



ΠΑΝΕΠΙΣΤΗΜΙΟ ΠΑΤΡΩΝ
ΤΜΗΜΑ ΜΗΧΑΝΟΛΟΓΩΝ ΚΑΙ ΑΕΡΟΝΑΥΠΗΓΩΝ ΜΗΧΑΝΙΚΩΝ
ΤΟΜΕΑΣ ΔΙΟΙΚΗΣΗΣ ΚΑΙ ΟΡΓΑΝΩΣΗΣ
ΣΠΟΥΔΑΣΤΗΡΙΟ ΔΙΟΙΚΗΣΗΣ ΕΡΓΟΣΤΑΣΙΩΝ ΚΑΙ
ΣΥΣΤΗΜΑΤΩΝ ΠΑΡΑΓΩΓΗΣ

ΣΠΟΥΔΑΣΤΙΚΗ ΕΡΓΑΣΙΑ

ΣΥΣΤΗΜΑΤΑ ΣΥΣΤΑΣΕΩΝ ΚΑΙ ΜΗΧΑΝΙΚΗ ΜΑΘΗΣΗ –
ΑΝΑΣΚΟΠΗΣΗ ΒΙΒΛΙΟΓΡΑΦΙΑΣ

ΤΣΟΥΝΗΣ ΓΕΩΡΓΙΟΣ

AM1054464

Καρακαπλίδης Νικόλαος, Καθηγητής

ΠΑΤΡΑ, ΝΟΕΜΒΡΙΟΣ 2023

Πανεπιστήμιο Πατρών, Τμήμα Μηχανολόγων και Αεροναυπηγών Μηχανικών
– Τομέας Διοίκησης και Οργάνωσης Τσουνής Γεώργιος
© 2023 – Με την επιφύλαξη παντός δικαιώματος

Η έγκριση της διπλωματικής εργασίας δεν υποδηλοί την αποδοχή των γνώμών του συγγραφέα.
Κατά τη συγγραφή τηρήθηκαν οι αρχές της ακαδημαϊκής δεοντολογίας.

ΠΕΡΙΛΗΨΗ

ΣΥΣΤΗΜΑΤΑ ΣΥΣΤΑΣΕΩΝ ΚΑΙ ΜΗΧΑΝΙΚΗ ΜΑΘΗΣΗ – ΑΝΑΣΚΟΠΗΣΗ

ΒΙΒΛΙΟΓΡΑΦΙΑΣ

ΤΣΟΥΝΗΣ ΓΕΩΡΓΙΟΣ

Στην εποχή που οι υπηρεσίες γίνονται ολοένα και πιο ανταγωνιστικές, η παροχή συστάσεων σε χρήστες αποτελεί μια κομβική διαδικασία. Ενώ τα συστήματα συστάσεων αναπτύσσονται εδώ και πολλές δεκαετίες, η σημερινή πραγματικότητα που χαρακτηρίζεται από την πληθώρα πληροφοριών και τις προκλήσεις αναφορικά με την ταχύτητα και την ακρίβεια δημιουργεί την αναγκαιότητα για την χρήση πιο εξελιγμένων αλγορίθμων. Στο πρώτο κεφάλαιο της παρούσας εργασίας γίνεται μια παρουσίαση των πιο βασικών κατηγοριών αλγορίθμων συστάσεων. Στην παρουσίαση αυτή, δομικές μονάδες είναι ο χρήστης και τα αντικείμενα, τα οποία διαφέρουν από σύστημα σε σύστημα, ενώ οι ίδιοι αλγόριθμοι μπορεί να χρησιμοποιούνται για να μεταχειρίζονται διαφορετικό τύπο αντικειμένων. Οι δύο κύριες οικογένειες αλγορίθμων που χρησιμοποιούνται στα συστήματα συστάσεων είναι οι αλγόριθμοι συλλογικού φιλτραρίσματος, όπου οι συστάσεις παράγονται με βάση το τι προτιμούν οι χρήστες που θεωρούνται παρόμοιοι με τον χρήστη για τον οποίον παράγονται συστάσεις, και το φιλτράρισμα βάση του περιεχόμενου των αντικειμένων, όπου εκείνα αναλύονται με κάποια μέθοδο, και ύστερα προτείνονται στον χρήστη τα εκτιμώμενα ως πιο χρήσιμα για εκείνον. Υπάρχουν και συμπληρωματικές μέθοδοι που χρησιμοποιούνται σε χαμηλότερης ακρίβειας συστήματα, ή συμπληρωματικά για να καλύψουν αδυναμίες άλλων αλγορίθμων, καθώς και υβριδικές υλοποιήσεις οι οποίες επίσης σκιαγραφούνται σε αυτήν την εργασία.

Στο δεύτερο κεφάλαιο γίνεται αρχικά μια εισαγωγή στην έννοια της μηχανικής μάθησης, και έπειτα παρουσιάζονται κάποιες πιο συγκεκριμένες περιπτώσεις όπου εκείνη χρησιμοποιείται μέσα σε συστήματα συστάσεων.

Λέξεις κλειδιά

Συστήματα Συστάσεων, Συλλογικό Φιλτράρισμα, Φιλτράρισμα Βάση Περιεχομένου, Νευρωνικά Δίκτυα

ABSTRACT

RECOMMENDER SYSTEMS AND MACHINE LEARNING – A LITERATURE

REVIEW

TSOUNIS GEORGIOS

During modern times, when services keep getting increasingly competitive in terms of performance, providing recommendations to users has come to be an important procedure. While recommender systems have been developing for decades, the overwhelming amount of information present today, and the increased demands create the need for more sophisticated algorithms.

The first chapter of this piece of work consists of a presentation of the most basic categories of recommendation algorithms. In this presentation, the basic units are the users and the items with which they interact, which of course differ between systems. However, the same types of algorithms are often used to recommend completely diverse kinds of items. The main families of algorithms are the ones that use collaborative filtering, where recommendations are generated based on the preferences of the users similar to the user in question, and content-based filtering, where the content of the items is analyzed and then compared to what is described as the user's preferences. There also exist other algorithms suitable for lower accuracy, fast-paced systems, or that exist complementary to the main algorithms, and many hybrid implementations. All of these are also presented.

The second chapter is dedicated to describing the concept of machine learning, and afterwards presenting some cases of implementation of machine learning technology and neural networks in recommender systems.

Λέξεις κλειδιά

Recommender Systems, Collaborative Filtering, Content Based Filtering, Neural Networks

ΚΑΤΑΛΟΓΟΣ ΠΙΝΑΚΩΝ

Πίνακας 1. Αξιολογήσεις για το παράδειγμα 1.....σελ. 28	28
Πίνακας 2. Ομοιότητες Χρηστών για το παράδειγμα 1..... σελ. 29	29
Πίνακας 3. Αξιολογήσεις για το παράδειγμα 2σελ. 31	31
Πίνακας 4. Ομοιότητες Αντικειμένων για το Παράδειγμα 2.....σελ. 32	32
Πίνακας 5. Χαρακτηριστικά Αντικειμένων για το Παράδειγμα 3.....σελ. 47	47

ΚΑΤΑΛΟΓΟΣ ΣΧΗΜΑΤΩΝ ΚΑΙ ΕΙΚΟΝΩΝ

Εικόνα 1. Σχηματική απεικόνιση ενός τυπικού συστήματος που λειτουργεί με Συλλογικό Φιλτράρισμα.....σελ. 23	
Εικόνα 2. Σχηματική απεικόνιση ενός Bayesian Network. Φαίνεται πως το ενδεχόμενο L(Win Lottery) είναι ανεξάρτητο από τα άλλα δύο, ενώ τα ενδεχόμενα R(Rain),W(Wet Ground) είναι εξαρτημένα.....σελ. 36	
Εικόνα 3. Διάγραμμα Ροής Δεδομένων για ένα τυπικό CBF Σύστημασελ. 41	
Εικόνα 4. Παράδειγμα Μικτών Προτάσεων στην Πλατφόρμα "Youtube".....σελ. 58	

