειονομικής περίθαλψης κ.α. Δεν πραγματοποιήσαν αξιολόγηση του συστήματος.

5.2.2.5 Resume Recommendation Using Machine Learning (Σύσταση Βιογραφικών με τη Χρήση της Μηχανικής Μάθησης)

Ο σκοπός των Naik και Dhotre (2022) είναι να δημιουργήσουν ένα μοντέλο Μηχανικής Μάθησης που να βρίσκει αυτόματα τα δέκα έως είκοσι καλύτερα βιογραφικά σημειώματα από έναν μεγάλο όγκο βιογραφικών βασιζόμενο σε προσωπικές πληροφορίες, στα προσόντα, στις δεξιότητες και την εμπειρία των υποψηφίων. Η εξέλιξη της τεχνολογίας έχει οδηγήσει στις ολοένα και περισσότερες online αιτήσεις, είτε σε πλατφόρμες ευρέσεως εργασίας και είτε βεβαίως μέσω ηλεκτρονικού ταχυδρομείου. Όλες αυτές οι αιτήσεις, οι οποίες σπανίως έχουν ίδια μορφή και δομή, ενδέχεται να υποβάλλονται σε μορφή PDF ενώ πρέπει να

⁵⁶ Το Networkx είναι ένα εργαλείο για τη δημιουργία, τον χειρισμό και τη μελέτη της δομής, της δυναμικής και των λειτουργιών πολύπλοκων δικτύων. Στο άρθρο, χρησιμοποιήθηκε για την ανάλυση γραφημάτων. [NetworkX — NetworkX documentation](#)

⁵⁷ <https://ubiai.tools/>

ταξινομηθούν γρήγορα, χωρίς λάθη και να αντιστοιχιστούν στην κατάλληλη θέση στην παραγόμενη κατάταξη.

Αρχικά πραγματοποιείται καθαρισμός λέξεων (αφαίρεση αριθμών, ειδικών χαρακτήρων, μονόφθογγες λέξεις), tokenization (κατηγοριοποίηση λεκτικών μονάδων), εξαγωγή και χαρτογράφηση των χαρακτηριστικών. Ακολουθεί η εκπαίδευση αλγορίθμων Μηχανικής Μάθησης και οι προτάσεις των βιογραφικών που ταιριάζουν στη συγκεκριμένη θέση εργασίας. Το Python Jupyter Notebook⁵⁸, το οποίο είναι μια διαδραστική διαδικτυακή εφαρμογή ανοικτού κώδικα για τη δημιουργία και την κοινή χρήση υπολογιστικών εγγράφων, χρησιμοποιήθηκε για αυτό το μοντέλο.

Μετά τον “καθαρισμό” των κειμένων και το tokenization, εφαρμόζεται ανάλυση και κατηγοριοποίηση των βιογραφικών από το Natural Language Toolkit (NLTK). Ο Naive Bayes είναι αλγόριθμος που βασίζεται στην Θεωρία Πιθανοτήτων του Bayes. Επειδή υποθέτει ότι όλα τα χαρακτηριστικά των δεδομένων είναι ανεξάρτητα μεταξύ τους, ονομάζεται “naive”. Είναι κατάλληλος για την ανάλυση συναισθήματος (sentiment analysis) και για εφαρμογές Επεξεργασίας Φυσικής Γλώσσας (NLP) αν και χρησιμοποιήθηκε κυρίως για εξαγωγή χαρακτηριστικών. Οι αλγόριθμοι που εκπαιδεύτηκαν είναι οι SVM και K-Nearest Neighbor (KNN). Το Cosine Similarity, επίσης κατάλληλο για την NLP, χρησιμοποιήθηκε στους δύο αλγορίθμους για τον υπολογισμό των ομοιοτήτων μεταξύ κειμένων, εγγράφων και προτάσεων. Το K-Nearest Neighbor, αλγόριθμος ταξινόμησης και παλινδρόμησης, βρίσκει τους “K” πλησιέστερους γείτονες των δεδομένων. Οι δύο τελευταίοι αλγόριθμοι εντοπίζουν τις ομοιότητες μεταξύ των προσόντων και των θέσεων εργασίας. Ο K-Nearest Neighbor λειτούργησε καλύτερα αφού είχε πιστότητα (accuracy) 99%, δηλαδή οι αντιστοιχίσεις των βιογραφικών και των θέσεων εργασίας ήταν σχεδόν τέλειες.

Δεν μπορεί να βρεθεί κάτι αρνητικό στο μοντέλο που παρουσίασαν οι (Naik & Dhotre, 2022) καθώς είναι ξεκάθαρο, γρήγορο και εύχρηστο. Η πρόταση των δέκα με είκοσι καλύτερων υποψηφίων διενεργείται αυτόματα και η ορθότητά του αγγίζει το 100%. Δεν αναφέρεται στο άρθρο εάν η πιστότητα αφορά μόνο την ορθή αντιστοίχιση των βιογραφικών με τις κατάλληλες θέσεις ή αν αφορά και τη σειρά κατάταξης των βιογραφικών. Φαίνεται ότι οι συγγραφείς επικεντρώνονται περισσότερο στην σωστή αντιστοίχιση, παρά στην κατάταξη των βιογραφικών.

⁵⁸ [Jupyter Notebook](#)

5.2.2.6 Resume Screening using Machine Learning and NLP: A Proposed System (Διαλογή Βιογραφικών με τη Χρήση της Μηχανικής Μάθησης και της Επεξεργασίας Φυσικής Γλώσσας: Ένα Προτεινόμενο Σύστημα)

Οι Kinge και συν. (2022) προτείνουν ένα σύστημα που συνδυάζει την Μηχανική Μάθηση και την Επεξεργασία της Φυσικής Γλώσσας (NLP) για τον έλεγχο των βιογραφικών σημειωμάτων ώστε να βοηθήσουν τις επιχειρήσεις και τις κυβερνήσεις, που διαχειρίζονται μεγάλο όγκο βιογραφικών, να επιλέξουν τον κατάλληλο υπάλληλο. Επίσης, δίνονται συστάσεις στους υποψηφίους για να βελτιώσουν το βιογραφικό τους. Το μοντέλο συστήματος που παρουσιάζουν έχει ως εισερχόμενο το βιογραφικό σε κείμενο, που έπειτα αποθηκεύεται σε μία βάση δεδομένων SQL. Τα δεδομένα (βιογραφικά) με τα οποία εκπαιδεύτηκαν οι αλγόριθμοι, συλλέχθηκαν από την ανοιχτή πλατφόρμα Kaggle, χωρίς να παρέχεται συγκεκριμένος αριθμός. Η εξαγωγή στοχευμένων στοιχείων από το GitHub και το LinkedIn αποτελεί καινοτόμα πρακτική.

Το GitHub είναι μία cloud-based υπηρεσία, που χρησιμοποιείται κυρίως για την διαχείριση και αποθήκευση του πηγαίου κώδικα. Όμως έχει και δεύτερη υπόσταση. Αποτελεί ένα παγκόσμιο κοινωνικό δίκτυο προγραμματιστών, μέσω του οποίου οι προγραμματιστές μπορούν να προωθήσουν την δουλειά τους και να τους δημιουργηθούν επαγγελματικές ευκαιρίες. Μέσω του GitHub, κάποιος μπορεί παρουσιάσει το έργο του ή να παρακολουθήσει τις αλλαγές που πραγματοποιούνται σε κάποιον κώδικα, όπως και να συνεργαστεί με άλλους προγραμματιστές σε έργα λογισμικού ανοιχτού κώδικα. Οπότε, μέσα από αυτά τα έργα, παρουσιάζονται οι τεχνικές δεξιότητες των προγραμματιστών. Το GitHub αφορά κυρίως άτομα που ασχολούνται με το τεχνικό πεδίο της πληροφορικής (Simply Great Resumes, 2023).

Το LinkedIn είναι το μεγαλύτερο παγκόσμιο επαγγελματικό μέσο κοινωνικής δικτύωσης καθώς παρέχει τη δυνατότητα εκτός από το να μεταφορτώσεις το βιογραφικό σου, επιπλέον να συνδεθείς με άλλους επαγγελματίες, επιχειρήσεις και αρμόδιους προσλήψεων. Οι χρήστες του ξεπερνούν πλέον τα 500 εκατομμύρια και πολλοί υπεύθυνοι προσλήψεων προσεγγίζουν υποψηφίους αποκλειστικά μέσω της εφαρμογής αυτής (professionaid, 2024). Το LinkedIn παρέχει πληροφορίες βιογραφικού χαρακτήρα και τις επαγγελματικές διασυνδέσεις ενός ατόμου. Σε αντίθεση με το LinkedIn, το GitHub δεν παρέχει πληροφορίες σχετικά με π.χ. την

εκπαίδευση και την προϋπηρεσία του ατόμου. Το GitHub επικεντρώνεται στο τεχνικό επίπεδο και τις προγραμματιστικές γνώσεις του χρήστη. Ο συνδυασμός λοιπόν των δύο αυτών εφαρμογών, ειδικά για τους υποψηφίους στον κλάδο της τεχνολογίας, προσδίδει τεράστια ώθηση για την επιτυχή έκβαση ευρέσεως εργασίας αφού παρουσιάζονται τόσο τα βιογραφικά στοιχεία, όσο και οι πρακτικές δεξιότητες των υποψηφίων.

Η διαδικασία που επιτελείται στο προτεινόμενο σύστημα έχει ως εξής: ο υποψήφιος υποβάλλει το βιογραφικό του στο front-end, μετά αυτό καταλήγει στον αναλυτή βιογραφικών (resume parser), ο οποίος με τη χρήση της Επεξεργασία της φυσικής γλώσσας (NLP) και του NER, το οποίο χρησιμοποιείται στην αναγνώριση οντοτήτων, εξάγει τις χρήσιμες πληροφορίες. Παράλληλα, το σύστημα επισκέπτεται το προφίλ του υποψηφίου στις πλατφόρμες GitHub και LinkedIn για να συγκεντρώσει περισσότερες πληροφορίες. Ακολουθώντας, η χρήση του Cosine Similarity χρησιμοποιείται για τον εντοπισμό των ομοιοτήτων των κειμένων, ώστε η αντιστοίχιση να παράγεται αυτόματα από το εκπαιδευμένο μοντέλο Μηχανικής Μάθησης. Οι αλγόριθμοι που χρησιμοποιήθηκαν είναι ο K-Nearest Neighbor (KNN) και ο SVM, οι οποίοι προβλέπουν ποιό βιογραφικό ταιριάζει καλύτερα με την ελεύθερη θέση εργασίας. Στο τέλος της διαδικασίας οι υποψήφιοι ανατροφοδοτούνται για την καλύτερη αντιστοίχιση του βιογραφικού τους με συγκεκριμένη θέση εργασίας.

Το μειονέκτημα αυτής της έρευνας έγκειται στο ότι δεν εξηγείται σαφώς ποιά είναι η πιστότητα (accuracy) για το κάθε μοντέλο MM που έχει παραχθεί. Ενώ περιγράφεται η διαδικασία, στο τέλος δίνονται ποσοστά από 78% μέχρι 98%, όμως αόριστα. Οι συγγραφείς γράφουν ότι η πιστότητα ποικίλει ανάλογα με τα σύνολα δεδομένων που χρησιμοποιήθηκαν, την πολυπλοκότητα των μεθόδων εκμάθησης και το μέγεθος του συνόλου δεδομένων. Η τεκμηρίωση των αποτελεσμάτων είναι ανεπαρκής. Θετική είναι η προσθήκη των δύο κοινωνικών δικτύων/πλατφορμών, μέσω των οποίων συλλέγονται επιπλέον κριτήρια για την πρόβλεψη πρόσληψης υποψηφίων. Τα επιπλέον κριτήρια που συλλέγονται, όπως αναφέρουν οι συγγραφείς, είναι η συνεισφορά των υποψηφίων σε διάφορους τομείς χωρίς να αναφέρονται σε ποιους συγκεκριμένα. Θεωρείται ότι από το GitHub μπορεί να συλλέγουν τις τεχνικές δεξιότητες και από το LinkedIn την επαγγελματική σταδιοδρομία. Οπότε, φαίνεται ότι θέλουν να εστιάσουν περισσότερο σε κλάδους τεχνολογίας όπου οι τεχνικές

δεξιότητες είναι απαραίτητες και αξιοποιούνται τα δεδομένα του GitHub. Ο κύριος λόγος που μελετήθηκε αυτή η έρευνα είναι η καινοτόμα προσέγγιση συνδυασμού αυτών των δύο κοινωνικών δικτύων για επιπλέον εξαγόμενα χαρακτηριστικά. Η ανατροφοδότηση προς τους υποψηφίους για την καλύτερη αντιστοίχισή τους με θέσεις εργασίας είναι κάτι πολύ χρήσιμο για αυτούς. Τέλος, θα πρέπει να υπάρχει συγκατάθεση των υποψηφίων για τη συλλογή των δεδομένων τους από το σύστημα, διαφορετικά μπορεί να προκύψουν νομικά ζητήματα.

5.2.2.7 Design and Development of Machine Learning Based Resume Ranking System (Σχεδιασμός και Ανάπτυξη Συστήματος Κατάταξης Βιογραφικών Βάσει της Μηχανικής Μάθησης)

Στο πείραμά τους οι Tejaswini και συν. (2022) ασχολούνται με την διαλογή, το ταίριασμα και την κατάταξη των καταλληλότερων βιογραφικών σε σχέση με συγκεκριμένη θέση εργασίας που πραγματοποιείται μέσω ενός λογισμικού ελέγχου βιογραφικών που παίρνει τη μορφή ενός συστήματος συστάσεων. Ένα σύστημα συστάσεων (recommendation engine) φιλτράρει πληροφορίες και προσπαθεί να προβλέψει τις προτιμήσεις ή την βαθμολόγηση για οτιδήποτε. Τέτοια συστήματα χρησιμοποιούν π.χ. το Netflix, το Tik Tok, e-shop κ.α. για συστάσεις ταινιών, βίντεο και προϊόντων. Υπάρχουν δύο είδη συστημάτων συστάσεων, αυτά που βασίζονται στο περιεχόμενο των αντικειμένων και τις προτιμήσεις του ίδιου καταναλωτή (Content-based recommender system) και αυτά που βασίζονται στο συνεργατικό φιλτράρισμα, δηλαδή σε παρόμοιες προτιμήσεις διαφορετικών ανθρώπων (Collaborative filtering recommender system). Εδώ χρησιμοποιείται το πρώτο είδος συστήματος συστάσεων, βασιζόμενο στο περιεχόμενο. Τέλος, οι συγγραφείς διάλεξαν πενήντα βιογραφικά από τη βάση δεδομένων Kaggle σε μορφή DOC και DOCX και τα μετέτρεψαν σε μορφή PDF για να διαχειριστούν ομοιόμορφα δεδομένα.