ΣΥΜΒΟΛΙΣΜΟΙ

$r_{c,s}$	αξιολόγηση ενός αντικείμενου s από τον χρήστη c
C	το σύνολο που περιέχει όλους τους χρήστες c
S	το σύνολο που περιέχει όλα τα αντικείμενα s
C_{xy}	το σύνολο των χρηστών που έχουν αξιολογήσει τόσο το αντικείμενο x όσο και το y
S_{xy}	το σύνολο των αντικειμένων που έχουν αξιολογήσει τόσο ο χρήστης x όσο και ο y $\text{sim}(x,y)$ η ομοιότητα των χρηστών x,y
k	παράγοντας κανονικοποίησης που χρησιμοποιείται για τον υπολογισμό της χρησιμότητας ενός αντικειμένου
$x_{d,i}$	η συνάρτηση βάρους για τον όρο i στο έγγραφο d
$f_{d,i}$	Η συχνότητα με την οποία εμφανίζεται ο όρος i στο έγγραφο d
L_c	η λίστα προτάσεων για τον χρήστη c

ΣΥΝΤΜΗΣΕΙΣ

CF	Collaborative Filtering
SVD	Singular Value Decomposition
CBF	Content-Based Filtering
TF-IDF	Term Frequency – Inverse Document Frequency
ML	Machine Learning
DML	Deep Machine Learning
AE	Autoencoder
DBM	Deep Boltzmann Machine
DBN	Deep Believe Network
CNN	Convolutional Neural Network
RNN	Recurrent Neural Network
CRNN	Convolutional Recurrent Neural Network
STFT	Short Time Fourier Transform
ACR	Article Content Representation
NAR	Next-Article Recommendation

ΠΕΡΙΕΧΟΜΕΝΑ

ΠΕΡΙΛΗΨΗ.....	IV
ABSTRACT	VII
ΚΑΤΑΛΟΓΟΣ ΠΙΝΑΚΩΝ	X
ΚΑΤΑΛΟΓΟΣ ΣΧΗΜΑΤΩΝ ΚΑΙ ΕΙΚΟΝΩΝ	XI
ΣΥΜΒΟΛΙΣΜΟΙ.....	XII
ΣΥΝΤΜΗΣΕΙΣ	XIII
ΠΕΡΙΕΧΟΜΕΝΑ.....	XIV
1. ΕΙΣΑΓΩΓΗ.....	2
1.1 ΣΥΣΤΗΜΑΤΑ ΣΥΣΤΑΣΕΩΝ.....	2
1.2 ΜΗΧΑΝΙΚΗ ΜΑΘΗΣΗ	3
2. ΣΥΛΛΟΓΙΚΟ ΦΙΛΤΡΑΡΙΣΜΑ - COLLABORATIVE FILTERING (CF).....	6
2.1 ΕΙΣΑΓΩΓΗ	6
2.2 MEMORY BASED CF	8
2.2.1 User Based.....	8
2.2.2 Item Based.....	12
2.3 MODEL BASED (CLUSTER MODELS, BAYESIAN NETWORKS)	15
2.3.1 Το πρόβλημα της ελαχιστοποίησης των διαστάσεων – Αλγόριθμος SVD	16
2.3.2 Cluster Models	18
2.3.3 Bayesian Networks	19

2.4 ΣΧΟΛΙΑ-ΑΔΥΝΑΜΙΕΣ-ΣΥΜΠΕΡΑΣΜΑΤΑ	20
3. ΦΙΛΤΡΑΡΙΣΜΑ ΒΑΣΗ ΠΕΡΙΕΧΟΜΕΝΟΥ - CONTENT BASED FILTERING (CBF)	23
3.1 ΕΙΣΑΓΩΓΗ	23
3.2 ΑΝΑΛΥΣΗ ΠΕΡΙΕΧΟΜΕΝΟΥ	25
3.2.1 TF-IDF	25
3.2.2 Naïve Bayes Classifier	26
3.3 ΔΗΜΙΟΥΡΓΙΑ ΠΡΟΦΙΛ ΧΡΗΣΤΗ	28
3.4 ΔΗΜΙΟΥΡΓΙΑ ΣΥΣΤΑΣΕΩΝ	28
3.5 ΣΧΟΛΙΑ-ΑΔΥΝΑΜΙΕΣ-ΣΥΜΠΕΡΑΣΜΑΤΑ	31
4. ΣΥΜΠΛΗΡΩΜΑΤΙΚΟΙ ΑΛΓΟΡΙΘΜΟΙ	34
4.1 ΕΙΣΑΓΩΓΗ	34
4.2 STEREOTYPING	34
4.3 GLOBAL RELEVANCE	36
4.4 CO-OCCURRENCE	37
5. ΥΒΡΙΔΙΚΕΣ ΜΕΘΟΔΟΙ	38
5.1 ΕΙΣΑΓΩΓΗ	38
5.2 ΣΥΝΔΥΑΣΤΙΚΑ ΜΟΝΤΕΛΑ – ΣΥΓΚΕΡΑΣΜΟΣ ΠΡΟΤΑΣΕΩΝ	38
5.2.1 Συνδυασμός Προτάσεων με συντελεστές βαρύτητας	39
5.2.2 Συνδυασμός Προτάσεων με Εναλλαγή Αλγορίθμων	40
5.2.3 Παρουσίαση Μικτών Προτάσεων	41
5.3 ΕΝΣΩΜΑΤΩΣΗ ΧΑΡΑΚΤΗΡΙΣΤΙΚΩΝ – ΠΡΑΓΜΑΤΙΚΑ ΥΒΡΙΔΙΑ	42

5.3.1	Ενσωμάτωση CF χαρακτηριστικών σε CBF Υλοποιήσεις	42
5.3.2	Ενσωμάτωση CBF χαρακτηριστικών σε CF Υλοποιήσεις	43
5.3.3	Αλληλουχίες Αλγορίθμων	44
5.4	ΣΧΟΛΙΑ – ΑΔΥΝΑΜΙΕΣ – ΣΥΜΠΕΡΑΣΜΑΤΑ	44
6.	ΣΥΣΤΗΜΑ ΣΥΣΤΑΣΕΩΝ ΜΟΥΣΙΚΗΣ ΜΕ ΤΗΝ ΧΡΗΣΗ ΕΠΑΝΑΛΑΜΒΑΝΟΜΕΝΟΥ ΝΕΥΡΩΝΙΚΟΥ ΔΙΚΤΥΟΥ ΣΥΝΕΛΙΞΗΣ	46
6.1	ΕΠΑΝΑΛΑΜΒΑΝΟΜΕΝΑ ΝΕΥΡΩΝΙΚΑ ΔΙΚΤΥΑ ΣΥΝΕΛΙΞΗΣ	47
6.2	ΠΕΡΙΓΡΑΦΗ ΣΥΣΤΗΜΑΤΟΣ	48
6.3	ΑΞΙΟΛΟΓΗΣΗ ΑΠΟΤΕΛΕΣΜΑΤΩΝ	48
7.	ΣΥΣΤΗΜΑ ΣΥΣΤΑΣΕΩΝ ΕΙΔΗΣΕΩΝ ΒΑΣΙΣΜΕΝΟ ΣΕ ΜΕΜΟΝΩΜΕΝΕΣ ΣΥΝΕΔΡΙΕΣ ΜΕ ΧΡΗΣΗ ΒΑΘΙΩΝ ΝΕΥΡΩΝΙΚΩΝ ΔΙΚΤΥΩΝ	50
7.1	ΣΥΝΑΡΤΗΣΕΙΣ ΠΟΥ ΧΡΗΣΙΜΟΠΟΙΗΘΗΚΑΝ	50
7.1.1	Η συνάρτηση softmax	50
7.1.2	Η λογαριθμική απώλεια διασταυρούμενης εντροπίας (cross-entropy log loss) 51	
7.3	ΠΕΡΙΓΡΑΦΗ ΣΥΣΤΗΜΑΤΟΣ	51
7.3.1	Σύστημα αναπαράστασης άρθρου (ACR)	52
7.3.2	Σύστημα πρόβλεψης επόμενου άρθρου (NAR)	52
7.4	ΑΞΙΟΛΟΓΗΣΗ ΑΠΟΤΕΛΕΣΜΑΤΩΝ	53
8.	ΣΥΝΟΨΗ – ΣΥΜΠΕΡΑΣΜΑΤΑ	55
	ΒΙΒΛΙΟΓΡΑΦΙΑ	60

Στην Παλλού

1. ΕΙΣΑΓΩΓΗ

1.1 ΣΥΣΤΗΜΑΤΑ ΣΥΣΤΑΣΕΩΝ

Τα συστήματα συστάσεων είναι συστήματα τα οποία στις μέρες μας βρίσκονται παντού. Χρησιμοποιούνται κυρίως για διαφημιστικούς σκοπούς, όπου προτείνουν στους χρήστες ενός συστήματος αντικείμενα τα οποία είναι πιο πιθανό να θέλουν να αγοράσουν, αλλά η χρήση τους δεν περιορίζεται εκεί. Αποτελούν ένα σημαντικό πεδίο έρευνας από τα μέσα της δεκαετίας του 1990, όπου δημοσιεύτηκαν τα πρώτα επιστημονικά άρθρα μέχρι και σήμερα, καθώς υπάρχει πληθώρα προβλημάτων τα οποία λύνονται με την χρήση τέτοιων συστημάτων[1]. Τα συστήματα αυτά είναι ιδιαίτερα προηγμένα και παράγουν υψηλής ακρίβειας συστάσεις, σε σημείο που οι χρήστες πολύ συχνά νιώθουν πως παραβιάζονται τα προσωπικά τους δεδομένα[2] [3]

Ένα τυπικό σύστημα συστάσεων περιέχει τους χρήστες του και τα αντικείμενα του, και αποθηκεύει στοιχεία και για τους δύο. Σκοπός του είναι να κρίνει για κάθε χρήστη του ξεχωριστά το ποια αντικείμενα θα του είναι χρήσιμα, να παράγει δηλαδή συστάσεις για εκείνον. Τα αντικείμενα αυτά μπορεί να είναι πολλά διαφορετικά πράγματα, από επιστημονικά άρθρα και βιβλία μέχρι ταινίες[4], σειρές και προϊόντα, και να περιγράφονται ή να αναλύονται με πολλούς διαφορετικούς τρόπους.

Μεγάλη ποικιλία υπάρχει επίσης στον τρόπο που παράγονται οι συστάσεις. Αυτό θα είναι και το αντικείμενο της παρούσας εργασίας, όπου θα παρουσιαστούν οι διαφορετικοί τύποι αλγορίθμων που τις παράγουν. Στην πράξη, σχεδόν ποτέ δεν υλοποιούνται οι αλγόριθμοι στην «καθαρή» τους μορφή – αντ’ αυτού, προτιμώνται υβριδικές υλοποιήσεις. Στην εργασία αυτή υπάρχει ένα κεφάλαιο το οποίο αφιερώνεται σε αυτές.

Οι κύριες κατηγορίες αλγορίθμων είναι οι αλγόριθμοι *Συλλογικού Φιλτραρίσματος* ή αλλιώς *Collaborative Filtering*, όπου οι συστάσεις παράγονται με βάση τις προτιμήσεις των χρηστών με τους οποίους ο χρήστης σχετίζεται, και οι αλγόριθμοι *Φιλτραρίσματος με Βάση το Περιεχόμενο*

των αντικειμένων, ή αλλιώς *Content Based Filtering*, όπου παράγονται με βάση τα χαρακτηριστικά των αντικειμένων, τα οποία εξάγονται με διάφορους τρόπους[5]. Σε κάθε περίπτωση, συνηθίζεται να χρησιμοποιούνται συναρτήσεις ομοιότητας, η οποίες συγκρίνουν αντικείμενα και χρήστες προκειμένου να υπάρχει ένα μέτρο για το ποιοι χρήστες είναι παρόμοιοι μεταξύ τους, ποια αντικείμενα ταιριάζουν στις ανάγκες των χρηστών, κ.α..

1.2 ΜΗΧΑΝΙΚΗ ΜΑΘΗΣΗ

Η Μάθηση αποτελεί την διαδικασία της συσσώρευσης γνώσης. Οι άνθρωποι μαθαίνουν φυσικά, αλλά οι υπολογιστές βασίζονται σε αλγορίθμους, καθότι δεν διαθέτουν τους μαθησιακούς μηχανισμούς του ανθρώπου, και τα ένστικτα του που λειτουργούν σαν φίλτρα μέσα από τα οποία αντιλαμβάνεται τον κόσμο, από την πρώτη στιγμή που αναπνέει.

Η Μηχανική Μάθηση (Για την υπόλοιπη εργασία θα αναφερόμαστε στον όρο και ως ML, ή αλλιώς Machine Learning) αποτελεί έναν κλάδο της τεχνητής νοημοσύνης, ο οποίος αποσκοπεί στην δημιουργία συστημάτων που θα μιμούνται κατά μία έννοια την ανθρώπινη διαδικασία της μάθησης. Τα συστήματα ML εμπεριέχουν συνήθως πολλές παραμέτρους οι οποίες πρέπει να ρυθμιστούν ώστε να επιτελούνται συγκεκριμένες εργασίες, και αυτές οι παράμετροι καθορίζονται από την «εμπειρία» του συστήματος, την χρήση δηλαδή πραγματικών δεδομένων. Υπάρχουν τέσσερις διαφορετικές κατηγοριοποιήσεις των ML συστημάτων βάση της διαδικασίας μάθησης: τα επιβλεπόμενα, τα ημί-επιβλεπόμενα, τα μη επιβλεπόμενα συστήματα, και η ενισχυμένη μάθηση [6].

Τα συστήματα μηχανικής μάθησης υλοποιούνται συνήθως μέσω νευρωνικών δικτύων. Τα δίκτυα αυτά αποτελούνται από νευρώνες που διαβιβάζουν πληροφορίες και εκτελούνε απλές πράξεις, δίνοντας τελικά τιμές σαν αποτελέσματα βάση της εκάστοτε εισόδου. Υπάρχουν πολλές διαφορετικές αρχιτεκτονικές νευρωνικών δικτύων.

Στην περίπτωση των επιβλεπόμενων συστημάτων, στον αλγόριθμο παρέχονται δεδομένα (datasets) τα οποία περιλαμβάνουν κάποιου είδους αντιστοίχιση μεταξύ των δεδομένων και της «σωστής απάντησης», μιας περιγραφής δηλαδή του τι αναπαριστά το κάθε δεδομένο. Ένα τέτοιο σύνολο θα ήταν για παράδειγμα, μια σειρά από συγκεκριμένου μεγέθους εικόνες που απεικονίζουν χειρόγραφα νούμερα, και μια αντιστοίχιση για το ποιο νούμερο συμβολίζει το καθένα. Σκοπός του αλγορίθμου είναι να διαμορφώσει τις παραμέτρους του δικτύου του μέσω αυτών των δεδομένων, και ύστερα να τα εφαρμόσει σε πραγματικά δεδομένα τα οποία θα πρέπει να κατηγοριοποιήσει.

Στην μη επιβλεπόμενη μάθηση, τα δεδομένα που δίνονται στον αλγόριθμο δεν είναι οργανωμένα, δεν έχουν δημιουργηθεί δηλαδή για τον σκοπό της εκπαίδευσης του αλγορίθμου και μόνο, οπότε το σύστημα πρέπει να ανιχνεύσει μοτίβα μέσα στα δεδομένα, και ύστερα να τα κατηγοριοποιήσει σε κλάσεις. Ημί-επιβλεπόμενο μπορεί να χαρακτηριστεί ένα σύστημα όπου ναι μεν τα δεδομένα που του δίνονται είναι οργανωμένα, αλλά κάποιες πληροφορίες λείπουν.

Τέλος, στην περίπτωση που υπάρχει ανάμειξη του χρήστη στην εκπαίδευση του συστήματος, εάν δηλαδή τα αποτελέσματα που δίνει αξιολογούνται από κάποιον, όπως για παράδειγμα ένας αλγόριθμος που εκπαιδεύεται να παίζει ένα παιχνίδι στρατηγικής καθώς το παίζει, οπότε και παραμετροποιείται βάση των παρτίδων που υπήρξαν νικηφόρες, τότε μιλάμε για

την περίπτωση της επαυξημένης μάθησης [6]. Όλα αυτά τα μοντέλα μπορούνε σαφώς να υποστούνε προσαρμογές και τροποποιήσεις.