Ο υποψήφιος εισέρχεται στο προτεινόμενο σύστημα των Tejaswini και συν., (2022), με τα διαπιστευτήρια του λογαριασμού του στην Google. Στη συνέχεια, επιλέγει απαντήσεις από μία φόρμα ερωτήσεων πολλαπλής επιλογής που εξετάζει τις βασικές γνώσεις του και είτε θα αποτυγχάνει να προχωρήσει στην εφαρμογή, είτε θα εισέρχεται στην ιστοσελίδα προκειμένου να υποβάλλει το βιογραφικό του που θα αποθηκεύεται στη βάση δεδομένων. Όταν το βιογραφικό γίνει αποδεκτό και έχει γίνει

η προεπεξεργασία των λέξεων ώστε αυτές να λάβουν τη ριζική τους μορφή, εφαρμόζεται το TF-IDF (Term Frequency-Inverse Document Frequency). Με αυτό τον τρόπο, δημιουργούνται τα χαρακτηριστικά που χρειάζονται οι αλγόριθμοι.

Οι ερευνητές, χρησιμοποίησαν μία βιβλιοθήκη ανοικτού κώδικα, το Gensim (από το “generate similar”) που είναι γραμμένη στη γλώσσα προγραμματισμού Python, η οποία επεξεργάζεται τη φυσική γλώσσα (NLP) και συνοψίζει τα κείμενα. Έπειτα, με το Cosine Similarity εντοπίζονται οι ομοιότητες μεταξύ των βιογραφικών και της θέσης εργασίας, αντιστοιχώντας στην αναμενόμενη λειτουργία ενός συστήματος συστάσεων με βάση το περιεχόμενο. Ακολούθως, εκπαιδεύεται ο αλγόριθμος KNN για να εντοπίσει τα δέκα βιογραφικά που ταιριάζουν καλύτερα με την τρέχουσα θέση εργασίας και τα ταξινομεί, βάσει των αποτελεσμάτων του Cosine Similarity. Το σύστημα έχει πολύ καλή πιστότητα ανάλυσης κειμένου (parsing accuracy), στα δέκα βιογραφικά ανέλυσε και τα δέκα ορθά (100%), στα δεκαπέντε τα δεκατέσσερα (93.3%) και στα είκοσι τα δεκαεπτά (85%). Αυτό αναφέρεται στο κατά πόσο προέβλεψε ορθά την συσχέτιση των βιογραφικών με τη θέση εργασίας. Η πιστότητα κατάταξης (ranking accuracy) ήταν: από τα δέκα βιογραφικά, βρήκε ορθά τη σωστή σειρά και των δέκα (100%), από τα δεκαπέντε τα δώδεκα (80%) και από τα είκοσι τα δεκαπέντε (75%). Συνεπώς, τα τελευταία αποτελέσματα αναφέρονται στο κατά πόσο είναι σωστή η σειρά κατάταξης των βιογραφικών. Παρατηρείται ότι όσο αυξάνεται ο αριθμός των βιογραφικών, τα ποσοστά πιστότητας (τόσο ανάλυσης κειμένου όσο και κατάταξης) μειώνονται.

Το καινοτόμο στοιχείο σε αυτό το μοντέλο είναι ότι κάνει μία αρχική διαλογή των βιογραφικών. Η απόρριψη μετά από ανεπιτυχείς απαντήσεις στο ερωτηματολόγιο πολλαπλής επιλογής μειώνει τον όγκο των βιογραφικών που δέχεται και πρέπει να επεξεργαστεί το σύστημα. Οι συγγραφείς εξηγούν ότι οι ερωτήσεις βοηθούν στο να εντοπίσουν υποψήφιους με ανεπαρκείς δεξιότητες. Όμως, δεν διευκρινίζεται εάν αυτό αφορά την συγκεκριμένη θέση. Ο υποψήφιος πρέπει να επιτύχει μία ελάχιστη βαθμολογία για να δεχτούν το βιογραφικό του. Κατά την επεξεργασία με την βιβλιοθήκη Gensim και τη χρήση του TF-IDF, υπάρχει φόβος μήπως χάνεται σημαντική πληροφορία που τελικά εν δυνάμει μπορεί να αλλοιώσει το αποτέλεσμα, για αυτό και πρέπει οι ερευνητές να πειραματιστούν με διαφορετικές μεθόδους και εργαλεία, όπως έγινε στις προηγούμενες προσεγγίσεις. Ο αριθμός των βιογραφικών που εξετάστηκαν παραμένει και σε αυτό το άρθρο μη ικανοποιητικός.

5.2.2.8 Recruitment Prediction using Machine Learning (Πρόβλεψη Πρόσληψης με τη Χρήση της Μηχανικής Μάθησης)

Υπάρχουν περιπτώσεις στις οποίες μόλις επιλεγθεί ένας υποψήφιος, αποφασίζει για δικούς του λόγους να μην προχωρήσει στην αποδοχή της θέσης/πρόσληψης. Αυτό έχει κάποιο κόστος χρόνου και χρήματος για τις εταιρείες και τις κυβερνήσεις, καθώς δεν είναι πάντα εύκολο να αναπληρωθούν άμεσα οι θέσεις αυτές και να επαναληφθούν οι διαδικασίες πρόσληψης. Ο σκοπός της έρευνας των Jagan και συν. (2020) είναι η πρόβλεψη σίγουρης ανάληψης υπηρεσίας υποψηφίων βάσει ποιοτικών και ποσοτικών γνωρισμάτων, όπως η ηλικία, το φύλο, η προϋπηρεσία, το μηνιαίο εισόδημα, η εμπειρία, η απόσταση από τον χώρο εργασίας, η οικογενειακή κατάσταση και η αύξηση του μισθού. Δημιουργήθηκε ένα σύνολο εικονικών δεδομένων για τη δημιουργία βιογραφικών, βασισμένο σε περιγραφές θέσεων εργασίας. Το δείγμα δεδομένων (dataset) είναι διαθέσιμο στο GitHub.

Τα βιογραφικά εισέρχονται στο προτεινόμενο σύστημα και μετά από προεπεξεργασία (π.χ. αφαίρεση μη επιθυμητών στοιχείων και μηδενικών τιμών - null values), αποκτούν συγκεκριμένη μορφή και δομή καθώς και ταξινομούνται. Χρησιμοποιήθηκε ο Συντελεστής Συσχέτισης του Pearson, ο οποίος επιλέγει τα κατάλληλα χαρακτηριστικά και μετρά τις συσχετίσεις ανάμεσα σε μεταβλητές όπως το κατά πόσο σχετίζεται η απόσταση της εργασίας από την κατοικία του υποψηφίου. Το υποσύνολο εκπαίδευσης ήταν 70% και δοκιμής 30%. Με τη μετρική της πιστότητας (accuracy) μετρήθηκε το ποσοστό της σωστής πρόβλεψης για πρόσληψη ή όχι. Με τη μετρική της ακρίβειας (precision) το σύστημα προέβλεψε το ποσοστό των υποψηφίων που όντως θα αναλάβουν καθήκοντα, σε σχέση με τον συνολικό αριθμό των υποψηφίων, ενώ με την μετρική της ανάκλησης (recall) προβλέφθηκε το ποσοστό αυτών που θα αναλάβουν τη δουλειά και θα παραμείνουν σε αυτή, σε σχέση με αυτούς που θα έπρεπε να προταθούν.

Οι αλγόριθμοι που χρησιμοποιήθηκαν για εκπαίδευση είναι οι: Decision Tree, Random Forest, Gaussian Naive Bayes (GNB) και ο K-Nearest Neighbor (KNN). Για τις παραπάνω μετρικές τα αποτελέσματα κατά αντίστοιχη σειρά είναι: πιστότητα 96%, 95%, 99% και 93%, ακρίβεια 75%, 100%, 100% και 44%, και τέλος για την ανάκληση 70%, 17%, 88% και 23%. Συνεπώς, ο GNB έφερε τα καλύτερα ποσοστά.

Η πρωτότυπη προσέγγιση αυτής της έρευνας έγκειται στο ότι η χρήση της Μηχανικής Μάθησης προσφέρει όχι μόνο την πρόβλεψη της πρόσληψης του υποψηφίου, αλλά και την πιθανότητα παραμονής του στη θέση εργασίας που έχει προσληφθεί. Αυτό είναι ένα μείζον ζήτημα που αντιμετωπίζουν οι οργανισμοί καθώς χάνουν χρόνο και πόρους με την επανάληψη της διαδικασίας πρόσληψης. Η δημιουργία του συνόλου δεδομένων που είναι διαθέσιμο στο GitHub, καθιστά δυνατή την αναπαραγωγή της έρευνάς τους από ειδικούς στον τομέα αυτό. Η προεπεξεργασία και η εξαγωγή στοιχείων δεν εξηγείται διεξοδικά, οπότε θα μπορούσε να ειπωθεί ότι υπάρχει κενό στην μεθοδολογία που ακολούθησαν. Τα χαμηλά ποσοστά, ειδικά στην μετρική της ανάκλησης δημιουργεί ερωτηματικά για το κατά πόσο είναι αποτελεσματικό το προτεινόμενο σύστημα. Η έρευνα παρουσιάζει τα αποτελέσματα των μετρικών, χωρίς όμως να έχουν εφαρμοστεί στην πράξη.

5.2.2.9 Differential Hiring using a Combination of NER and Word Embedding (Διαφορική Πρόσληψη με τη Χρήση Συνδυασμού Αναγνώρισης Οντοτήτων και Ενσωμάτωσης Λέξεων)

Στην έρευνά τους, οι Suhas και Manjunath (2020) προτείνουν μία νέα μέθοδο για την αντιστοίχιση του κατάλληλου υποψήφιου με μία θέση εργασίας. Αυτή η μέθοδος προτείνει στους εργοδότες, τους υποψηφίους με τα καλύτερα βιογραφικά, και στους υποψηφίους τις θέσεις εργασίας που ταιριάζουν καλύτερα με τα βιογραφικά τους. Τα βιογραφικά είναι υπαρκτά και με την έγκριση των υποψηφίων, χρησιμοποιήθηκαν για την έρευνα. Όμως, οι συγγραφείς δεν αναφέρουν τον αριθμό και την πηγή τους, όπως ούτε και τον αριθμό των αξιολογητών που εξέτασαν τα βιογραφικά.

Μετά την συγκέντρωση βιογραφικών των υποψηφίων και την εισαγωγή τους στο σύστημα, αρχίζει η Επεξεργασία της Φυσικής Γλώσσας (NLP) με την εξαγωγή των απαραίτητων στοιχείων. Το βήμα της NLP αποτελείται: από το Stop Word Removal, που αφαιρεί τις ανούσιες λέξεις, το stemming που μειώνει τις λέξεις μέχρι τη ρίζα τους και αφαιρεί καταλήξεις και παραλλαγές. Τέλος, από το NER που αναγνωρίζει και κατηγοριοποιεί τις οντότητες των κειμένων, ώστε να παραμείνουν οι κρίσιμες πληροφορίες για την εκπαίδευση των αλγορίθμων. Ειδικότερα, προς αυτό τον σκοπό, επισημαίνονται οι δεξιότητες, τα εκπαιδευτικά ιδρύματα που σπουδών και άλλα χαρακτηριστικά, για να συσχετιστούν στη συνέχεια με τις περιγραφές θέσεων

εργασίας με τη βοήθεια του Cosine Similarity και να γίνει η κατάταξη των υποψηφίων.

Επίσης, χρησιμοποιήθηκε και το νευρωνικό δίκτυο Word embedding model (word2vec) που εντόπισε τις λέξεις που έχουν παρόμοιο νόημα, από τις προηγούμενες εξαγόμενες λέξεις. Το σύνολο δεδομένων των δεξιοτήτων που χρησιμοποίησαν οι Suhas & Manjunath (2020) για την εκπαίδευση του word2vec ανακτήθηκε από την ιστοσελίδα “stack exchange”, στην οποία οι χρήστες κάνουν ερωτήσεις για την πληροφορική. Το τελικό βήμα ήταν η εφαρμογή της μετρικής Cosine Similarity που χρησιμοποιήθηκε για τον εντοπισμό των ομοιοτήτων μεταξύ των βιογραφικών και των θέσεων εργασίας ώστε να μπορεί το σύστημα έπειτα να προτείνει τους καλύτερους πέντε με δέκα υποψηφίους. Δεν διασαφηνίζεται στο άρθρο συγκεκριμένος αριθμός από προτεινόμενους υποψηφίους. Μία υπόθεση που μπορεί να γίνει είναι ότι οι υποψήφιοι είναι κατ’ ελάχιστο πέντε και κατά το μέγιστο δέκα. Ο αριθμός πέντε ενδέχεται να αναφέρεται στην περίπτωση που δεν υπάρχουν περισσότερα κατάλληλα βιογραφικά.

Οι συγγραφείς, συνέκριναν τα αποτελέσματα (outputs) του συστήματος με τις ανθρώπινες αποδόσεις. Οι μετρικές για το παραπάνω είναι οι: recall, precision και F-score. Όμως, ενώ οι συγγραφείς είχαν γραφήματα στη δημοσίευσή τους, δεν αποτύπωσαν τα ακριβή ποσοστά αυτών καθώς η κλίμακα των γραφημάτων τους ήταν πολύ μεγάλη και αυτό δεν επιτρέπει στον αναγνώστη να διαβάσει τα ακριβή ποσοστά, αλλά μόνο να τα προσεγγίσει. Αυτό που προκύπτει από τα γραφήματα, είναι ότι το recall και το F-score προσεγγίζουν το 80%. Το recall υποδεικνύει το ποσοστό των ορθά αναγνωρισμένων οντοτήτων από το NER σε σχέση με την χειρωνακτική αντιστοίχιση που όντως αναγνωρίστηκαν. Η καλύτερη συσχέτιση των αποτελεσμάτων του NER με τα ανθρώπινα αποτελέσματα, (δηλαδή ανά περίπτωση κατά πόσο ταίριαζαν τα αποτελέσματα του NER με τα ανθρώπινα αποτελέσματα), ως προς τη μετρική F-score είναι οι δεξιότητες με 84% και ο εργοδότης του υποψηφίου με 81%. Αυτά που συσχετίστηκαν λιγότερο είναι η προϋπηρεσία με 75% και το εκπαιδευτικό ίδρυμα με 70%. Το precision, με ποσοστό στο 82%, μέτρησε τις οντότητες που αναγνωρίστηκαν ορθά.

Ο νεωτερισμός σε αυτή την έρευνα αφορά τη χρήση του νευρωνικού δικτύου Word embedding (Word2Vec), το οποίο οι ερευνητές πιστεύουν ότι κάνει πιο περιεκτική τη

διαδικασία εξαγωγής των κατάλληλων λέξεων. Η χρήση του NER φέρει υψηλά αποτελέσματα ιδίως στην κατηγορία των δεξιοτήτων. Ιδιαίτερη έμφαση δίνουν οι συγγραφείς στη χρήση του Cosine Similarity καθώς εντοπίζει τις ομοιότητες μεταξύ των βιογραφικών και της περιγραφής θέσης εργασίας. Για ακόμα μία φορά οι ερευνητές επικεντρώνονται στην καινοτομία που προτείνουν, εδώ στον συνδυασμό του NER με το Word embedding. Επίσης, δίνουν τιμές με έμμεσο τρόπο, τις οποίες είναι δύσκολο να τις προσδιορίσει κανείς διότι είναι μέρος ενός γραφήματος με ακατάλληλη κλίμακα. Οι συγγραφείς δεν αναζητούν την ακρίβεια στην αντιστοίχιση βιογραφικών, αλλά μόνο την ακρίβεια στον εντοπισμό/αναγνώριση οντοτήτων. Επιπλέον, η συλλογή δεδομένων από το Stack Exchange περιορίζεται στον τομέα της πληροφορικής, οπότε τα δεδομένα δεν αντιπροσωπεύουν όλους τους κλάδους εργασίας.

5.2.2.10 A Machine Learning Approach for Automation of Resume Recommendation System (Μια Προσέγγιση Μηχανικής Μάθησης για την Αυτοματοποίηση του Συστήματος Σύστασης Βιογραφικών)

Οι Roy και συν. (2020) προτείνουν ένα σύστημα για την αυτοματοποιημένη σύσταση βιογραφικών σημειωμάτων μέσα από την ταξινόμηση και την αντιστοίχιση τους σε προσφερόμενες θέσεις εργασίας. Τα δεδομένα των βιογραφικών συγκεντρώθηκαν από το Kaggle σε μορφή Excel, με τρεις στήλες (ID, Category, Resume) δηλαδή τον αριθμό του βιογραφικού, τον κλάδο εργασίας και το βιογραφικό του υποψηφίου (ως περιεχόμενο). Αρχικά, με το Natural Language Toolkit (NLTK), υλοποιούνται, το Stop Word Removal, το stemming και το lemmatization, διαμορφώνονται, μειώνονται και αφαιρούνται οι μη απαραίτητες λέξεις και το κείμενο χωρίζεται σε οντότητες (tokens). Μετά από την προεπεξεργασία των λέξεων χρησιμοποιούν το Data Masking, μία τεχνική απόκρυψης δεδομένων που μεταμορφώνει τα δεδομένα σε δυσανάγνωστα ή μη αναγνωρίσιμα. Αυτό προστατεύει τα ευαίσθητα δεδομένα από μη εξουσιοδοτημένους χρήστες, από τυχόν διαρροές ή κακόβουλη χρήση.

Τα tokens έπειτα αντιστοιχίζονται σε τιμές μέσω της εφαρμογής του TF-IDF (Term Frequency-Inverse Document Frequency) στο πλαίσιο της προσέγγισης BoW (Bag of Words), με τη βοήθεια της βιβλιοθήκης scikit learn. Πριν την εφαρμογή του KNN χρησιμοποίησαν την βιβλιοθήκη Gensim. Η βιβλιοθήκη Gensim δημιουργεί

περιλήψεις των κειμένων των βιογραφικών και των περιγραφών θέσεων εργασίας. Η επόμενη κίνηση είναι η σύγκριση των βιογραφικών με τις περιγραφές θέσεων εργασίας και η εύρεση των πιο ταιριαστών αντιστοιχίσεων, με τον αλγόριθμο KNN και με το Cosine Similarity, εναλλακτικά. Έπειτα, εφαρμόζονται αλγόριθμοι MM που ταξινομούν τα βιογραφικά στην κατάλληλη κατηγορία εργασίας και προτείνουν τα δέκα επικρατέστερα βιογραφικά. Στους παρακάτω αλγόριθμους υπολογίστηκαν τα δέκα πλησιέστερα βιογραφικά με βάση τον αλγόριθμο KNN σε συνδυασμό με την Gensim και με το Cosine Similarity σε συνδυασμό με το BoW. Οι αλγόριθμοι που δοκιμάστηκαν είναι οι: Random Forest, Multinomial Naive Bayes, Logistic Regression και Linear Support Vector Classifier. Καλύτερα ποσοστά πιστότητας (accuracy) έφερε ο Linear Support Vector Classifier με τον KNN με 78.53%. Αυτό το ποσοστό δείχνει το πόσο καλά ταξινόμησε τα βιογραφικά στην ορθή κατηγορία κλάδου εργασίας.

Η τεχνική απόκρυψης Data Masking προστατεύει τα ευαίσθητα δεδομένα και διασφαλίζει τη συμμόρφωση με Κανονισμούς, όπως το GDPR. Αυτό μειώνει τον κίνδυνο νομικών συνεπειών λόγω διαρροών ευαίσθητων δεδομένων. Όμως, αν και το προτεινόμενο σύστημα προστατεύει τα ευαίσθητα δεδομένα, η τροποποίηση των δεδομένων ενδέχεται να δημιουργήσει προβλήματα στην διαδικασία εξαγωγής των χαρακτηριστικών, σε περίπτωση που δεν είναι προσαρμοσμένη στα αλλοιωμένα δεδομένα. Επιπλέον, το μειονέκτημα σε αυτό το σύστημα είναι ότι περιορίζεται η συμβατότητά του σε βιογραφικά μόνο με τη μορφή CSV, ενώ θα έπρεπε να υποστηρίζει και άλλες δημοφιλείς μορφές αρχείων, όπως το DOC, DOCX και το PDF, και να τα μετατρέπει από την εκάστοτε μορφή στην υποστηριζόμενη.

5.3 Συγκριτική Αξιολόγηση Άρθρων

5.3.1 Κριτήρια Σύγκρισης

Από την ανάλυση που πραγματοποιήθηκε έγινε φανερό πως τα επιλεγμένα άρθρα, αν και έχουν ορισμένες ομοιότητες, διαφοροποιούνται σημαντικά με βάση το είδος των προσόντων που πραγματεύονται. Με αυτό τον τρόπο, η σύγκριση που θα

πραγματοποιήσουμε στη συνέχεια θα γίνει κυρίως ξεχωριστά για τα άρθρα που ασχολούνται με τα ουσιαστικά σε σχέση με αυτά που ασχολούνται με τα τυπικά προσόντα. Όμως, στο τέλος, θα γίνει και μια σύγκριση των άρθρων σε καθολικό επίπεδο με βάση τα κοινά τους χαρακτηριστικά.

Υπενθυμίζουμε πως τα ουσιαστικά προσόντα αναφέρονται σε μη απτά προσόντα, αυτά που συμπεραίνονται από στοιχεία της προσωπικότητας του ατόμου όπως ο τρόπος ομιλίας. Οι υπο-κατηγορίες αυτών των άρθρων χωρίζονται σε λεκτικά, μη λεκτικά, προσωδιακά, παραλεκτικά, κοινωνικά, γνωστικά και προσωπικά γνωρίσματα. Τα άρθρα ουσιαστικών προσόντων αντλούν τα δεδομένα τους από συνεντεύξεις. Τα τυπικά προσόντα αναφέρονται σε μετρήσιμα και αντικειμενικά προσόντα, όπως οι προϋπηρεσίες. Τα άρθρα που αφορούν τα τυπικά προσόντα, αντλούν τα δεδομένα κυρίως από βιογραφικά σημειώματα.

Τα κριτήρια σύγκρισης αποτελούν αναπόσπαστο κομμάτι για τη διαδικασία αξιολόγησης και τα αποτελέσματα. Για αυτό τον λόγο, τα κριτήρια που εφαρμόσαμε για την σύγκριση των άρθρων επιλέχθηκαν προσεκτικά για να καλύψουν τόσο τις ομοιότητες όσο και τις ιδιαιτερότητες των δύο προαναφερόμενων κατηγοριών άρθρων. Ειδικότερα, τα άρθρα τυπικών και ουσιαστικών προσόντων συγκρίνονται ξεχωριστά με βάση ορισμένα κοινά αλλά και ορισμένα ειδικά ως προς την κατηγορία τους κριτήρια. Η χρήση ειδικών κριτηρίων στα άρθρα ουσιαστικών προσόντων αφορά την ιδιαιτερότητά τους ως προς την ανάγκη ύπαρξης συνεντεύξεων και αξιολογητών. Από την άλλη μεριά, στα άρθρα των τυπικών προσόντων χρησιμοποιούνται δύο ειδικά κριτήρια που εστιάζουν κυρίως στη πηγή των δεδομένων τους και το πλήθος των δεδομένων αυτών. Σημειώνουμε πως η σύγκριση των άρθρων γίνεται και σε καθολικό επίπεδο με βάση τα κοινά τους κριτήρια, όπως αναφέρθηκε και προηγουμένως. Στη συνέχεια θα αναλύσουμε τα κοινά και τα ειδικά ως προς τα είδη προσόντων κριτήρια.

Ένα πρώτο κοινό κριτήριο σύγκρισης είναι ο σκοπός του κάθε άρθρου / προσέγγισης, το οποίο σχετίζεται με τον βασικό στόχο που προσπαθεί να επιτύχει το κάθε άρθρο. Το κριτήριο αυτό επιλέχθηκε διότι παρατηρήθηκε μια σημαντική διαφοροποίηση στους στόχους των άρθρων, η οποία έπρεπε να καταγραφεί και να αναλυθεί.

Ένα δεύτερο κοινό κριτήριο σύγκρισης αποτελούν οι εφαρμογές και τα εργαλεία που χρησιμοποιήθηκαν για την εξαγωγή, την ανάλυση και την επεξεργασία των δεδομένων. Κάθε εργαλείο και εφαρμογή διαχειρίζεται διαφορετικά δεδομένα και εξυπηρετεί ένα διαφορετικό σκοπό ή λειτουργία επεξεργασίας δεδομένων. Μέσω του κριτηρίου αυτού επομένως προσπαθούμε να καλύψουμε τόσο τα κοινά εργαλεία που χρησιμοποιούνται αλλά και να αναδείξουμε την ποικιλία τους ως προς την διαχείριση των δεδομένων στις επιλεγμένες προσεγγίσεις.

Ένα τρίτο κοινό κριτήριο σύγκρισης αποτελούν οι αλγόριθμοι MM και τα Νευρωνικά Δίκτυα (ΝΔ) που χρησιμοποιήθηκαν για εκπαίδευση. Μέσω του κριτηρίου αυτού και λαμβάνοντας υπ'όψιν την κατηγοριοποίηση των αλγορίθμων MM και των ΝΔ, θέλουμε να εξετάσουμε την ύπαρξη αλγορίθμων MM και ΝΔ που χρησιμοποιούνται από τα περισσότερα άρθρα καθώς και εκείνες τις κατηγορίες αλγορίθμων MM / ΝΔ που γίνονται εκμεταλλεύσιμες περισσότερο.

Το κοινό κριτήριο των αποτελεσμάτων αποσκοπεί να αναδείξει τα αποτελέσματα από την αξιολόγηση της απόδοσης της κάθε προσέγγισης, τις μετρικές αξιολόγησης που χρησιμοποιούνται καθώς και να ανιχνεύσει μετρικές αξιολόγησης που χρησιμοποιούνται σε πολλές προσεγγίσεις αλλά και προσεγγίσεις που έχουν μια βέλτιστη απόδοση.

Κατά την εξέταση των διαφόρων άρθρων, παρατηρήσαμε ορισμένες καινοτομίες που ήταν αρκετά ενδιαφέρουσες. Οι καινοτομίες συνδέονται με πρωτοποριακές ιδέες, προσεγγίσεις, εφαρμογές και μεθόδους, οι οποίες βελτιώνουν τις ήδη υπάρχουσες διαδικασίες. Επιπλέον, οι καινοτομίες κρίνονται για την βελτιστοποίηση των διαδικασιών, την περισσότερη αξιοπιστία, ακρίβεια και ταχύτητα που προσφέρουν. Επομένως, σκοπός του κριτηρίου των καινοτομιών είναι να αναδείξει όλες αυτές τις καινοτομίες και να τις συσχετίσει, όπου αυτό είναι εφικτό, καθώς και να εξετάσει αν κάποιες από αυτές προτείνονται από κοινού από περισσότερα του ενός άρθρα. Κατά την εξέταση των άρθρων, παρατηρήθηκε πως υπάρχουν καινοτομίες που αφορούν τους υποψηφίους, τους εργοδότες και την εξέλιξη της έρευνας.

Το τελευταίο κοινό κριτήριο αφορά τις αρνητικές πτυχές των προσεγγίσεων. Οι πτυχές αυτές αφορούν τους περιορισμούς, τις αδυναμίες και τα κενά των ερευνών, των διαδικασιών και των αποτελεσμάτων, τα οποία χρήζουν επίλυσης. Συνεπώς, επηρεάζουν την αξιοπιστία, την ακρίβεια, την απόδοση και την αμεροληψία των

ερευνών. Επομένως, οι αρνητικές πτυχές δείχνουν τα σημεία που χρήζουν βελτίωσης και για τα οποία απαιτείται περαιτέρω ανάπτυξη. Ως εκ τούτου, σκοπός του κριτηρίου αυτού είναι όχι μόνο να αναδείξει αυτές τις πτυχές αλλά και να εντοπίσει αν αυτές επαναλαμβάνονται συχνά, επισημαίνοντας μια δυνητικά σημαντική πρόκληση που θα πρέπει να αντιμετωπιστεί στο άμεσο μέλλον.

Ένα ειδικό ως προς την κατηγορία άρθρων κριτήριο αποτελούν οι συνεντεύξεις, οι οποίες αφορούν μόνο τα ουσιαστικά προσόντα, καθώς για τα τυπικά προσόντα, τα δεδομένα εξάγονται από τα βιογραφικά σημειώματα (οπότε δεν απαιτείται η διενέργεια συνεντεύξεων διότι το σύστημα προχωρά σε μια αυτοματοποιημένη αξιολόγηση των υποψηφίων με βάση τα σημειώματα αυτά). Οι συνεντεύξεις αποτελούν τη πηγή δεδομένων για τα ουσιαστικά προσόντα, από τις οποίες τα δεδομένα αρχικά συλλέγονται και έπειτα αξιολογούνται οι υποψήφιοι, είτε από το σύστημα είτε από αξιολογητές. Οι συνεντεύξεις χωρίζονται σε πραγματικές και εικονικές. Συνήθως, έχει παρατηρηθεί πως το είδος των συνεντεύξεων επηρεάζει την απόδοση των υποψηφίων. Επομένως, το κριτήριο αυτό εξετάζει πρωτίστως το είδος των συνεντεύξεων που πραγματοποιούνται. Επίσης, εκτός από τον τύπο της συνέντευξης, μας ενδιαφέρει και ο αριθμός των συνεντεύξεων. Όσο πιο μεγάλος είναι ο αριθμός, τόσο περισσότερα δεδομένα έχουμε και λογικά τόσο περισσότερη ακρίβεια μπορούμε να επιτύχουμε στην κατάταξη/αξιολόγηση των υποψηφίων. .