Η Βαθιά Μηχανική Μάθηση (Για την υπόλοιπη εργασία θα αναφερόμαστε στον όρο ως DML, ή αλλιώς Deep Machine Learning) ανήκει στο ευρύτερο σύνολο της ML, και αποτελεί την «νέα γενιά» των νευρωνικών δικτύων[7] [8], τα οποία εκμεταλλεύονται πολλά επίπεδα νευρώνων (οπότε και στάδια της διαδικασίας) για να παράγουν αξιοσημείωτα αποτελέσματα[9] [10], [11]

Τα συστήματα DML χωρίζονται σε Παραγωγικές και Διαχωριστικές Αρχιτεκτονικές. Στην πρώτη περίπτωση, η διαδικασία είναι μη-επιβλεπόμενη, και αποσκοπεί να βρει συσχετίσεις υψηλής τάξης μέσα στα δεδομένα και μοτίβα μέσα σε αυτά. Τυπικές περιπτώσεις είναι οι Autoencoders(AE), οι Deep Boltzmann Machines (DBMs) , τα Deep Believe Networks (DBNs), κ.α. Οι Διαχωριστικές αρχιτεκτονικές επικεντρώνονται στο να προσδιορίσουν μια συνάρτηση που καθορίζει το όριο μεταξύ των κλάσεων και να μπορεί να διαχωρίσει τα δεδομένα σε αυτές, αντί να «συνθέτει» μοτίβα μέσα σε αυτές. Τυπικές Αρχιτεκτονικές τέτοιου τύπου είναι τα Νευρωνικά Δίκτυα Συνέλιξης(Convolutional Neural Networks, CNN), τα επαναλαμβανόμενα Νευρωνικά Δίκτυα(Recurrent Neural Networks, RNN), κ.α.[7].

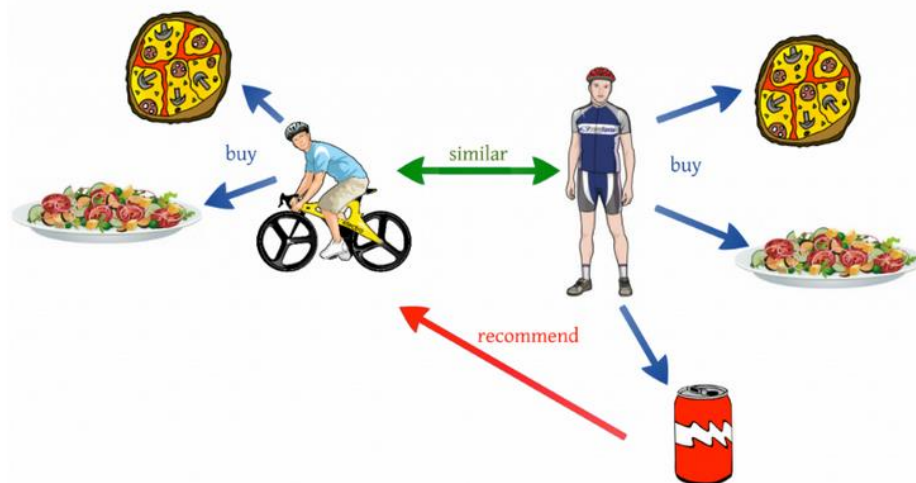
2. ΣΥΛΛΟΓΙΚΟ ΦΙΛΤΡΑΡΙΣΜΑ - COLLABORATIVE FILTERING (CF)

2.1 ΕΙΣΑΓΩΓΗ

Το Συλλογικό Φιλτράρισμα, ή όπως αποκαλείται στην διεθνή ορολογία Collaborative Filtering (CF για λόγους συντομίας), είναι η μια από τις δύο βασικές κατηγορίες που παρουσιάζονται σε αυτήν την εργασία.

Όπως υποδηλώνει το όνομα της μεθόδου, η βασική αρχή της μεθόδου είναι η δημιουργία προτάσεων για κάθε χρήστη με βάση τις προτιμήσεις άλλων χρηστών με τους οποίους εκείνος σχετίζεται με ορισμένα κριτήρια.[12]

Η μέθοδος μελετάται από την δεκαετία του 90, και έχουν υπάρξει πολλές προσαρμογές και διαφοροποιήσεις.



Εικόνα 1. Σχηματική απεικόνιση ενός τυπικού συστήματος που λειτουργεί με Συλλογικό Φιλτράρισμα

Υπάρχουν διαφορετικές προσεγγίσεις για την συγκεκριμένη μέθοδο, οι οποίες παρουσιάζονται παρακάτω. Κατά βάση, οι αλγόριθμοι χωρίζονται σε εκείνους που βασίζονται στην ομοιότητα των χρηστών ή των αντικειμένων (Memory Based), και σε εκείνους που αποσκοπούν στο να σχηματίσουν ένα μοντέλο για κάθε χρήστη μέσω μηχανικής μάθησης (Model Based)[1],[13]

Σε κάθε μία από τις δύο περιπτώσεις, χρησιμοποιούνται οι αξιολογήσεις των χρηστών. Οι αξιολογήσεις [14]αυτές μπορούν να είναι είτε έμμεσες είτε άμεσες[15] Στην δεύτερη περίπτωση, ο χρήστης έχει μπει στην διαδικασία να αξιολογήσει αντικείμενα ενεργά(μέσω π.χ. μιας κλίμακας από το 1 μέχρι το 5, ή ενός μηχανισμού like/dislike), ενώ στην πρώτη ως αξιολόγηση λαμβάνεται το να έχει έρθει σε επαφή με το αντικείμενο(δηλαδή να έχει περάσει χρόνο όπου εκείνο βρίσκεται στην οθόνη του εάν αυτό είναι ένα άρθρο, να το έχει παρακολουθήσει εάν αυτό είναι ένα βίντεο, κ.λπ.).

Σε ορισμένες περιπτώσεις, οποιαδήποτε επαφή με το αντικείμενο μπορεί να θεωρηθεί ως μια θετική αξιολόγηση εκείνου, ή μπορούν άμεσες και έμμεσες κριτικές να χρησιμοποιηθούν και οι δύο. Οποσδήποτε, βέβαια, θα υπάρχει πάντοτε το ζήτημα έλλειψης και αραιότητας των δεδομένων, καθώς κανένας χρήστης δεν θα έχει αλληλοεπιδράσει με όλα τα αντικείμενα [15].

Είναι πολλές φορές χρήσιμο να εισάγονται στο σύστημα εικονικοί χρήστες (filterbots), με σκοπό να βελτιώνονται οι προτάσεις προς τους χρήστες που μοιάζουν οι προτιμήσεις τους με εκείνους.

2.2 MEMORY BASED CF

Σε αυτήν την περίπτωση υλοποίησης του αλγορίθμου, παράγονται προβλέψεις για την εκτιμώμενη χρησιμότητα ενός αντικειμένου για έναν χρήστη, μέσω μιας ευριστικής μεθόδου με βάση τις αξιολογήσεις άλλων χρηστών οι οποίες αντλούνται από μια βάση δεδομένων [15].

Σε αυτό το σημείο εισάγουμε τον συμβολισμό $r_{c,s}$ για να συμβολίσουμε την αξιολόγηση ενός αντικειμένου από έναν χρήστη, όπου

c : ο χρήστης, και C το σύνολο όλων των χρηστών ($c \in C$)

s : το αντικείμενο, και S το σύνολο όλων των αντικειμένων ($s \in S$)

Η αξιολόγηση ή χρησιμότητα αυτή αποτελεί το ζητούμενο μέγεθος από τον αλγόριθμο εφόσον δεν είναι γνωστή.

Θα αναφερόμαστε στην προαναφερθείσα βάση ως έναν πίνακα με στοιχεία $r_{c,s}$, που αντιστοιχούν σε κάθε αξιολόγηση r ενός χρήστη c για ένα αντικείμενο s . Ο συγκεκριμένος πίνακας εμπεριέχει σαφώς κενές θέσεις, εφόσον δεν παρέχουν ποτέ όλοι οι χρήστες αξιολογήσεις για όλα τα αντικείμενα.

2.2.1 User Based

Η πιο συνηθισμένη υλοποίηση αυτής μεθόδου, βασίζεται στις αξιολογήσεις των k πιο όμοιων χρηστών με τον χρήστη i . (Κοντινότεροι γείτονες).

Εισάγουμε την έννοια της ομοιότητας $\text{sim}(x,y)$, μεταξύ δύο χρηστών [1].

Αρχικά, ας προσδιορίσουμε το σύνολο S_{xy} ως το σύνολο όλων των αντικειμένων s για τα οποία έχουν δώσει κάποια αξιολόγηση τόσο ο χρήστης x όσο και ο χρήστης y . Δηλαδή

$$S_{xy} = \{s \in S | r_{x,s} \neq \emptyset, r_{y,s} \neq \emptyset\} \text{ (Εξίσωση 1)}$$

Ορίζουμε λοιπόν, τη συνάρτηση $\text{sim}(x,y)$ (pearson correlation score) ως την ομοιότητα των χρηστών x,y βάση των αξιολογήσεων που έχουν για τα αντικείμενα που ανήκουν στο σύνολο S_{xy} , ως εξής:

$$\text{sim}_s(x,y) = \frac{\sum_{s \in S_{xy}} (r_{x,s} - \bar{r}_x)(r_{y,s} - \bar{r}_y)}{\sqrt{\sum_{s \in S_{xy}} (r_{x,s} - \bar{r}_x)^2 \sum_{s \in S_{xy}} (r_{y,s} - \bar{r}_y)^2}} \text{ (Εξίσωση 2)}$$

Όπου $\bar{r}_x, \bar{r}_y$ ο μέσος όρος των αξιολογήσεων που έχουν δώσει x,y

Οι παράγοντες $\bar{r}_x$ και $\bar{r}_y$ αντιπροσωπεύουν τον μέσο όρο των αξιολογήσεων του κάθε χρήστη [16] και εισέρχονται στη συνάρτηση ώστε να κανονικοποιούνται οι αξιολογήσεις του κάθε χρήστη για το κάθε αντικείμενο s . Αντιλαμβανόμαστε πως αυτό είναι αναγκαίο, επειδή εάν π.χ. ο χρήστης x αξιολογήσει την ταινία A με 6/10 αλλά ο μέσος όρος αξιολογήσεων του είναι 5/10, σημαίνει πως του αρέσει, ενώ ο χρήστης B ο οποίος έχει μέσο όρο αξιολογήσεων 8/10, μάλλον δεν έμεινε πολύ ευχαριστημένος. (όσο πιο κοντά στην μονάδα είναι το αποτέλεσμα της συνάρτησης, τόσο περισσότερο μοιάζουν οι δύο χρήστες)

Εναλλακτικά, μπορούμε να χρησιμοποιήσουμε την έννοια της συνημιτονοειδούς ομοιότητας μεταξύ δύο αντικειμένων. Σε αυτή την περίπτωση, οι χρήστες αντιμετωπίζονται σαν διανύσματα σε έναν m -διάστατο χώρο, όπου $m = |S_{xy}|$.

Η ομοιότητα μεταξύ των δύο χρηστών υπολογίζεται λοιπόν σαν το συνημίτονο της γωνίας που σχηματίζουν τα διανύσματα που αντιστοιχούν σε εκείνους, ως εξής:

$$sim_s(x, y) = \cos(\vec{x}, \vec{y}) = \frac{\vec{x} \cdot \vec{y}}{||\vec{x}||_2 \times ||\vec{y}||_2} = \frac{\sum_{s \in S_{xy}} r_{x,s} r_{y,s}}{\sqrt{\sum_{s \in S_{xy}} r_{x,s}^2} \sqrt{\sum_{s \in S_{xy}} r_{y,s}^2}} \quad (Εξίσωση 3)$$

Αντίστοιχα με τον προηγούμενο ορισμό για την ομοιότητα, όσο πιο κοντά είναι το αποτέλεσμα στην μονάδα, τόσο περισσότερο μοιάζουν οι δύο χρήστες.

Μια συνήθης στρατηγική είναι να υπολογίζονται όλες οι ομοιότητες μεταξύ των χρηστών εκ των προτέρων, και να χρησιμοποιούνται όποτε αυτές χρειάζονται, καθώς οι σχέσεις αυτές μεταξύ των χρηστών δεν μεταβάλλονται δραστικά με την πάροδο του χρόνου[1].

Εφόσον γνωρίζουμε την ομοιότητα του ζητούμενου χρήστη με κάθε άλλο χρήστη για τον οποίον είναι εφικτό εκείνη να υπολογιστεί, εξάγουμε την ζητούμενη αξιολόγηση για το αντικείμενο s από τις αξιολογήσεις των N **κοντινότερων γειτόνων** με τον χρήστη c . Το σύνολο αυτών των χρηστών ονομάζεται $\hat{C}$ και ορίζεται ως εξής:

$$\hat{C} = \{c' \in C \mid [1 - sim_s(c, c')] < n, r_{c',s} \neq \emptyset\} \quad (Εξίσωση 4)$$

Όπου n είναι η απόκλιση της ομοιότητας από την μονάδα για την οποία θεωρούμε αποδεκτή για να συμπεριλάβουμε κάποιον χρήστη στο σύνολο.

Μια απλοϊκή προσέγγιση είναι να πούμε πως η εκτιμώμενη τιμή της $r_{c,s}$, είναι απλά ο μέσος όρος των υπολοίπων N αξιολογήσεων τις οποίες λαμβάνουμε υπόψιν μας, δηλαδή:

$$r_{c,s} = \frac{1}{N} \sum_{c' \in \hat{C}} r_{c',s} \quad (Εξίσωση 5)$$

Εναλλακτικά, μπορούμε να χρησιμοποιήσουμε και συντελεστές βαρύτητας ανάλογους με την ομοιότητα του χρήστη c' στον οποίον ανήκει η αξιολόγηση με τον χρήστη c .

$$r_{c,s} = k \sum_{c' \in \hat{C}} sim_s(c, c') \times r_{c',s} \quad (Εξίσωση 6)$$

$$\text{ή } r_{c,s} = \bar{r}_c + k \sum_{c' \in \hat{C}} sim_s(c, c') \times (r_{c',s} - \bar{r}_c') \quad (Εξίσωση 7)$$

Όπου το k είναι παράγοντας κανονικοποίησης, που πρέπει να εισαχθεί καθώς στο σύνολο υπολογίζεται το γινόμενο της εκάστοτε ομοιότητας με την αντίστοιχη αξιολόγηση.

$$k = 1 / \sum_{c' \in \hat{c}} sim_s(c, c') \times (r_{c',s} - r_c') \quad (\text{Εξίσωση 8})$$

Παρακάτω παρουσιάζεται ένα παράδειγμα χρήσης της μεθόδου.

Παράδειγμα 1

Έστω ότι στον παρακάτω πίνακα καταγράφονται οι αξιολογήσεις των χρηστών Α,Β,Γ,Δ,Ε για τις ταινίες T_1, T_2, T_3, T_4 .

Πίνακας 1. Αξιολογήσεις για το παράδειγμα 1

	T_1	T_2	T_3	T_4	$\bar{r}_c$
A	9	6	2	-	5,6
B	-	4	-	7	5,3
Γ	8	7	-	7,5	7,5
Δ	-	8	6	7	7
E	4	10	5	-	6,3

Ζητούμενο είναι να υπολογίσουμε την εκτιμώμενη αξιολόγηση της T_4 από τον χρήστη Α ($r_{A,4}$). Μέσω της εξίσωσης 1, θα υπολογίσουμε την ομοιότητα του χρήστη Α με τους υπόλοιπους.

Τα αποτελέσματα:

Πίνακας 2. Ομοιότητες Χρηστών για το παράδειγμα 1

$S_{AB} = T_2$	$\text{sim}_s(A,B) = -1$
$S_{A\Gamma} = T_1, T_2$	$\text{sim}_s(A,\Gamma) = 0,757$
$S_{A\Delta} = T_2, T_3$	$\text{sim}_s(A,\Delta) = 0,9157$
$S_{AE} = T_1, T_2, T_3$	$\text{sim}_s(A,E) = 0,2457$

Τώρα θα υπολογίσουμε την εκτιμώμενη τιμή.

Θα επιλέξουμε να χρησιμοποιήσουμε την εξίσωση 6, αφού πρώτα έχουμε θέσει $n=0.8$

Άρα $k = 1,6727$ (ο χρήστης E δεν βρίσκεται στο $\hat{C}$)

$$r_{c,s} = 12,54525/1,6727 = 7,5$$

Εάν χρησιμοποιούσαμε την εξίσωση 2.1.1, με το ίδιο n θα λαμβάναμε την τιμή 7,25

Παρατηρούμε επίσης, πως η εξίσωση 1 παρουσιάζει εσφαλμένα την τιμή -1 (ή 1 αντίστοιχα), εάν δοθεί μόνο μια κοινή αξιολόγηση.

2.2.2 Item Based

Σε αυτήν την περίπτωση, η εκτιμώμενη αξιολόγηση ή αλλιώς χρησιμότητα του αντικειμένου s για τον χρήστη c εξάγεται από τις αξιολογήσεις των αντικειμένων που είναι κοντινότεροι γείτονες με το αντικείμενο s .