Από τα βίντεο που υπάρχουν, εκ των πραγμάτων, μόνο στα άρθρα των ουσιαστικών προσόντων εξάγονται όλες οι πληροφορίες των προσόντων αυτών. Αφορά την κύρια πηγή δεδομένων, το οπτικοακουστικό υλικό, για την καταγραφή και την ανάλυση των λεκτικών, μη λεκτικών, προσωδιακών κ.α. δεδομένων. Όσο περισσότερα είναι τα βίντεο, τόσα περισσότερα δεδομένα υπάρχουν για επεξεργασία, ανάλυση και εκπαίδευση, τα οποία είναι κρίσιμα για την ανάλυση της συμπεριφοράς και των συναισθημάτων. Επομένως, αυτά παίζουν ουσιώδη ρόλο για πιο αξιόπιστα αποτελέσματα. Το ίδιο ισχύει και για τους συμμετέχοντες, καθώς από αυτούς συλλέγονται και αναλύονται τα χαρακτηριστικά. Οπότε ο αριθμός τους πρέπει να είναι μεγάλος για να ενισχύεται η στατιστική δύναμη των ερευνών. Επίσης, η ποικιλομορφία των συμμετεχόντων οδηγεί σε αμερόληπτα συμπεράσματα. Συνεπώς, τόσο το κριτήριο του αριθμού των βίντεο όσο και το κριτήριο του αριθμού των συμμετεχόντων είναι εξίσου σημαντικά και για αυτό τον λόγο επιλέχθηκαν για τη σύγκριση των άρθρων των ουσιαστικών προσόντων.

Όσον αφορά το κριτήριο των αξιολογητών, ο αριθμός τους καθορίζει το πόση αντικειμενικότητα ή υποκειμενικότητα υπάρχει στις βαθμολογήσεις και στις σωστές κρίσεις των χαρακτηριστικών που καλούνται να αξιολογήσουν, είτε σε πραγματικό χρόνο είτε ασύγχρονα. Η ύπαρξή τους στις συνεντεύξεις επηρεάζει την συμπεριφορά των συμμετεχόντων και επιπλέον, οι υποψήφιοι νιώθουν εκτίμηση από τον εργοδότη με την παρουσία τους. Η εξειδίκευση και η εμπειρία των αξιολογητών επίσης επηρεάζει την ακρίβεια της αξιολόγησης. Ακόμα, ο μεγάλος αριθμός των αξιολογητών, αυξάνει το κόστος αλλά και τον χρόνο της διαδικασίας, κάτι που τα συστήματα πραγματοποιούν αυτόματα.

Τέλος, όπως προαναφέρθηκε, είχαμε δύο κριτήρια ειδικά ως προς την κατηγορία των άρθρων που εστιάζουν στα τυπικά προσόντα. Το πρώτο κριτήριο αφορά τον αριθμό των βιογραφικών με βάση τα οποία γίνεται η εκπαίδευση των αλγορίθμων MM. Προφανώς, όσο μεγαλύτερος είναι ο αριθμός των βιογραφικών αυτών, τόσο καλύτερη θα είναι η εκπαίδευση και θα οδηγεί σε περισσότερο αξιόπιστα αποτελέσματα. Επιπλέον, μπορεί να καλύπτονται και περισσότεροι τομείς με αυτό τον τρόπο. Από την άλλη, ένας μεγάλος αριθμός βιογραφικών μπορεί να οδηγεί σε μια καθυστέρηση στην εκπαίδευση, η οποία όμως δεν είναι τόσο μεγάλη σε σχέση με την περίπτωση των δεδομένων ουσιαστικών προσόντων.

Το δεύτερο ειδικό κριτήριο αφορά τη πηγή των δεδομένων που χρησιμοποιήθηκε για την άντληση των δεδομένων. Μέσω αυτού του κριτηρίου προσπαθούμε να εξετάσουμε δύο πτυχές: (α) ποια είναι η πηγή που χρησιμοποιείται περισσότερο, ειδικά ως προς την άντληση των βιογραφικών & (β) αν χρησιμοποιούνται εναλλακτικές πηγές και αν αυτές μπορούν να καλύψουν πλευρές των τυπικών προσόντων των υποψηφίων που δεν καλύπτονται από τα βιογραφικά σημειώματα.

5.3.2. Αποτελέσματα Σύγκρισης

Η σύγκριση των άρθρων πραγματοποιήθηκε με βάση τα ουσιαστικά και τα τυπικά προσόντα και στηρίχθηκε στα προαναφερόμενα κριτήρια. Παρακάτω παρατίθενται τρεις πίνακες σύγκρισης που συνοψίζουν τα αποτελέσματα: οι δύο πρώτοι πίνακες αναφέρονται στα ουσιαστικά προσόντα, ενώ ο τρίτος στα τυπικά προσόντα. Κάθε πίνακας περιλαμβάνει βασικές πληροφορίες για τα κοινά κριτήρια, όπως τον σκοπό

κάθε άρθρου, τους αλγορίθμους και τα εργαλεία που χρησιμοποιήθηκαν και τα αποτελέσματά τους. Στον πίνακα των ουσιαστικών προσόντων, όταν δεν αναφέρεται εάν οι συνεντεύξεις είναι πραγματικές ή εικονικές, ή συγκεκριμένος αριθμός τους, είναι επειδή αυτό δεν προσδιορίζεται στο αντίστοιχο άρθρο. Επίσης, στη στήλη του αριθμού των αξιολογητών (raters), δεν αναγράφεται πάντα αριθμός. Επιπλέον, υπάρχει ένα άρθρο όπου δεν προσδιορίζεται το είδος των αξιολογητών που χρησιμοποιήθηκαν. Ο συγκριτικός πίνακας των τυπικών προσόντων αποτελείται από λιγότερα κριτήρια σε σχέση με αυτόν με τα ουσιαστικά. Αυτό οφείλεται στο γεγονός πως για όλες τις κατηγορίες άρθρων έχουμε ορισμένα κοινά κριτήρια αλλά για την κατηγορία των άρθρων των ουσιαστικών προσόντων ο αριθμός των ειδικών κριτηρίων είναι μεγαλύτερος σε σχέση με αυτόν στην κατηγορία των άρθρων των τυπικών προσόντων.

Πίνακες 1 και 2 - Οι δύο πίνακες σύγκρισης απεικονίζουν τις πληροφορίες των άρθρων που εστιάζουν στα ουσιαστικά προσόντα.

ΑΡΘΡΟ	ΣΚΟΠΟΣ	ΕΡΓΑΛΕΙΑ	ΑΛΓΟΡΙΘΜΟΙ / ΝΔ	ΚΑΙΝΟΤΟΜΙΕΣ
2016 Rasipuram	Να μετρηθεί η ευχέρεια του λόγου	VoiceBase Automatic Speech Recognition- OPENSmile - PRAAT	Linear Regression, SupportVectorMachine, Logistic Regression, relief	Πραγματικού χρόνου ανατροφοδότηση
2016 Chen	Αυτόματη αξιολόγηση χρησιμοποιώντας πολυτροπικά (λεξιλογικά, ομιλία, οπτικά) γνωρίσματα για καλύτερη απόδοση πρόβλεψης	Linguistic Inquiry Word Count - Automated Scoring - ETS SpeechRater - Visage SDK Face Track - gazeDirectionGlobal - Emotient FACET SDK	SupportVectorMachine - LASSO - Ridge Regression	Χρήση του SSP
2017 Chen	Πρόβλεψη χαρακτηριστικών προσωπικότητας	ASR - pyAudioAnalysis - OpenFace - BoW	SVM, Random Forest	Χρήση του SSP, μεγάλος αριθμός δεδομένων, ποικιλομορφία
2018 Sawla	Απόδοση βασισμένη σε φωνητική ανάλυση υποψηφίου	PRAAT - Baruto package in R	RF, SVM, ANN	Χρήση των 7 Χαρακτηριστικών Προσωπικότητας σε συνδυασμό 4 προσωδιακών
2019 Hemamou	Σύγκριση των δειγμάτων attention and sample slices για την πρόβλεψη της δυνατότητας πρόσληψης	Gated Recurrent Unit (GRU) - OpenFace - Google Speech-to-Text API - Density-Based Spatial Clustering of Applications with Noise (DBSCAN)	HireNet - LASSO - RF - SVM	Οι υποψήφιοι λαμβάνουν τα αποτελέσματα, μέτρηση των attention slices
2020 Agrawal	Αυτόματη αξιολόγηση σε συνεντεύξεις εργασίας	Discrete Fourier Transform (DFT) - Inverse DFT - PRAAT - pyAudioAnalysis - OpenCV - LeNet - Google Cloud Speech-to-Text API - Natural Language Toolkit (NLTK) - Tone Analyzer - Stanford Named Entity Recognizer (NER)	RF, SVM, MLP, LASSO	Ανατροφοδότηση στους υποψηφίους. μεγάλος αριθμός δεδομένων
2021 Hickman	Ανάπτυξη αλγορίθμων MM για την εξαγωγή συμπερασμάτων της προσωπικότητας των υποψηφίων	LIWC - IBM Watson Speech-to-Text - R-CARET	Ridge Regression - LASSO	Ανάλυση για: «πώς» λέγεται κάτι και όχι «τι» λέγεται
2021 Booth	Προκατάληψη και δικαιοσύνη σε ένα πολυτροπικό πλαίσιο MM που αξιολογεί τη δυνατότητα προσλήψεως (και την αφαίρεση χαρακτηριστικών φύλου)	LIWC - OpenSmile - Emotient FACET - Motion Tracker	RF - elastic net	Αφαίρεση χαρακτηριστικών φύλου
2023 Anju Raj	Ένα σύστημα που αναλύει τις εκφράσεις του προσώπου των υποψηφίων και τα γλωσσικά και προσωδιακά στοιχεία από βιντεοσκοπημένες συνεντεύξεις	PRAAT - Haar Cascade - Valence Aware Dictionary and sentiment Reasoner (VADER)	Μη αναφερόμενο	Χρήση γραφήματος χαμόγελου σε πραγματικό χρόνο από τους αξιολογητές

ΑΡΘΡΑ	ΣΥΝ/ΞΕΙΣ	ΑΡΙΘΜΟΣ ΣΥΝ/ΞΕΩΝ	ΑΤΟΜΑ	ΒΙΝΤΕΟ	ΑΡΙΘΜΟΣ ΑΞΙΟΛΟΓΗΤΩΝ	ΕΙΔΗ ΑΞΙΟΛΟΓΗΤΩΝ	ΑΠΟΤΕΛΕΣΜΑΤΑ	ΑΡΝΗΤΙΚΑ
2016 Rasipuram	Εικονικές	106	106	106	3	Εξωτερικοί παρατηρητές	Πρόβλεψη για ακρίβεια ακρίβειας: προσωπικά 83% & γλωσσικά 78%	Λίγες και εικονικές συνεντεύξεις
2016 Chen	Μη ορισμένο	419	36	419	9	9 Σύμβουλοι & απροσδιόριστος αριθμός Amazon Mechanical Turk Workers	Ο συνδυασμός στοιχείων προβλέπει καλύτερα τις απαντήσεις των αξιολογητών: SVM 44.6%, LASSO 45,2% & RR 45,8%	Εικονικές συνεντεύξεις, λίγα άτομα, χαμηλά ποσοστά αποτελεσμάτων, απώλεια δεδομένων
2017 Chen	Εικονικές	1891	260	1891	5	Μη ορισμένη	Precision, recall&F-1 (78% -86%) για πρόβλεψη των 5 Μεγάλων Χαρακτηριστικών Προσωπικότητας από κείμενο και 63%-67% για πρόσληψη	Εικονικές συνεντεύξεις
2018 Sawla	Εικονικές	138	69	138	1	Σύστημα	RF καλύτερος με ακρίβεια πρόβλεψης 70,52% - 79,76% στις απαντήσεις των αξιολογητών	Λίγες συνεντεύξεις, «θόρυβος» στα βίντεο
2019 Hemamou	Πραγματικές	7938	7938	7938	1	Σύστημα	Τα Attention slices προβλέπουν καλύτερα τη δυνατότητα προσλήψεως από τα random slices (RF 0.554, LASSO 0.550, SVM Linear 0.543) & HireNet 61.8%	Χαμηλά ποσοστά
2020 Agrawal	Εικονικές	138	69	138	Μη ορισμένη	Amazon Mechanical Turk Workers	Ο RF καλύτερη ακρίβεια με όλα τα χαρακτηριστικά/παραμέτρους έως και 96,43%	Λίγες και εικονικές συνεντεύξεις
2021 Hickman	Εικονικές	441	411	441	Μη ορισμένη	Προπτυχιακοί βοηθοί ερευνητές	Οι προβλέψεις MM συσχετίζονται μέτρια με τους αξιολογητές	Εικονικές συνεντεύξεις, μοναδική μετρική το Pearson correlation, χαμηλά αποτελέσματα
2021 Booth	Μη ορισμένο	524	733	4255	Μη ορισμένη	Εκπαιδευμένοι αξιολογητές	Correlational accuracy .46% για όλα μαζί τα χαρακτηριστικά (λεκτικά, παραλεκτικά και οπτικά)	Χαμηλά αποτελέσματα
2023 Anju Raj	Εικονικές	Μη ορισμένη	Μη ορισμένη	Μη ορισμένη	6	Εκπαιδευμένοι επαγγελματίες	Παρουσίαση του συστήματος	Μη ορισμένα αποτελέσματα

Πίνακας 3 - Πίνακας σύγκρισης τυπικών προσόντων.