Αρχικά, ως προσδιορίζουμε το σύνολο C_{xy} ως το σύνολο όλων των χρηστών c οι οποίοι έχουν αξιολογήσει τόσο το αντικείμενο x όσο και το y .

$$C_{xy} = \{c \in C | r_{c,x} \neq \emptyset, r_{c,y} \neq \emptyset\} \text{ (Εξίσωση 9)}$$

Ορίζουμε λοιπόν, τη συνάρτηση $sim_c(x,y)$ ως την ομοιότητα των αντικειμένων x,y βάση των αξιολογήσεων που από τους χρήστες που ανήκουν στο σύνολο C_{xy} , ως εξής:

$$sim_c(x,y) = \frac{\sum_{c \in C_{xy}} (r_{c,x} - \bar{r}_x)(r_{c,y} - \bar{r}_y)}{\sqrt{\sum_{c \in C_{xy}} (r_{c,x} - \bar{r}_x)^2 \sum_{c \in C_{xy}} (r_{c,y} - \bar{r}_y)^2}} \quad (Εξίσωση 10)$$

όπου $\bar{r}_x, \bar{r}_y$ ο μέσος όρος των αξιολογήσεων που λαμβάνουν τα αντικείμενα x,y .

Σε πλήρη αντιστοιχία με την User Based μεθοδολογία, θα ακολουθήσουμε μία μέθοδο όπου θα εξάγουμε την ζητούμενη αξιολόγηση από τις αξιολογήσεις των κοντινότερων γειτόνων του s (των αντικειμένων δηλαδή που θεωρήθηκαν επαρκώς όμοια με το αντικείμενο s).

Το σύνολο αυτών των χρηστών ονομάζεται $\hat{S}$ και ορίζεται ως εξής:

$$\hat{S} = \{s' \in s \mid [1 - sim_c(s, s')] < n, r_{s',c} \neq \emptyset\} \quad (Εξίσωση 11)$$

όπου n είναι η απόκλιση της ομοιότητας από την μονάδα για την οποία θεωρούμε αποδεκτή για να συμπεριλάβουμε κάποιο αντικείμενο στο σύνολο.

Μια απλοϊκή προσέγγιση είναι να πούμε πως η εκτιμώμενη τιμή της $r_{c,s}$, είναι απλά ο μέσος όρος των υπολοίπων N αξιολογήσεων τις οποίες λαμβάνουμε υπόψιν μας, δηλαδή:

$$r_{c,s} = \frac{1}{N} \sum_{s' \in \hat{S}} r_{c',s} \quad (Εξίσωση 12)$$

Εναλλακτικά, μπορούμε να χρησιμοποιήσουμε και συντελεστές βαρύτητας ανάλογους με την ομοιότητα του αντικειμένου s' με το αντικείμενο s .

$$r_{c,s} = k \sum_{s' \in \hat{S}} sim_c(s, s') \times r_{c,s'} \quad (Εξίσωση 13)$$

ή

$$r_{c,s} = \bar{r}_s + k \sum_{s' \in \hat{S}} sim_c(s, s') \times (r_{c,s'} - \bar{r}_s') \quad (Εξίσωση 14)$$

όπου το k είναι παράγοντας κανονικοποίησης, που πρέπει να εισαχθεί καθώς στο σύνολο υπολογίζεται το γινόμενο της εκάστοτε ομοιότητας με την αντίστοιχη αξιολόγηση.

$$k = 1 / \sum_{s \in \mathcal{S}} sim_c(s, s') \times (r_{c,s'} - r_s') \quad (Εξίσωση 15)$$

Παρακάτω λύνεται το παράδειγμα της προηγούμενης υποενότητας χρησιμοποιώντας την Item Based μέθοδο:

Παράδειγμα 2

Έστω ότι στον παρακάτω πίνακα καταγράφονται οι αξιολογήσεις των χρηστών Α,Β,Γ,Δ,Ε για τις ταινίες T₁,T₂ ,T₃,T₄.

Πίνακας 3. Αξιολογήσεις για το παράδειγμα 2

	T ₁	T ₂	T ₃	T ₄
A	9	6	2	-
B	-	4	-	7
Γ	8	7	-	7,5
Δ	-	8	6	7
E	4	10	5	-
$\bar{r}_s$	7	7	4,33	7,16

Ζητούμενο είναι να υπολογίσουμε την εκτιμώμενη αξιολόγηση της T₄ από τον χρήστη Α (Γ_{A,4})

Μέσω της εξίσωσης 7, θα υπολογίσουμε την ομοιότητα του αντικειμένου 4 με τα υπόλοιπα.

Τα αποτελέσματα παρουσιάζονται στον ακόλουθο πίνακα:

Πίνακας 4. Ομοιότητες Αντικειμένων για το Παράδειγμα 2

$C_{41}=\Gamma$	$\text{sim}_s(4,1) = 1$
$C_{42}=B,\Gamma,\Delta$	$\text{sim}_s(4,2) = 0,433$
$C_{43}=\Delta$	$\text{sim}_s(4,3) = 1$

Τώρα θα υπολογίσουμε την εκτιμώμενη τιμή.

Θα επιλέξουμε να χρησιμοποιήσουμε την εξίσωση 8, θέτοντας $n=0.2$

Άρα $r_{c,s}=7$

Παρόλο που στο συγκεκριμένο παράδειγμα υφίσταται ανακρίβεια λόγω έλλειψης δεδομένων, είναι ενδιαφέρον να παρατηρήσουμε την απόκλιση με το αποτέλεσμα της User Based μεθόδου.

2.3 MODEL BASED (CLUSTER MODELS, BAYESIAN NETWORKS)

Η συγκεκριμένη προσέγγιση αντιμετωπίζει πολλές από της ελλείψεις της memory based προσέγγισης, και επί της ουσίας στοχεύει στη δημιουργία ενός μοντέλου το οποίο δύναται να προβλέπει τις αξιολογήσεις του χρήστη, εκμεταλλευόμενο τα δεδομένα μέσω μηχανικής μάθησης. Υπάρχουν πάρα πολλές διαφορετικές υλοποιήσεις της μεθόδου, οι οποίες βασίζονται κατά κύριο λόγο στις ίδιες αρχές.[15][3][18]

Γενικά σε τέτοιου τύπου μοντέλα, το πρόβλημα προσεγγίζεται από μια πιθανολογική σκοπιά, όπου υπολογίζεται η πιθανότητα μια αξιολόγηση να λάβει μια συγκεκριμένη τιμή, και η εκτιμώμενη αξιολόγηση για τον χρήστη είναι εκείνη για την οποία μεγιστοποιείται η

προαναφερθείσα πιθανότητα. Η εκτιμώμενη αξιολόγηση αυτή για έναν χρήστη c και ένα αντικείμενο s , το οποίο αξιολογείται σε μια ακέραιη κλίμακα από 0 έως m εκφράζεται ως εξής:

$$r_{c,s} = \sum_{i=0}^m P(r_{c,s} = i | r_{c,k}, k \in S_c) i \text{ (Εξίσωση 16)}$$

Στην παρούσα εργασία, παρουσιάζονται υλοποιήσεις οι οποίες έχουν χρησιμοποιηθεί ευρέως, είναι αντιπροσωπευτικές της μεθόδου, και περιγραφόντουσαν επαρκώς στην βιβλιογραφία.

2.3.1 Το πρόβλημα της ελαχιστοποίησης των διαστάσεων – Αλγόριθμος SVD

Θα ξεκινήσουμε περιγράφοντας ένα πολύ συνηθισμένο πρόβλημα στις model based μεθόδους, καθώς και σε πολλές άλλες εφαρμογές, και παρουσιάζοντας τον κύριο τρόπο που αυτό επιλύεται.

Τα συστήματα τα οποία μεταχειρίζονται αυτοί οι αλγόριθμοι περιέχουν κατά κανόνα πάρα πολλούς χρήστες και πάρα πολλά αντικείμενα, οπότε η βάση δεδομένων τους είναι ένας πίνακας με πάρα πολλά στοιχεία. Το διάνυσμα που περιγράφει έναν χρήστη, αποτελείται από τις αξιολογήσεις του για κάθε αντικείμενο του συστήματος. Αυτό είναι ασύμφορο, καθώς οι υπολογισμοί με τόσο μεγάλο όγκο πληροφοριών είναι χρονοβόροι, και επειδή όπως έχει ήδη αναφερθεί, οι περισσότερες θέσεις σε αυτό το διάνυσμα είναι κενές. Επίσης, στόχος είναι να μπορούν να γίνουν υπολογισμοί και για χρήστες που δεν έχουν αξιολογήσει κατά ανάγκη τα ίδια αντικείμενα, διαφορετικά υπονοείται πως δύο χρήστες που δεν έχουν αλληλοεπιδράσει με τα ίδια αντικείμενα, έχουν διαφορετικά χαρακτηριστικά [19].

Προκύπτει λοιπόν η αναγκαιότητα να περιγραφεί ο χρήστης ή το αντικείμενο με την χρήση λιγότερων, «χαμηλού επιπέδου» χαρακτηριστικών[20][21].

Λύση σε αυτό το πρόβλημα προσφέρει η μέθοδος SVD (Singular Value Decomposition-Ανάλυση σε Ιδιάζουσες Τιμές) , η οποία αναλύει έναν πίνακα A σε τρεις μικρότερους πίνακες (U , Σ , V^T) που τον περιγράφουν. Δηλαδή:

$$A = U\Sigma V^T \text{ (Εξίσωση 17)}$$

(Για τα προβλήματα που περιγράφονται, σε ένα σύστημα n χρηστών και m αντικειμένων θεωρούμε πως υφίσταται ένας πίνακας με $n \times m$ λογικές μεταβλητές, όπου κάθε χαρακτηριστικό του χρήστη (δηλαδή κάθε αντικείμενο που υπάρχει στο σύστημα μπορεί να λάβει την τιμή 1 για θετική αξιολόγηση, 0 για αρνητική αξιολόγηση, ή να μην έχει τιμή αν ο χρήστης δεν έχει αλληλοεπιδράσει με το αντικείμενο)

Οι πίνακες U, V είναι ορθογώνιοι πίνακες που αποτελούνται από n ορθοκανονικά διανύσματα, ενώ ο πίνακας Σ είναι διαγώνιος και αποτελείται από n στοιχεία. Οι γραμμές του πίνακα U αντιστοιχούν με τις στήλες του πίνακα A , ενώ οι γραμμές του πίνακα V αντιστοιχούν με τις γραμμές του πίνακα A , ενώ οι τιμές επί της διαγώνιου του πίνακα Σ αποτυπώνουν την βαρύτητα των αντίστοιχων διανυσμάτων στην περιγραφή της πληροφορίας που βρίσκεται στον πίνακα A .

Αυτή η μορφή αναπαράστασης είναι ιδιαίτερα χρήσιμη, καθώς μπορούν πλέον να ιεραρχηθούν οι πληροφορίες που βρίσκονται στα διανύσματα, δηλαδή να λαμβάνονται υπόψιν μόνο εκείνα που είναι, σύμφωνα με τον πίνακα Σ , σημαντικά στην περιγραφή των δεδομένων.

2.3.2 Cluster Models

Στη συγκεκριμένη υλοποίηση, προβλέπεται η αξιολόγηση που θα δώσει ένας χρήστης για ένα αντικείμενο, με βάση τις αξιολογήσεις που έχουνε δώσει οι χρήστες που βρίσκονται στην ίδια κλάση (Cluster) με εκείνον. Με τον όρο κλάση, εννοούνται ομαδοποιήσεις χρηστών που μοιράζονται κοινά χαρακτηριστικά. Υπάρχουν πολλές διαφορετικές μέθοδοι για να χωριστεί ένα σύνολο από χρήστες σε κλάσεις βάση των αξιολογήσεων τους.

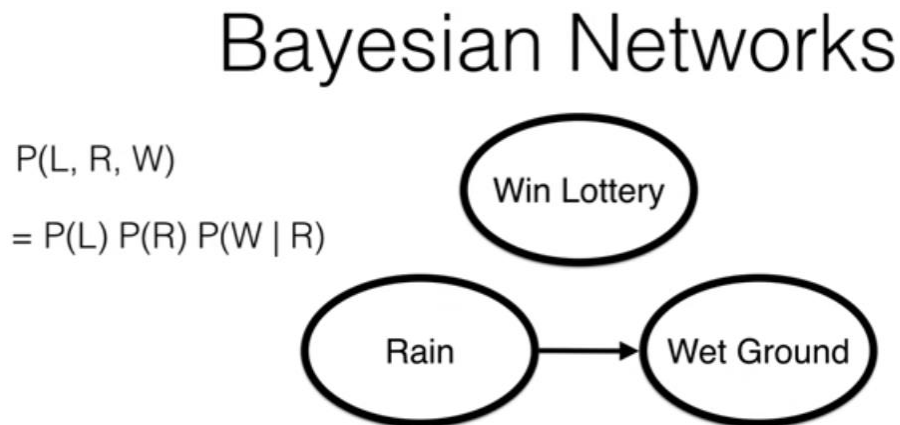
Για μία συγκεκριμένη κλάση, η πιθανότητα ένας οποιοσδήποτε χρήστης που ανήκει σε αυτήν να έχει ένα συγκεκριμένο σύνολο από αξιολογήσεις (οι οποίες μπορούν να λαμβάνουν διάφορες τιμές) για n αντικείμενα, ορίζεται ως εξής;

$$P(\Phi = c, r_1, ., r_2, \dots, r_n) = P(\Phi = c) \prod_{i=1}^n P(r_i | \Phi = c) \text{ (Εξίσωση 18)}$$

Όπου Φ η τυχαία μεταβλητή που συμβολίζει την πιθανότητα του χρήστη να ανήκει στην κλάση C , και $P(r_i | \Phi = c)$ συμβολίζει την πιθανότητα κάποιος χρήστης να έχει δώσει την αξιολόγηση r_i , δοσμένου ότι ανήκει στην κλάση c [15].

Οι παράμετροι του μοντέλου, δηλαδή οι πιθανότητες $P(C=c)$, $P(r_i | C=c)$ υπολογίζονται από την ίδια βάση δεδομένων που αναφέρονται στις προηγούμενες μεθόδους. Εφόσον οι μεταβλητές που καθορίζουν την κάθε κλάση δεν είναι γνωστές εκ των προτέρων, πρέπει να χρησιμοποιούνται μέθοδοι οι οποίες δύνανται να προσδιορίζουν τις παραμέτρους αυτές σε μοντέλα με «κρυφές» μεταβλητές [15].

2.3.3 Bayesian Networks



Εικόνα 2. Σχηματική απεικόνιση ενός Bayesian Network. Φαίνεται πως το ενδεχόμενο L(Win Lottery) είναι ανεξάρτητο από τα άλλα δύο, ενώ τα ενδεχόμενα R(Rain), W(Wet Ground) είναι εξαρτημένα.

Τα Bayesian Networks αποτελούν μία μέθοδο αναπαράστασης λογικών ενδεχομένων που σχετίζονται μεταξύ τους. Κάθε ενδεχόμενο αντιπροσωπεύεται σαν ένας κόμβος, και η ύπαρξη ενός βέλους από ένα ενδεχόμενο προς ένα άλλο υποδηλώνει την εξάρτηση της κατάστασης δεύτερου από το πρώτο.

Στην περίπτωση του Collaborative Filtering, κάθε κόμβος αντιπροσωπεύει ένα αντικείμενο, και μπορεί να βρίσκεται στην κατάσταση όπου ο χρήστης το αξιολογεί θετικά, αρνητικά, ή δεν το αξιολογεί καθόλου, σε αντιστοιχία με το πως περιγράψαμε το ίδιο πρόβλημα νωρίτερα στην ενότητα. Κάθε μία από αυτές τις καταστάσεις έχει την αντίστοιχη της πιθανότητα.

Εφαρμόζεται στο σύστημα των κόμβων ένας αλγόριθμος που το «εκπαιδεύει» σύμφωνα μέσω της βάσης δεδομένων, που προσπαθεί δηλαδή να εντοπίσει την δομή του μοντέλου αναφορικά με την εξάρτηση της κατάστασης του κάθε αντικειμένου από την κατάσταση ενός άλλου. Στόχος δηλαδή είναι να υφίσταται σε ένα «δέντρο», όπου θα είναι πλέον γνωστό το ποια αντικείμενα είναι πιθανό να αξιολογήσει θετικά ο χρήστης εφόσον έχει αξιολογήσει θετικά εκείνα που σχετίζονται μεταξύ τους με τον τρόπο που εξηγήθηκε [15].