ΑΡΘΡΟ	ΣΚΟΠΟΣ	ΑΡΙΘΜΟΣ ΒΙΟΓΡΑΦΙΚΩΝ	ΠΗΓΗ ΔΕΔΟΜΕΝΩΝ	ΕΡΓΑΛΕΙΑ	ΑΛΓΟΡΙΘΜΟΙ / ΝΔ	ΑΠΟΤΕΛΕΣΜΑ	ΚΑΙΝΟΤΟΜΙΑ	ΑΡΝΗΤΙΚΑ
2020 Roy	Σύστημα συστάσεων βιογραφικού	Μη ορισμένος	Kaggle	NLTK, Stop Word Removal, Stemming, Lemmatization, Data Masking, TF-IDF, BoVW, scikit learn, Cosine Similarity	KNN, RF, Logistic Regression, Linear Support Vector Classifier, Multinomial Naive Bayes	Καλύτερη Accuracy 78,53% Linear Support Vector Classifier με KNN (ταξινόμηση για ορθή θέση εργασίας)	Χρήση Data Masking για GDPR	Συμβατότητα μόνο για csv
2022 Tejaswini	Σύστημα κατάταξης βιογραφικών	50	Kaggle	TF-IDF, Gensim, NLP, Cosine Similarity	KNN	Parsing accuracy από 85% με 100%, Ranking accuracy από 75% με 100%	Χρήση μηχανής προτάσεων, φόρμα πολλαπλής επιλογής πριν δοθεί το βιογραφικό	Λίγα δεδομένα
2020 Suhas	Αντιστοίχιση Βιογραφικού με θέση εργασίας	Μη ορισμένος	Stack exchange	NLP, Stop Word Removal, NER, Word Embedding, Cosine Similarity	word2vec	Recall 80% (ορθές αντότητες), F-score 84% (ορθές δεξιότητες), precision 82% (ορθά αποτελέσματα)	Εκπαίδευση δεδομένων από "stack exchange"	Εκπαίδευση μόνο για IT θέσεις εργασίας
2020 Jagan	Πρόβλεψη Προσλήψεων	Μη ορισμένος	Δημιουργημένα από τους ίδιους	Pearson Correlation Coefficient	Decision Tree, Random Forest, Gaussian Naive Bayes (GNB), KNN	Ο GNB 99% (πρόβλεψη πρόσληψης), 100% (ανάληψη θέσης), 88% (παραμονή στη θέση)	Πρόβλεψη παραμονής στην θέση εργασίας	Απουσία μεθοδολογίας
2023 Farzana	Αυτοματοποιημένη πρόταση βιογραφικού	Μη ορισμένος	Kaggle, GitHub	MongoDB, HTML CSS, Flask, NLP, SpaCy	Μη καθορισμένοι	Λίστα με καλύτερα βιογραφικά	Χρήση ιστοσελίδας για κατάθεση βιογραφικών	Ασφάλεια ιστοσελίδας, απουσία αναφοράς κόστους
2023 Chavare	Αναλυτής Βιογραφικών	Μη ορισμένος	Μη καθορισμένη πηγή	NLP, Spacy NER, UBLAI	BERT, Joint NER, Relation Extraction, Connection Extraction	-	Knowledge Graph (σύνδεση δεξιοτήτων με εμπειρία)	Απουσία αναφοράς δεδομένων
2024 Naveed	Κατάταξη Βιογραφικών	20	Dataturks platform	NLP, Visual Studio Code, MERN Framework, tokenization, Stop Word Removal, NLTK, Text Blob, SpaCy, NER	Μη καθορισμένοι	F1 Score, Recall, Precision almost 100% (in finding correct tokens)	Ιστοσελίδα με ικανότητες στη διαχείριση βιογραφικών, χρήση της Dataturks platform	Ασφάλεια ιστοσελίδας, απουσία αναφοράς κόστους
2022 Naik	Οι 10-20 Προτάσεις Βιογραφικών	Μη ορισμένος	Μη καθορισμένη πηγή	tokenization, NLTK, NLP, Cosine Similarity	Naive Bayes, KNN, SVM	Ο KNN accuracy 99% για την αντιστοίχιση	Διαχείριση πολλών βιογραφικών	Βάση σε αντιστοίχιση και όχι κατάταξη
2022 Kinge	Έλεγχος Βιογραφικών	Μη ορισμένος	Kaggle, GitHub & LinkedIn	NLP, NER, Cosine Similarity, NLP, SpaCy, Stop Word Removal, Stemming, Lemmatization, TF-IDF, Cosine Similarity, EDA, Gensim	KNN, SVM	Από 78% με 98%	Διαχείριση πολλών βιογραφικών, ανατροφοδότηση σε υποψηφίους για καλύτερο βιογραφικό, εξαγωγή από GitHub & LinkedIn	Δοκιμή μόνο στον κλάδο Πληροφορικής
2023 Sowjanya	Έξυπνος Αναλυτής Βιογραφικών	50	Kaggle	NLP, NER, Cosine Similarity, NLP, SpaCy, Stop Word Removal, Stemming, Lemmatization, TF-IDF, Cosine Similarity, EDA, Gensim	ResumeNet, Latent Dirichlet Allocation (LDA), KNN	85% text parsing accuracy, 92% ranking accuracy	Δωρεάν ανατροφοδότηση & YouTube βίντεο για προετοιμασία, οπτικοποιήσεις	Λίγα δεδομένα

5.3.3 Συμπεράσματα Σύγκρισης

5.3.3.1 Συμπεράσματα Άρθρων που Εστιάζουν στα Ουσιαστικά Προσόντα.

Στα άρθρα που εστιάζουν στα ουσιαστικά προσόντα κοινός σκοπός είναι να βελτιστοποιηθεί η πρόβλεψη αυτόματης αξιολόγησης χρησιμοποιώντας τα προσόντα αυτά. Στα άρθρα αυτά χρησιμοποιήθηκαν πλείστα εργαλεία και εφαρμογές. Τα περισσότερα επικεντρώνονται στα παραλεκτικά, προσωδιακά, λεκτικά και μη λεκτικά χαρακτηριστικά για την αξιολόγηση των υποψήφιων εργαζόμενων που απορρέουν από τις συνεντεύξεις. Από αυτά τα χαρακτηριστικά εξάγουν τις πληροφορίες που είναι απαραίτητες για την εκπαίδευση των αλγορίθμων. Όμως, οι προσεγγίσεις των ερευνητών διαφέρουν και εξετάζουν διαφορετικά χαρακτηριστικά και στοιχεία (ουσιαστικών προσόντων), όπως την ορθή εξαγωγή χαρακτηριστικών, την προκατάληψη λόγω φύλου και την πιθανότητα πρόσληψης. Αρκετές μελέτες χρησιμοποίησαν ως βάση τα “Πέντε Μεγάλα χαρακτηριστικά προσωπικότητας” (Big Five personality traits) και έδωσαν έμφαση στη χρήση πολλαπλών στοιχείων δεδομένων (μη λεκτικά, λεκτικά, προσωδιακά). Επομένως, στις υπόλοιπες κατηγορίες ουσιαστικών προσόντων ουσιαστικά παρατηρείται ένα κενό.

Οι σκοποί των άρθρων ουσιαστικών προσόντων είναι: η πρόβλεψη της ευχέρειας του λόγου (fluency), η δημιουργία ενός συστήματος αυτόματης αξιολόγησης για βίντεο-συνεντεύξεις με μονόλογο με τη χρήση πολυτροπικών δεδομένων για καλύτερη απόδοση πρόβλεψης, η πρόβλεψη χαρακτηριστικών προσωπικότητας, η ρητορική απόδοση βασισμένη στη φωνητική ανάλυση του υποψηφίου, η σύγκριση των attention με τα sample slices για την πρόβλεψη της δυνατότητας πρόσληψης, η αυτόματη αξιολόγηση σε συνεντεύξεις εργασίας, η ανάπτυξη αλγορίθμων MM για την εξαγωγή συμπερασμάτων της προσωπικότητας του υποψηφίου, η προκατάληψη φύλου σε ένα πολυτροπικό πλαίσιο MM που αξιολογεί τη δυνατότητα πρόσληψης και τέλος, η δημιουργία ενός συστήματος που αναλύει τις εκφράσεις του προσώπου και τα γλωσσικά και προσωδιακά γνωρίσματα των υποψηφίων από βιντεοσκοπημένες συνεντεύξεις.

Τα εργαλεία με την περισσότερη χρήση στα ουσιαστικά προσόντα είναι ο LASSO, που αποτρέπει την υπερπροσαρμογή, το PRAAT, που επεξεργάζεται τη φωνή, οι εφαρμογές “speech-to-text”, το VoiceBase ASR (Automatic Speech Recognition) και το Linguistic Inquiry Word Count (LIWC), που κατηγοριοποιεί τα λεξιλογικά χαρακτηριστικά που περιγράφουν συναισθήματα. Συνεπώς, τα περισσότερα εργαλεία που χρησιμοποιήθηκαν φαίνεται να εστιάζουν στην ανάλυση λόγου και φωνής, πράγμα αρκετά λογικό λόγω της φύσης των ουσιαστικών προσόντων.

Οι αλγόριθμοι MM που χρησιμοποιούνται περισσότερο είναι οι Random Forest (RF) και ο Support Vector Machine. Την μεγαλύτερη πιστότητα (accuracy) πέτυχε ο Random Forest χρησιμοποιώντας ταυτόχρονα όλες τις παραμέτρους (ήχο, λεκτικά στοιχεία, βίντεο). Όμως, μια απόλυτη σύγκριση των άρθρων δεν μπορεί να πραγματοποιηθεί επειδή δεν ασχολούνται όλα τα άρθρα με τις ίδιες μετρικές (που να αξιολογούν το ίδιο ακριβώς πράγμα) και τα ίδια δεδομένα. Παρόλα αυτά, η μετρική της ακρίβειας (precision) χρησιμοποιήθηκε περισσότερο στα άρθρα αυτά.

Καινοτομίες στα άρθρα των ουσιαστικών προσόντων είναι: η ποικιλομορφία σε κάποιες έρευνες όσον αφορά το φύλο, την εθνικότητα κ.α., η ανατροφοδότηση για τους υποψηφίους σε πραγματικό χρόνο, η χρήση των 7 Χαρακτηριστικών Προσωπικότητας σε συνδυασμό με τέσσερα προσωδιακά χαρακτηριστικά, η μέτρηση των attention slices, η αφαίρεση χαρακτηριστικών φύλου, η χρήση γραφήματος χαμόγελου σε πραγματικό χρόνο από τους αξιολογητές και η χρήση της Επεξεργασίας του Κοινωνικού Σήματος, Social Signal Processing (SSP).

Στα αρνητικά συγκαταλέγονται η χρήση εικονικών συνεντεύξεων και ο μικρός αριθμός δεδομένων (συνεντεύξεων & βίντεο) με αποτέλεσμα η πλειοψηφία των προβλέψεων να έχει μικρά ποσοστά ακρίβειας. Επιπλέον, πρέπει να αναλογιστεί κάποιος ότι η σύγκριση των μετρήσεων έγινε σε σχέση με τους αξιολογητές και αυτό σημαίνει ότι υπήρχε το στοιχείο της υποκειμενικότητας, καθώς τα ουσιαστικά προσόντα δεν είναι κάτι το εύκολο και αντικειμενικά μετρήσιμο, όπως είναι τα τυπικά προσόντα. Επίσης, ο αριθμός των αξιολογητών στη πλειοψηφία των άρθρων είναι μικρός. Αυτό επηρεάζει την αντικειμενικότητα, αφού με περισσότερους αξιολογητές ενδέχεται να είχαμε διαφορετική αξιολόγηση.

Με βάση το κριτήριο των δεδομένων, θα ξεχώριζε το άρθρο των Hematou και συν. (2019) λόγω του όγκου δεδομένων (7938 πραγματικές συνεντεύξεις) που χρησιμοποίησαν, όμως απορρίπτεται εξαιτίας των χαμηλών ποσοστών ακρίβειας. Το άρθρο που τελικά ξεχώρισε συνολικά, λαμβάνοντας υπόψη κυρίως τα κριτήρια των δεδομένων, της καινοτομίας και των αποτελεσμάτων, είναι αυτό των Chen και συν. (2017), λόγω των τριών μετρικών που χρησιμοποιήθηκαν (Precision, Recall, F1-score) και με ποσοστά από 78% μέχρι και 86% για την πρόβλεψη των Πέντε Μεγάλων Χαρακτηριστικών και από 63% μέχρι και 67% για την πρόβλεψη πρόσληψης, όπως επίσης και λόγω του μεγάλου αριθμού συνεντεύξεων (1891), αν και εικονικών, καθώς και του αριθμού των καινοτομιών.

5.3.3.2 Συμπεράσματα Άρθρων που Εστιάζουν στα Τυπικά Προσόντα.

Οι σκοποί των άρθρων τυπικών προσόντων είναι: η δημιουργία συστημάτων προτάσεων, διαλογής, κατάταξης, ανάλυσης βιογραφικών σημειωμάτων, η αντιστοίχιση βιογραφικών με θέσεις εργασίας, η πρόβλεψη πρόσληψης και τέλος η δημιουργία ενός έξυπνου αναλυτή βιογραφικών. Οπότε, μπορούμε να συμπεράνουμε πως τα περισσότερα άρθρα εστίασαν στην κατάταξη και το ταίριασμα των βιογραφικών με θέσεις εργασίας.

Στα άρθρα που ασχολήθηκαν με τα τυπικά προσόντα επίσης χρησιμοποιήθηκαν αρκετά εργαλεία. Στην μεγάλη πλειοψηφία των άρθρων τυπικών προσόντων χρησιμοποιήθηκαν εργαλεία Επεξεργασίας Φυσικής Γλώσσας (NLP). Ειδικότερα, το Named Entity Recognition (NER), που αναγνωρίζει και ταξινομεί οντότητες και το Stop Word Removal που αφαιρεί τις ανούσιες λέξεις, χρησιμοποιήθηκαν σε μεγάλο αριθμό άρθρων τυπικών προσόντων. Η μεγαλύτερη διαφορά συγκριτικά με τα προηγούμενα άρθρα έγκειται στο ότι σε εκείνα υπήρχαν άνθρωποι αξιολογητές, ενώ σε αυτά τα άρθρα η αξιολόγηση πραγματοποιείται από ένα σύστημα-μοντέλο. Αυτό οφείλεται στο γεγονός πως η αξιολόγηση των τυπικών προσόντων μπορεί να γίνει πιο εύκολα και με αντικειμενικό τρόπο.

Οι βασικότεροι αλγόριθμοι που χρησιμοποιήθηκαν είναι ο KNN και ο Random Forest, πράγμα το οποίο παρατηρήσαμε και για τις προσεγγίσεις ουσιαστικών προσόντων. Σε κάποιες περιπτώσεις χρησιμοποιήθηκαν ειδικοί αλγόριθμοι, όπως ResumeNet και

Latent Dirichlet Allocation (LDA). Τα ποσοστά αποτελεσμάτων στα άρθρα με τα τυπικά προσόντα κυμαίνονται από 75% μέχρι 100%. Και εδώ δεν μπορεί να υπάρξει απόλυτη σύγκριση αφού δεν αξιολογούν όλοι οι ερευνητές τα ίδια πεδία/αντικείμενα και δεν χρησιμοποιούν τις ίδιες μετρικές. Η μετρική που χρησιμοποιήθηκε περισσότερο είναι η πιστότητα (accuracy). Κάποιοι εστιάζουν στην ορθή εύρεση των οντοτήτων, άλλοι στο εάν ο υποψήφιος θα παραμείνει στη νέα του θέση, άλλοι στην αντιστοίχιση των βιογραφικών με τη θέση εργασίας και άλλοι στην ακρίβεια των εξαγόμενων στοιχείων από τα κείμενα. Όμως, εστιάζοντας στο βασικό αντικείμενο της διπλωματικής εργασίας, η πιστότητα (accuracy) 99% στην εργασία των Naik & Dhotre (2022) φαίνεται να είναι η καλύτερη καθώς πέτυχε αυτό το ποσοστό στο ταίριασμα των βιογραφικών με συγκεκριμένη θέση εργασίας.

Καινοτομίες που παρατηρήθηκαν στα άρθρα των τυπικών προσόντων αποτελούν: η χρήση της μεθόδου Data Masking για την διαφύλαξη των ευαίσθητων προσωπικών δεδομένων, η συμπλήρωση φόρμας ερωτήσεων πριν την αποδοχή των υποψηφίων από το σύστημα, η έρευνα για την πρόβλεψη της παραμονής του εργαζόμενου στην εργασία του μετά από την πρόσληψη, η διαδικασία της πρόσληψης εξ' ολοκλήρου μέσα από μια ιστοσελίδα δημιουργούμενης από τους ερευνητές για την διευκόλυνση της διεξαγωγής της διαδικασίας εύρεσης θέσης της σωστής εργασίας (και από την πλευρά του υποψηφίου και από την πλευρά του εργοδότη), τα οπτικοποιημένα αποτελέσματα των υποψηφίων για διευκόλυνση των εργοδοτών, η χρήση στοιχείων / δεδομένων από άλλες πλατφόρμες εκτός του Kaggle (όπως LinkedIn & Github) και τέλος η ανατροφοδότηση των υποψηφίων για βελτίωση των βιογραφικών σημειωμάτων τους και της απόδοσής τους σε συνεντεύξεις, όπως και η πρόταση για παρακολούθηση βίντεο σχετικών με τη βελτίωση δημιουργίας του βιογραφικού και με τη προετοιμασία συνέντευξης στο YouTube.

Στα αρνητικά των άρθρων με τα τυπικά προσόντα συγκαταλέγονται: το μικρό δείγμα δεδομένων από τα βιογραφικά σημειώματα, η απώλεια και η αλλοίωση δεδομένων, τα χαμηλά ποσοστά των μετρικών αξιολόγησης σε κάποιες περιπτώσεις, η χρήση ενός μόνο εργαλείου ή αλγορίθμου, η απουσία αποτελεσμάτων λόγω εστίασης των ερευνητών στη λειτουργία του προτεινόμενου συστήματος, η απουσία αποτελεσμάτων λόγω εστίασης των ερευνητών στη λειτουργία του προτεινόμενου συστήματος και η αμφίβολη ασφάλεια των ιστοσελίδων. Τα βιογραφικά σημειώματα

πάρθηκαν κυρίως από πλατφόρμες (π.χ. Kaggle) και ο αριθμός τους ήταν μη ικανοποιητικός. Ενδιαφέρον θα είχε τα βιογραφικά να ήταν όλα πραγματικά και για πραγματικές συνθήκες. Τα βιογραφικά που πάρθηκαν από το Kaggle ήταν μεν πραγματικά, αλλά τυχαία και όχι για τη συγκεκριμένη θέση εργασίας. Η χρήση επιπλέον πλατφορμών όπως Github & LinkedIn θεωρείται επιβεβλημένη για την άντληση επιπλέον πραγματικών στοιχείων τυπικών προσόντων που μπορεί να μην καλύπτονται από τα υποβεβλημένα βιογραφικά.

Τα άρθρα που ξεχωρίζουν είναι τρία, αυτό των Tejaswini και συν. (2022), των Naveed και συν. (2024) και των Sowjanya και συν. (2023) επειδή όλα αναφέρονται στον αριθμό των βιογραφικών σημειωμάτων που χρησιμοποιούν, στην πηγή δεδομένων, επίσης χρησιμοποιούν αρκετά εργαλεία επεξεργασίας (από τέσσερα - Tejaswini και συν., έως και 9 - Sowjanya & Naveed), και έχουν ποσοστά από 75% μέχρι 100% (ανάλογα με το τι μετρούν). Αν λάβουμε όμως υπ'όψιν επιπρόσθετα τον αριθμό των εργαλείων και τον αριθμό των αλγορίθμων MM που χρησιμοποιήθηκαν, τότε μπορούμε να πούμε με σιγουριά πως η προσέγγιση των Sowjanya και συν. (2023) μπορεί να θεωρηθεί ως η καλύτερη.

5.3.3.3 Καθολικά Συμπεράσματα

Και οι δύο κατηγορίες των άρθρων αποσκοπούν στην αυτόματη αξιολόγηση και στην αντιστοίχιση ενός υποψήφιου εργαζόμενου με μία θέση εργασίας χωρίς την εμπλοκή του ανθρώπινου παράγοντα στην όλη τη διαδικασία. Τα αποτελέσματα δείχνουν να είναι καλύτερα για τα άρθρα που εστιάζουν στα τυπικά προσόντα, ενώ τα άρθρα που εστιάζουν στα ουσιαστικά προσόντα δείχνουν ότι έχουν μεγαλύτερα περιθώρια βελτίωσης. Αυτό είναι βέβαια λογικό αφού τα πρώτα είναι μετρήσιμα και ποσοτικά γνωρίσματα ενώ τα δεύτερα μη μετρήσιμα και ποιοτικά, αν όχι υποκειμενικά, οπότε είναι δυσκολότερο να βγουν ορθά αποτελέσματα. Τα ουσιαστικά προσόντα επικεντρώνονται στην πρόβλεψη προσωπικών χαρακτηριστικών για την περαιτέρω σύνδεσή τους με μία θέση εργασίας, ενώ τα τυπικά προσόντα δίνουν έμφαση στην καλύτερη αντιστοίχιση ενός βιογραφικού σημειώματος με μία θέση εργασίας.

Η πλειοψηφία Νευρωνικών Δικτύων που χρησιμοποιήθηκαν ανήκουν στα βαθιά. Η μετρική του accuracy (πιστότητα) είναι αυτή που και στις δύο κατηγορίες

χρησιμοποιείται περισσότερο, με την μετρική του precision (ακρίβεια) να έρχεται δεύτερη. Τέλος, η βιβλιοθήκη/εργαλείο NLTK και η βιβλιοθήκη SpaCy είναι οι δύο βιβλιοθήκες/εργαλεία που χρησιμοποιήθηκαν περισσότερο στα παραπάνω άρθρα.

Στα αρνητικά των άρθρων συγκαταλέγονται τα εξής: το μικρό πλήθος δεδομένων για εκπαίδευση, είτε από τις συνεντεύξεις είτε από τα βιογραφικά σημειώματα, η απώλεια δεδομένων, η κακή ποιότητα των δεδομένων, τα χαμηλά ποσοστά των μετρικών αξιολόγησης και η χρήση ενός μόνο εργαλείου ή αλγορίθμου.

6. ΣΥΜΠΕΡΑΣΜΑΤΑ

6.1 Συνεισφορά Διπλωματικής & Σημαντικότερα Αποτελέσματα

Η παρούσα διπλωματική εργασία επιδιώκει να συμβάλει στη διερεύνηση της αυτοματοποίησης διαδικασιών πρόσληψης προσωπικού με τη συμβολή της Μηχανικής Μάθησης. Προϋπόθεση για επιτυχή αποτελέσματα είναι η ορθή εξαγωγή χαρακτηριστικών και δεξιοτήτων των υποψηφίων με βάση τα οποία έπειτα μπορεί να γίνει φιλτράρισμα και κατάταξη των υποψηφίων σε σχέση με μια θέση εργασίας. Η εργασία επικεντρώθηκε αρχικά στην ανάλυση και ακολούθως στη σύγκριση (με βάση ένα καλά σχεδιασμένο σύνολο από κριτήρια) των ερευνητικών άρθρων που εστιάζουν είτε σε ουσιαστικά είτε σε τυπικά προσόντα. Σκοπός ήταν η αξιολόγηση της αποτελεσματικότητας των σύγχρονων προσεγγίσεων που στηρίζονται στη Μηχανική Μάθηση και εφαρμόζονται σε αυτή τη διαδικασία, ώστε αυτή να αποτελέσει τη βάση για πρόσθετη σχεδίαση και ανάπτυξη των ήδη υπάρχουσών προσεγγίσεων και υιοθέτηση τους, ιδιαίτερα στο δημόσιο τομέα. Η συγκριτική ανάλυση ενός μεγάλου αριθμού από σχετικές έρευνες από την παρούσα διπλωματική εργασία δίνει ώθηση προς αυτή τη κατεύθυνση. Μάλιστα, λόγω των πολλών συμπερασμάτων που προκύπτουν από την ανάλυση αυτή, η διπλωματική αυτή εργασία μπορεί να αποτελέσει έναν σημαντικό οδηγό για την μελλοντική βελτίωση των διαδικασιών πρόσληψης προσωπικού.

Η ανάλυση των άρθρων ανέδειξε ότι τα αποτελέσματα, με ποσοστό έως και 100% αλλά σε μικρό δείγμα δεδομένων, είναι καλύτερα όταν τα στοιχεία προέρχονται από

τα τυπικά προσόντα καθώς αυτά είναι μετρήσιμα. Επιπλέον, αναδείχθηκε πως η χρήση της επεξεργασίας της φυσικής γλώσσας (NLP), ταυτόχρονα όμως με τα κατάλληλα εργαλεία και τις εφαρμογές, επιτρέπει την βέλτιστη αντιστοίχιση υποψηφίων με μία θέση εργασίας. Αν και στα ουσιαστικά προσόντα υπάρχει ποσοστό ακρίβειας που ανέρχεται στο 96.43%, το οποίο αφορά την πρόβλεψη βαθμολογίας των αξιολογητών με τον αλγόριθμο Random Forest, τα περισσότερα αποτελέσματα είναι μέτρια και χρειάζονται βελτίωση. Αυτό οφείλεται στο γεγονός πως η πολυπλοκότητα των ουσιαστικών προσόντων και η ποιοτική φύση τους επηρεάζουν την αξιοπιστία των αποτελεσμάτων και δυσκολεύουν ακόμα περισσότερο την ποσοτικοποίησή τους.

6.2 Κρίσιμα Ερευνητικά Κενά & Κατευθύνσεις

Κενά των ερευνών είναι η έλλειψη πολλών, ποιοτικών, ποικιλόμορφων και πραγματικών δεδομένων και συνθηκών για την ορθότερη και αποδοτικότερη εκπαίδευση των αλγορίθμων. Με περισσότερα, πιο ποικιλόμορφα, ποιοτικά, αλλά και πραγματικά σύνολα δεδομένων, οι αλγόριθμοι θα γίνουν πιο αξιόπιστοι. Αδυναμία στα περισσότερα άρθρα ήταν ο περιορισμένος χρόνος των συνεντεύξεων καθώς και η μη εξειδίκευση, σε ορισμένες περιπτώσεις, των αξιολογητών. Και στις δύο κατηγορίες των άρθρων, πρέπει να γίνουν περαιτέρω έρευνες σε πραγματικές συνθήκες, με πραγματικά και περισσότερα δεδομένα, με την προσθήκη επιπλέον πηγών δεδομένων, με την συμπερίληψη περισσότερων δημογραφικών ομάδων, με την προσθήκη περισσότερων ξένων γλωσσών, καθώς και για υπαρκτές θέσεις εργασίας. Η πρόταση για δημιουργία μιας βιβλιοθήκης/εργαλείο βασισμένη σε λεξιλόγιο και ορολογία που αφορά τις συνεντεύξεις εργασίας, εκφράστηκε από κάποιους συγγραφείς, ώστε να εξαγονται χαρακτηριστικά με περισσότερη εξειδίκευση. Γενικότερα, στα άρθρα με τα τυπικά προσόντα φαίνεται ότι οι έρευνες είναι πιο ώριμες και πιο κοντά στην εφαρμογή τους από ολόένα και περισσότερους εργοδότες ανά τον κόσμο.

Το ιδεατό είναι οι προσλήψεις να πραγματοποιούνται αυτόματα, αποτελεσματικά, δίκαια, αντικειμενικά, χωρίς σφάλματα, προκαταλήψεις και κόστος. Τα συστήματα ενισχύουν τις διαφανείς και δίκαιες διαδικασίες αξιολόγησης και πρόσληψης. Αυτό είναι κάτι που μπορεί να επιτευχθεί με την αμερόληπτη ανάλυση δεδομένων. Βέβαια, δεν πρέπει να λησμονείται ότι οι αλγόριθμοι Μηχανικής Μάθησης εκπαιδεύονται από

ανθρώπους και οι συγκρίσεις των αποτελεσμάτων (των αξιολογητών και των συστημάτων) στα άρθρα πραγματοποιήθηκαν σε σχέση με την υποκειμενικότητα του αξιολογητή, ειδικά στις μετρήσεις των ουσιαστικών προσόντων. Η σημασία της αντικειμενικότητας σε αυτήν τη διαδικασία δεν μπορεί να παραβλεφθεί. Μία κατεύθυνση μελλοντικών ερευνών θα ήταν η εισαγωγή των δεδομένων από μεγάλο αριθμό ανθρώπων, από ποικίλες εθνότητες, πολιτισμούς, φυλές, ηλικίες, υπόβαθρο, δηλαδή από πιο αντιπροσωπευτικό δείγμα ανθρώπων. Με αυτόν τον τρόπο, θα μπορούσε να βελτιωθεί η αντικειμενικότητα. Η χρήση περισσότερων αξιολογητών είναι μια άλλη σχετική κατεύθυνση για την μείωση της υποκειμενικότητας. Η χρήση μοντέλων LLMs (Large Language Models) είναι μια άλλη ενδιαφέρουσα περίπτωση διότι τα μοντέλα αυτά έχουν εκπαιδευτεί με έναν τεράστιο όγκο δεδομένων και μπορούν να υλοποιήσουν πολύ εύκολα λειτουργίες NLP και αξιολόγησης. Μάλιστα, θα μπορούσαν να παίξουν τον ρόλο ενός “αντικειμενικού” αξιολογητή που χρησιμοποιείται επιπρόσθετα των πραγματικών.

Η ικανότητα ενός συστήματος να αποδίδει ακόμα και με ανακριβή ή ελλιπή δεδομένα είναι κάτι που το θέτει αξιόπιστο. Όμως, παρατηρήθηκε σε κάποια άρθρα αλλοίωση των αποτελεσμάτων λόγω χαμηλής ποιότητας δεδομένων ή λανθασμένης ερμηνείας συναισθημάτων και χαρακτηριστικών από τα δεδομένα. Επίσης, αλλοίωση λόγω απώλειας δεδομένων σημειώθηκε από ερευνητές κατά τη διάρκεια της περίληψης βιογραφικών σημειωμάτων από τη βιβλιοθήκη Gensim. Τέτοιες αλλοιώσεις θα μπορούσαν να επηρεάσουν τη σωστή αναγνώριση των κατάλληλων υποψηφίων με αποτέλεσμα να απορριφθούν. Ειδικά για τους συνεντευξιαζόμενους που η γλώσσα της συνέντευξης δεν είναι η μητρική τους, είναι λογικό να υπάρξει ποσοστό σφάλματος καθώς δεν εκφράζονται σωστά και πλήρως όπως στη μητρική τους. Μία πρόταση είναι η πραγματοποίηση της συνέντευξης στη μητρική γλώσσα και έπειτα μετάφραση υψηλού επιπέδου του περιεχομένου στην αγγλική γλώσσα, καθώς τα περισσότερα εργαλεία/βιβλιοθήκες εξειδικεύονται στην αγγλική. Η βελτίωση των εργαλείων/εφαρμογών εξαγωγής, ανίχνευσης, επεξεργασίας και ανάλυσης είναι αναπόφευκτη για αποδοτικότερα, διαφανή, αξιόπιστα αποτελέσματα, ώστε οι εργοδότες αλλά και οι υποψήφιοι να εμπιστεύονται τις διαδικασίες. Αν και σε αρκετά άρθρα έγινε προεπεξεργασία, φαίνεται ότι δεν είναι αρκετή για να ανιχνεύονται όλα τα μη σχετικά δεδομένα και να αφαιρούνται τα διάφορα είδη θορύβων. Θα μπορούσε η προεπεξεργασία να βελτιωθεί ώστε να υπάρχουν μόνο αντιπροσωπευτικά δεδομένα υψηλής ποιότητας προς χρήση και ανάλυση.

Επίσης, σε ορισμένα άρθρα εξετάστηκαν συγκεκριμένοι κλάδοι εργασίας. Κάθε σύστημα πρέπει να εκπαιδευτεί κατάλληλα ώστε να είναι ευέλικτο και να προσαρμόζεται σε διαφορετικές συνθήκες και ανάγκες πρόσληψης σε ποικίλους κλάδους. Η προσαρμοστικότητα του συστήματος και των μεθόδων που χρησιμοποιούνται σε νέες ανάγκες και δεδομένα, λόγω των ραγδαίων τεχνολογικών εξελίξεων, είναι απαραίτητη καθώς οι απαιτήσεις διαρκώς αλλάζουν. Η δημιουργία πύο ολοκληρωμένων συστημάτων, τα οποία θα συνδυάζουν πολλαπλές πηγές δεδομένων ώστε να πραγματοποιείται μια πολυδιάστατη επεξεργασία και αξιολόγηση, θα ήταν καθοριστική για ουσιαστική ανάπτυξη των διαδικασιών πρόσληψης που θα ανταποκρίνεται στις σύνθετες ανάγκες της αγοράς εργασίας. Μία πρόταση των ερευνητών υπήρξε η εξόρυξη δεδομένων από τα κοινωνικά δίκτυα των υποψηφίων, ώστε οι εργοδότες να εξετάζουν εάν η προσωπικότητα του υποψηφίου ταιριάζει με τις απαιτήσεις της θέσης εργασίας. Η χρήση των LLMs ενδείκνυται και σε αυτή την περίπτωση διότι έχουν εκπαιδευτεί βάσει δεδομένων διαφορετικών κλάδων εργασίας. Αλλά θα πρέπει να συνδυαστούν και με μοντέλα ειδικά προς ορισμένους κλάδους, ειδικότερα για εκείνους τους κλάδους για τους οποίους δεν είναι διαθέσιμα αρκετά σχετικά δεδομένα.

Περαιτέρω, οι μελλοντικές κατευθύνσεις της έρευνας θα μπορούσαν να περιλαμβάνουν την σχεδίαση και ανάπτυξη πολυδιάστατων μοντέλων αξιολόγησης, δηλαδή τον συνδυασμό των ουσιαστικών και των τυπικών προσόντων. Ειδικότερα, μελλοντική κατεύθυνση έρευνας θα μπορούσε να επικεντρωθεί στη δημιουργία μιας ενιαίας μεθοδολογίας αξιολόγησης, η οποία θα συνδυάζει ποσοτικά και ποιοτικά χαρακτηριστικά.

Η πλήρης και αντικειμενική αξιολόγηση των προσεγγίσεων θα μπορούσε να αποτελέσει μία σημαντική κατεύθυνση έρευνας που θα μπορούσε να αναπτυχθεί. Αν και οι περισσότερες έρευνες εστιάζουν στη δημιουργία καινοτόμων μοντέλων, παρατηρείται ότι απουσιάζει η συστηματική και συγκριτική ανάλυση των κριτηρίων και των αποτελεσμάτων. Απαιτείται ενδελεχής αξιολόγηση των συστημάτων με τυποποιημένα και μετρήσιμα κριτήρια ώστε η εφαρμογή τους να γίνει αποδεκτή σε διεθνές επίπεδο και να αποτελεί πολύτιμο εργαλείο στις διαδικασίες πρόσληψης.

Στην έρευνα των Langer και συν. (2020) παρουσιάζονται ενδιαφέρουσες πληροφορίες σχετικά με την επίδραση που έχει η αυτόματη αξιολόγηση των συνεντεύξεων εργασίας στους υποψηφίους. Οι υποψήφιοι, κατά τη διάρκεια των

συνεντεύξεων προσπαθούν, μέσω εντυπώσεων να “κερδίσουν” αυτόν που τους αξιολογεί, αυτό όμως δεν συμβαίνει όταν έχουν απέναντί τους ένα σύστημα. Οι ερευνητές πραγματεύονται το πώς και αν, οι υποψήφιοι διαχειρίζονται αυτόν τον μηχανισμό των εντυπώσεων (impression management) όταν αξιολογούνται αυτόματα και όχι από κάποιον άνθρωπο. Το άρθρο αναφέρει ότι έρευνες έχουν δείξει ότι η συμπεριφορά και η απόδοση των υποψηφίων διαφέρει, προς το καλύτερο, στις διαζώσεις, από τις τηλεδιασκέψεις και τις ασύγχρονες συνεντεύξεις. Η συγκεκριμένη έρευνα είναι ενδιαφέρουσα επειδή δεν είναι σίγουρο αν η αξιολόγηση από τα συστήματα είναι ορθή εφόσον υπάρχουν στοιχεία ότι η πλειοψηφία των υποψηφίων επηρεάζεται από το γεγονός ότι νιώθουν μόνοι τους, οπότε δεν καταβάλλουν προσπάθεια να είναι όπως αρμόζει σε μία συνέντευξη. Πιστεύουν ότι τα συστήματα δεν θα αναγνωρίσουν τις δυνατότητές τους, οπότε δεν επιδεικνύουν σε ικανοποιητικό βαθμό τις δυνατότητές τους. Ενδιαφέρουσα κατεύθυνση είναι η διερεύνηση των αντιδράσεων των υποψηφίων σε περιβάλλοντα που δεν γνωρίζουν εκ των προτέρων εάν αξιολογούνται από σύστημα ή από άνθρωπο. Αυτό θα μπορούσε να εξαλείψει τη μειωμένη προσπάθεια που παρατηρήθηκε από τους υποψηφίους, ώστε να μην είναι προκατειλημμένοι και να μην οδηγούνται σε μειωμένη προσπάθεια την ώρα της συνέντευξης.

Η συγκεκριμένη έρευνα ήταν η πρώτη που ασχολήθηκε με το θέμα της Διαχείρισης Εντυπώσεων (Impression Management), τις αντιδράσεις των υποψηφίων, στη διαδικασία πρόσληψης με τη χρήση της MM και έδειξε ότι πολλοί συμμετέχοντες ανησυχούν πρωτίστως για θέματα ιδιωτικότητας. Επίσης πολλοί πιστεύουν, λόγω ιδιοσυγκρασίας, ότι αφού η εταιρεία δεν επενδύει χρόνο για συνεντεύξεις, δεν τους αφιερώνει χρόνο, συνεπώς δεν τους εκτιμά, οπότε και οι προσδοκίες τους για δίκαιη μεταχείριση μειώνονται. Όσον αφορά το ζήτημα της ιδιωτικότητας, πρέπει να διασφαλίζεται με κάθε δυνατό τρόπο μέσω της εξασφάλισης προστασίας διαρροής δεδομένων και με την εφαρμογή αυστηρών πολιτικών απορρήτου και χωρίς να παραβλέπεται η ασφάλεια όπου εμπλέκονται πλατφόρμες και ιστοσελίδες. Επίσης, η ενίσχυση της ασφάλειας των συστημάτων (αξιολόγησης) με ισχυρά αντιαρκα και η λήψη κατάλληλων μέτρων για την πρόληψη παραβίασης δεδομένων είναι κάτι το απαραίτητο. Το αίσθημα υποτίμησης που ενδέχεται να αισθάνονται οι υποψήφιοι από τους εργοδότες, μπορεί να αντιμετωπιστεί με τον σωστό τρόπο επικοινωνίας από τους εργοδότες, αφουγκραζόμενοι τις ανησυχίες των υποψηφίων, με την αναλυτική

ενημέρωση του τρόπου λειτουργίας των συστημάτων και των μέτρων που λαμβάνουν για την προστασία της ιδιωτικότητάς τους καθώς και με ανατροφοδότηση μετά την συνέντευξη.

6.3 Προσωπική Εμπειρία

Κατά τη συγγραφή της παρούσας διπλωματικής εργασίας, απέκτησα πολύτιμες γνώσεις και δεξιότητες, τόσο σε θεωρητικό όσο και σε πρακτικό επίπεδο, ιδιαίτερα στον τομέα της Μηχανικής Μάθησης και της μεθοδολογίας της έρευνας. Η ικανότητά μου στη διαχείριση μεγάλου όγκου πληροφοριών βελτιώθηκε σημαντικά, μέσα από την αναζήτηση, την ανάλυση και τη σύγκριση βιβλιογραφικών πηγών. Η διαδικασία αυτή απαιτεί αφοσίωση, προγραμματισμό, διαχείριση του χρόνου, ενέργεια και αυτοπειθαρχία, ειδικά λαμβάνοντας υπόψη ότι αφορά έναν εργαζόμενο. Η διαχείριση της συγκριτικής αξιολόγησης και της κατηγοριοποίησης υπήρξε κάτι με το οποίο ήρθα αντιμέτωπη πρώτη φορά και αποτέλεσε το πιο απαιτητικό μέρος αυτής της προσπάθειας. Το θέμα της διπλωματικής εργασίας, αν και δεν μου ήταν γνώριμο, μου κέντρισε αμέσως το ενδιαφέρον λόγω της καινοτομίας του. Το εγχείρημα αυτό υπήρξε τεράστια πρόκληση καθώς η ακαδημαϊκή μου προέλευση είναι διαφορετικού κλάδου από αυτόν της Πληροφορικής. Ωστόσο, η πρόκληση αυτή με οδήγησε στο να αναπτύξω νέες δεξιότητες και να διευρύνω τους γνωστικούς μου ορίζοντες. Εκτός από τη συσσώρευση γνώσεων που αποκόμισα, ενισχύθηκε η κριτική μου σκέψη και αναπτύχθηκε η δεξιότητα της αναλυτικής προσέγγισης με τη δημιουργική σκέψη. Η αλλαγή τρόπου σκέψης και η κατανόηση πραγμάτων μέσα από νέο πρίσμα είναι κάτι που αναπόφευκτα αντιμετωπίζει κάθε σπουδαστής. Η εμπειρία αυτή είναι μοναδική και διαμορφώνει τον άνθρωπο, όχι μόνο σε επιστημονικό και γνωστικό επίπεδο, αλλά και σε προσωπικό. Η δια βίου μάθηση αποτελεί θεμέλιο λίθο για την προσωπική ανάπτυξη και η ανάγκη για διαρκή επιμόρφωση είναι απαραίτητη, ειδικά ζώντας σε έναν κόσμο που συνεχώς εξελίσσεται. Ο κάθε ένας από εμάς, έχει άπειρες δυνατότητες, φτάνει να τολμήσει να τις εξερευνήσει.

7. ΒΙΒΛΙΟΓΡΑΦΙΚΕΣ ΑΝΑΦΟΡΕΣ

1. Agrawal, A., George, R. A., Ravi, S., Kamath, S., & Kumar, A. (2020). *Leveraging Multimodal Behavioral Analytics for Automated Job Interview Performance Assessment and Feedback* (No. arXiv:2006.07909). arXiv. <http://arxiv.org/abs/2006.07909>. Association for Computational Linguistics.
2. Amazon Mechanical Turk Workers. [Amazon Mechanical Turk](#). 8-1-2025.
3. Anju Raj, E. & Farzana, T., (2023). *Interview Based Human Behavior Analysis with Machine Learning*. International Journal of Advances in Engineering and Management (IJAEM). Sarath Research Centre. V.5. Issue 6. 41-46.
4. Bajwa, J., Munir, U., Nori, A., Williams, B. (2021). *Artificial intelligence in healthcare: transforming the practice of medicine*. Future Healthc J. Jul;8(2):e188-e194. doi: 10.7861/fhj.2021-0095.
5. Baltrušaitis, T., Robinson, P., & Morency, L. (2016). *OpenFace: An open source facial behavior analysis toolkit*. 2016 IEEE Winter Conference on Applications of Computer Vision (WACV), 1-10. 10-1-2025.
6. Booth, B. M., Hickman, L., Subburaj, S. K., Tay, L., Woo, S. E., & D'Mello, S. K. (2021). *Bias and Fairness in Multimodal Machine Learning: A Case Study of Automated Video Interviews*. Proceedings of the 2021 International Conference on Multimodal Interaction, 268–277. <https://doi.org/10.1145/3462244.3479897>. International Commission on Mathematical Instruction (ACM).
7. Britannica. (2025). [Likert scale | Social Science Surveys & Applications | Britannica](#). 10-1-2025.
8. Brown, Sara. (2021). [Machine learning explained | MIT Sloan](#). 18-10-2024.
9. Brownlee1, Jason. (2020). [10 Clustering Algorithms With Python - MachineLearningMastery.com](#). 8-1-2025.
10. Brownlee2, Jason. (2020). [6 Dimensionality Reduction Algorithms With Python - MachineLearningMastery.com](#). 8-1-2025.

11. Chavare, D. S., & Patil, M. A. B. (2023). *Resume Parsing using Natural Language Processing*. Grenze International Journal of Engineering and Technology. (GIJET). Grenze Scientific Society. January Issue.
12. Chen, L., Feng, G., Martin-Raugh, M., Leong, C. W., Kitchen, C., Yoon, S.-Y., Lehman, B., Kell, H., & Lee, C. M. (2016). *Automatic Scoring of Monologue Video Interviews Using Multimodal Cues*. Interspeech 2016, International Speech Communication Association. (ISCA) 32–36. <https://doi.org/10.21437/Interspeech.2016-1453>.
13. Chen, L., Zhao, R., Leong, C. W., Lehman, B., & Feng, G. (2017). *Automated Video Interview Judgment on a Large-Sized Corpus Collected Online*. Seventh International Conference on Affective Computing and Intelligent Interaction (ACII). IEEE.
14. Cranfield University. (2024). [Snowballing and grey literature - Conducting your literature review - Cranfield Libraries at Cranfield University](#). 7-12-2024.
15. Decision Tree Algorithm. [What is a Decision Tree? | IBM](#). 8-1-2025.
16. Delgutte, Bertrand. Greenberg, Julie. (1999). The Discrete Fourier Transform. Biomedical Signal and Image Processing. HST582J/6.555J/16.456J. 2005. [Chapter 4 - THE DISCRETE FOURIER TRANSFORM](#). 8-1-2025.
17. Dharani. (2025). *A Step-By-Step Complete Guide to Principal Component Analysis | PCA for Beginners*. [Step-By-Step Guide to Principal Component Analysis With Example](#). Turing. 22-1-2025.
18. Dhumne, Shruti. (2023). *Elastic Net Regression Detailed Guide!*. [Elastic Net Regression detailed guide ! | by Shruti Dhumne | Medium](#). Medium. 8-1-2025.
19. Educational Testing Service - ETS. (2024). [ETS](#). 21-12-2024.
20. Farzana Khan, Hamdan Patel, Arshad Shaikh, Fawzah Sayed, & Abdul Rehman Soorya. (2023). *Resume Parser and Summarizer*. International Journal of Advanced Research in Science, Communication and Technology, (IJARSCT). Lambert Publications. 442–447. <https://doi.org/10.48175/IJARSCT-9064>.

21. gazeDirectionGlobal. Visage Technologies. (2025). [VisageTracker-Configuration-Manual.pdf](#). 10-1-2025.
22. Giannakopoulos T. (2015). *pyAudioAnalysis: An Open-Source Python Library for Audio Signal Analysis*. PLoS One. 11;10(12):e0144610. doi: 10.1371/journal.pone.0144610. [pyAudioAnalysis: An Open-Source Python Library for Audio Signal Analysis - PMC](#). 8-1-2025.
23. GitHub. (2024). [GitHub · Build and ship software on a single, collaborative platform · GitHub](#). 18-12-2024.
24. Google Cloud Speech-to-Text API. [Speech-to-Text documentation | Cloud Speech-to-Text Documentation | Google Cloud](#). 8-1-2025.
25. Hemamou, L., Felhi, G., Martin, J.-C., & Clavel, C. (2019). *Slices of Attention in Asynchronous Video Job Interviews*. 2019 8th International Conference on Affective Computing and Intelligent Interaction (ACII), 1–7. <https://doi.org/10.1109/ACII.2019.8925439>. IEEE.
26. Hickman, L., Saef, R., Ng, V., Woo, S. E., Tay, L., & Bosch, N. (2024). *Developing and evaluating language based machine learning algorithms for inferring applicant personality in video interviews*. Human Resource Management Journal, 34(2), 255–274. <https://doi.org/10.1111/1748-8583.12356>. WILEY.
27. Hoeting, Jennifer. (2020). *Caret package in R*. [Caret package in R](#). 10-1-2025.
28. Jagan Mohan Reddy, D., Regella, S., & Seelam, S. R. (2020). *Recruitment Prediction using Machine Learning*. 2020 5th International Conference on Computing, Communication and Security (ICCCS), IEEE. 1–4. <https://doi.org/10.1109/ICCCS49678.2020.9276955>.
29. Jang, B., Kim M., Harerimana G., and Kim J. W., "Q-Learning Algorithms: A Comprehensive Classification and Applications," in IEEE Access, vol. 7, pp. 133653-133667, 2019, doi: 10.1109/ACCESS.2019.2941229.
30. Kaggle. (2024). [Kaggle: Your Machine Learning and Data Science Community](#) 18-12-2024.
31. Kalra, Rakshit. (2024). [Understanding Model-Based Reinforcement Learning | by Rakshit Kalra | Medium](#). 9-1-225.

32. Kavlakoglu Eda & Winland Vanna. (2024). *What is k-means clustering?* [What is k-means clustering? | IBM](#). IBM. 22-1-2025.
33. Kinge Bhushan, Mandhare Shrinivas, Pranali Chavan, & Chaware, S., M.. (2022). *Resume Screening using Machine Learning and NLP: A proposed system*. International Journal of Scientific Research in Computer Science, Engineering and Information Technology, <https://www.google.com/url?q=https://technoscienceacademy.com/&sa=D&source=docs&ust=1734346111990160&usq=AOvVaw2ODOOBz2TZ4gfQzME m0H8>, 253–258. <https://doi.org/10.32628/CSEIT228240>. (IJSRCSE).
34. Kumar, Satyam. (2024). Updated by Urwin, Matthew (2024). [8 Anomaly Detection Algorithms to Know | Built In](#). 8-1-2025.
35. Langer, M., König, C. J., & Hemsing, V. (2020). *Is anybody listening? The impact of automatically evaluated job interviews on impression management and applicant reactions*. Journal of Managerial Psychology, (JMP), 35(4), 271–284. <https://doi.org/10.1108/JMP-03-2019-0156>. Emerald Publishing.
36. Least Absolute Shrinkage and Selection Operator (LASSO). [What is lasso regression? | IBM](#). 8-1-2025.
37. Lim, Annabelle. (2023). *Big Five Personality Traits: The 5-Factor Model of Personality*. [Big 5 Personality Traits: The 5-Factor Model of Personality](#). Simply Psychology. 10-1-2025.
38. Linear Regression. [Linear regression in R - IBM Developer](#). 8-1-2025.
39. LIWC - Linguistic Inquiry Word Count. (2025). [Welcome to LIWC-22](#). 10-1-2025.
40. Logistic Regression. [What Is Logistic Regression? | IBM](#). 8-1-2025.
41. Maklin, Cory. (2022). *Isolation Forest*. [Isolation Forest. Isolation Forest is an unsupervised... | by Cory Maklin | Medium](#). Medium. 22-1-2025.
42. Manifold Learning. [2.2. Manifold learning — scikit-learn 1.5.2 documentation](#). 8-1-2025.

43. Michelangelo. (2022). AlphaZero for dummies!. https://medium.com/@_michelangelo_/alphazero-for-dummies-5bcc713fc9c6. 22-1-2025.
44. Murel, Jacob & Kavlakoglu, Eda. (2024). *What is bag of words?*. [What is bag of words? | IBM](#). 10-1-2025.
45. Naik, R. S., & Dhotre, S. R. (2022). *Resume Recommendation Using Machine Learning*. 10(7). International Journal of Creative Research Thoughts. (IJCRT). Blue Eyes Intelligence Engineering and Sciences Publication (BEIESP).
46. Naive Bayes. (2025). [What Are Naïve Bayes Classifiers? | IBM](#). 9-1-2025.
47. Natural Language Processing. [What Is NLP \(Natural Language Processing\)? | IBM](#). 8-1-2025.
48. Naveed, Z., Nisar, B., Saifullah, D. M., & Baig, J. I. (2024). *Resume Ranking Using Natural Language Processing*. Journal of Computers and Intelligent Systems. (JCIS). Department of Computer Science, The Islamia University of Bahawalpur, Pakistan.
49. Named Entity Recognition (NER). [What Is Named Entity Recognition \(NER\)? | Definition from TechTarget](#). 8-1-2025.
50. Networkx. (2024). [NetworkX — NetworkX documentation](#) 18-12-2024.
51. Nickerson, Charlotte. (2024). *Impression Management: Erving Goffman Theory*. [Impression Management: Erving Goffman Theory](#). 10-1-2025.
52. Nilsson, Nils, J., (1998) *Introduction to Machine Learning an Early Draft of a Proposed Textbook*. Stanford University.
53. Natural Language Toolkit. [NLTK](#). 17-12-2024.
54. OpenCV. [OpenCV Face Recognition powered by Seventh Sense - OpenCV](#). 8-1-2025.
55. OPENSmile. [openSMILE: Audio Fingerprinting and Feature Extraction](#). 10-1-2025.
56. [Papers in 100 Lines of Code](#). (2024). [Value-Based vs Policy-Based Reinforcement Learning | by Papers in 100 Lines of Code | Nov. 2024 | Medium](#). 9-1-2025.
57. Polymer. (2025). [What is Hybrid Machine Learning?](#). 9-1-2025.

58. PRAAT. [Praat - Download](#). 8-1-2025.
59. ProfessionAid. (2024). [LinkedIn Profile: Γιατί είναι απαραίτητο](#), 14-11-2024.
60. Purovit, N., (2023). *LeNet: The Deep Learning Model That Revolutionized Image Recognition*. [LeNet: The Deep Learning Model That Revolutionized Image Recognition | by Neha Purohit | AI monks.io | Medium](#). 8-1-2024.
61. Python Jupyter Notebook. (2024). <https://jupyter.org/>. 18-12-2024.
62. Pyvis. (2024). [Interactive network visualizations — pyvis 0.1.3.1 documentation](#). 18-12-2024.
63. Rahman, Jeslur. (2024). *Simply understanding Artificial Neural Networks (ANN): Deep Learning*. [Simply understanding Artificial Neural Networks \(ANN\): Deep Learning | by Jeslur Rahman | Medium](#). 10-1-2025.
64. Random Forest algorithm. [What Is Random Forest? | IBM](#). 8-1-2025.
65. Rasipuram, S., Rao, S. B. P., & Jayagopi, D. B. (2016). *Automatic prediction of fluency in interface-based interviews*. 2016 Annual India Conference (INDICON), 1–6. <https://doi.org/10.1109/INDICON.2016.7838991>. IEEE.
66. RecordRTC. [RecordRTC: Home](#). 10-1-2025.
67. Ridge Regression. [What Is Ridge Regression? | IBM](#). 8-1-2024.
68. Roy, P. K., Chowdhary, S. S., & Bhatia, R. (2020). *A Machine Learning approach for automation of Resume Recommendation system*. Elsevier, 167, 2318–2327. <https://doi.org/10.1016/j.procs.2020.03.284>.
69. Sachinsoni. (2023). *Clustering Like a Pro: A Beginner's Guide to DBSCAN*. [Clustering Like a Pro: A Beginner's Guide to DBSCAN | by Sachinsoni | Medium](#). 22-1-2025.
70. Sawla, H., Kothari, V., Joshi, K., & Kelkar, K. (2018). *Computational Inference of Candidates' Oratory Performance in Employment Interviews Based on Candidates' Vocal Analysis*. 2018 Second International Conference on Intelligent Computing and Control Systems (ICICCS), 164–169. <https://doi.org/10.1109/ICCONS.2018.8662991>. IEEE.
71. Sentiment Analysis. (2025). [What Is Sentiment Analysis? | IBM](#). 9-1-2025.
72. Shai Shalev-Shwartz & Shai Ben-David. (2014). *Understanding Machine Learning: From Theory to Algorithms*. Cambridge University Press. <http://www.cs.huji.ac.il/~shais/UnderstandingMachineLearning>.

73. Simply Great Resumes. (2023). [How to Add Github to LinkedIn Quickly and Easily](#), 14-11-2024.
74. Siva, Gayathri. (2021). [BERT — Bidirectional Encoder Representations from Transformer | by Gayathri siva | Medium](#). 9-1-2025.
75. Sofeikov, Konstantin. (2023). *REINFORCE algorithm — Reinforcement Learning from scratch in PyTorch*. <https://medium.com/@sofeikov/reinforce-algorithm-reinforcement-learning-from-scratch-in-pytorch-41fccafa107>. Medium. 22-1-2025.
76. Sowjanya, Y., Keerthana, M., Suneeksha, P., Harsha, D., (2023). *Smart Resume Analyzer*. International Journal of Research in Engineering and Science (IJRES). Seventh Sense Research Group (SSRG). V.II., Issue 3, 409-418.
77. SpeechRater ETS. [SpeechRater Service: Advanced Spoken-response Scoring Application](#). 10-1-2025.
78. Suhas, H. E., Manjunath, A., E., (2020). *Differential Hiring using a Combination of NER and Word Embedding*. International Journal of Recent Technology and Engineering (IJRTE), Blue Eyes Intelligence Engineering and Sciences Publication (BEIESP). 9(1), 1344–1349. <https://doi.org/10.35940/ijrte.A2400.059120>.
79. Support Vector Machine (SVM). [What Is Support Vector Machine? | IBM](#). 8-1-2025.
80. Świechowski, M., Godlewski, K., Sawicki, B. *et al.* (2023). *Monte Carlo Tree Search: a review of recent modifications and applications*. *Artif Intell Rev* **56**, 2497–2562. [Monte Carlo Tree Search: a review of recent modifications and applications | Artificial Intelligence Review](#). Springer. 22-1-2025.
81. Tejaswini, K., Umadevi, V., Kadiwal, S. M., & Revanna, S. (2022). *Design and development of machine learning based resume ranking system*. Global Transitions Proceedings, Elsevier, 3(2), 371–375. <https://doi.org/10.1016/j.gltp.2021.10.002>
82. UBI AI. (2024). [Ubi ai](#). 18-12-2024.

83. Vinciarelli Alessandro. (2017). *Introduction: Social Signal Processing*. In: Burgoon JK, Magnenat-Thalmann N, Pantic M, Vinciarelli A, eds. *Social Signal Processing*. Cambridge University Press. [Introduction: Social Signal Processing \(Chapter 1\) - Social Signal Processing](#). 10-1-2025.
84. VoiceBase ASR. [How VoiceBase Automatic Speech Recognition \(ASR\) Works - VoiceBase](#). 10-1-2025.
85. Volk, A. A., Epps, R. W., Yonemoto, D. T., Masters, B. S., Castellano, F. N., Reyes, K. G., & Abolhasani, M. (2023). *AlphaFlow: Autonomous discovery and optimization of multi-step chemistry using a self-driven fluidic lab guided by reinforcement learning*. *Nature (Portfolio)*, 14(1), 1403. <https://doi.org/10.1038/s41467-023-37139-y>.
86. Vrankulj, Adam. (2013). *Emotient and iMotions partner for integrated facial expression recognition, bio sensor and eye-tracking solution*. [Emotient and iMotions partner for integrated facial expression recognition, bio sensor and eye-tracking solution | Biometric Update](#). 10-1-2025.
87. Wu Natalie. (2024). *Interactive AI Interviews are going to become a new norm, expert says: It's 'here, it's real, it's incorporated'*. [Interactive AI Interviews are going to become a new norm, expert says](#). CNBC. 16-12-2024.
88. Yadav, Amit. (2024). *An Introduction to Locally Linear Embedding (LLE)*. [An Introduction to Locally Linear Embedding \(LLE\) | by Amit Yadav | Biased-Algorithms | Medium](#). Medium. 22-1-2025.
89. ΑΣΕΠ1. Ανώτατο Συμβούλιο Επιλογής Προσωπικού. (2024). [Ο θεσμός του ΑΣΕΠ και η συμβολή του στη Δημόσια Διοίκηση | ΑΣΕΠ Ενημερωτική Πύλη](#), 18-10-2024.
90. ΑΣΕΠ2, Ανώτατο Συμβούλιο Επιλογής Προσωπικού. (2024). [Πώς μοριοδοτούνται τα προσόντα | ΑΣΕΠ Ενημερωτική Πύλη](#), 18-10-2024.
91. Ν.3614/Β΄/2022. *Διαδικασία καθορισμού του περιεχομένου των προκηρύξεων επιλογής προσωπικού με βάση προκαθορισμένα και αντικειμενικά κριτήρια*. (2022).

- 92.Π.Δ. 318/1992 (Α' 161). Προεδρικό Διάταγμα καθορισμού ουσιαστικών προσόντων. (1992). <http://elib.aade.gr/elib/view?d=/gr/pd/1992/318/art/7>. 17-12-2024.
- 93.ΑΔΑ: 6ΜΚΠ46ΜΤΛΚ-Η5Τ. Προκήρυξη πλήρωσης θέσεων ευθύνης. (2022). <https://ypergasias.gov.gr/wp-content/uploads/2022/07/6%CE%9C%CE%9A%CE%A046%CE%9C%CE%A4%CE%9B%CE%9A-%CE%975%CE%A4-%CE%A0%CE%A1%CE%9F%CE%9A%CE%97%CE%A1%CE%A5%CE%9E%CE%97-%CE%93%CE%94-%CE%94%CE%A5%CE%A0%CE%91.pdf>, 19-10-2024.
- 94.ΑΔΑ: 71Ψ6Χ – Μ5Λ. Τροποποιητική εγκύκλιος για τα ουσιαστικά προσόντα. (2014). [ΑΔΑ: 71Ψ6Χ – Μ5Λ](#). 17-12-2024.