2.4 ΣΧΟΛΙΑ-ΑΔΥΝΑΜΙΕΣ-ΣΥΜΠΕΡΑΣΜΑΤΑ

Η συγκεκριμένη μέθοδος έχει εφαρμοστεί ευρέως και έχει πολλά πλεονεκτήματα έναντι άλλων. Για αρχή, πρόκειται για μία μέθοδο που αξιοποιεί τις προτιμήσεις των χρηστών για να πετύχει τον σκοπό της, παράγοντας έτσι πιο ρεαλιστικά αποτελέσματα. Με το να εστιάζει στην ομοιότητα των χρηστών, και όχι των αντικειμένων, μπορεί να παράγει προτάσεις πραγματικά χρήσιμες στον χρήστη, αντί να του προτείνει άλλα αντικείμενα, τα οποία απλά μοιάζουν με εκείνα που αυτός έχει ήδη αξιολογήσει θετικά ή αγοράσει.

Υπάρχουν φυσικά και κάποια μειονεκτήματα. Ένα από τα πιο σοβαρά είναι το πρόβλημα "cold start", η αδυναμία δηλαδή του συστήματος να λειτουργήσει εάν δεν υπάρχουν επαρκείς αξιολογήσεις από χρήστες. Αυτό το πρόβλημα μπορεί να εμφανιστεί είτε όταν ένα νέο αντικείμενο εισάγεται στο σύστημα, είτε όταν ένας νέος χρήστης εισάγεται στο σύστημα, είτε, ακόμα χειρότερα, όταν όλο το σύστημα είναι καινούριο και δεν υπάρχουν επαρκή δεδομένα ώστε να υφίστανται συσχετισμοί. Επί της ουσίας, για να λειτουργήσει το σύστημα βέλτιστα, πρέπει οι χρήστες να αξιολογούν προϊόντα αρκετά συχνά, κάτι το οποίο τις περισσότερες φορές δεν

Τμήμα Μηχανολόγων και Αεροναυπηγών Μηχανικών – Τομέας Διοίκησης και Οργάνωσης 20

συμβαίνει[1]. Φυσικά, υπάρχουν προσπάθειες να λυθεί αυτό το πρόβλημα, όπως π.χ. το να μην λαμβάνονται υπόψιν μόνο οι απευθείας αξιολογήσεις, αλλά και άλλα δεδομένα όπως π.χ. το με ποια αντικείμενα ήρθε σε επαφή ο χρήστης, ή το πόσην ώρα διήρκεσε εκείνη.

Το πρόβλημα της αραιότητας είναι επίσης συχνό σε τέτοια συστήματα, καθώς συχνά λείπουν δεδομένα. Για αυτόν ακριβώς το λόγο, όπως προαναφέραμε, έχουν αναπτυχθεί τα model based συστήματα, τα οποία και χρησιμοποιούνται κατά κόρον[15].

Με βάση τα παραπάνω, είναι κατανοητό πως οι user/item based υλοποιήσεις των CF(Collaborative Filtering)[22] αλγόριθμων, είναι πιο κατάλληλες για συστήματα με πολλούς χρήστες και λίγα αντικείμενα, και όχι το αντίθετο. Σε μια περίπτωση με λίγα αντικείμενα και πολλούς χρήστες, είναι πολύ πιθανό για έναν χρήστη να μπορούν να βρεθούν λίγοι έως κανένας παρόμοιοι χρήστες.. Ένα άλλο ζήτημα φαίνεται να είναι εκείνο της κατανάλωσης πόρων, καθώς απαιτείται σχετικά μεγάλη υπολογιστική δύναμη για να τρέξει ένας τέτοιος αλγόριθμος.

Οι model based υλοποιήσεις χρησιμοποιούνται πολύ περισσότερο, και συνδυασμό με content based στοιχεία, αποτελούν την βάση για υβριδικούς αλγορίθμους, οι οποίοι, όπως περιγράφεται σε επόμενες ενότητες, αποτελούν την πιο σύγχρονη λύση στα συστήματα συστάσεων.[23]

3. ΦΙΛΤΡΑΡΙΣΜΑ ΒΑΣΗ ΠΕΡΙΕΧΟΜΕΝΟΥ - CONTENT BASED FILTERING (CBF)

3.1 ΕΙΣΑΓΩΓΗ

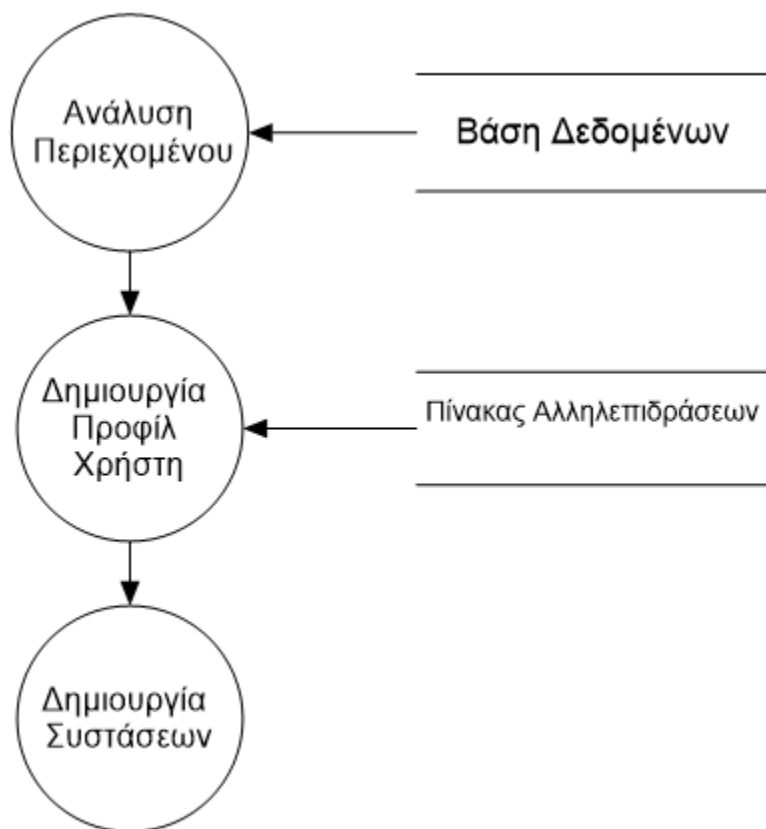
Πρόκειται για την δεύτερη μεγάλη κατηγορία που παρουσιάζεται.

Σε αυτήν την περίπτωση, η εκτιμώμενη χρησιμότητα ενός αντικειμένου για τον χρήστη, δεν εξάγεται με βάση την υπαρκτή χρησιμότητα αυτού του αντικειμένου για άλλους χρήστες, αλλά με βάση το πως έχει αξιολογήσει ο χρήστης άλλα, παρόμοια αντικείμενα[1].

Επί της ουσίας, τα αντικείμενα συγκρίνονται μεταξύ τους, και το σύστημα προσπαθεί να βρει αντικείμενα με κοινά χαρακτηριστικά που ταιριάζουν στο «προφίλ» του εκάστοτε χρήστη το οποίο έχει διαμορφωθεί από την αλληλεπίδραση του με άλλα αντικείμενα [24].

Η μέθοδος αυτή είναι επίσης αρκετά παλιά και ευρέως διαδεδομένη σε εφαρμογές που σχετίζονται με έγγραφα ή άλλα αντικείμενα σε γραπτή μορφή, καθώς εκείνα είναι σε σχέση με τα πολυμέσα, πολύ πιο εύκολο να αναλυθούνε αυτόματα ως προς τα περιεχόμενα και τα χαρακτηριστικά τους[1][6][25].

Για να παρουσιαστεί αυτή η μέθοδος, θα διαχωριστεί σε τρία ξεχωριστά βήματα [6]. Τα βήματα αυτά είναι η ανάλυση περιεχομένου, η δημιουργία προφίλ χρήστη, και η δημιουργία συστάσεων.



Εικόνα 3. Διάγραμμα Ροής Δεδομένων για ένα τυπικό CBF Σύστημα

Είναι σαφές πως δεν ακολουθούν όλοι οι αλγόριθμοι που βρίσκονται στην κατηγορία CBF ακριβώς αυτά τα βήματα. Ενδεχομένως να υπάρχουν επιπλέον διαδικασίες, ή κάποιες από εκείνες που έχουν ήδη περιγραφεί στην παρούσα εργασία να εκτελούνται με διαφορετικό τρόπο ή σειρά.

3.2 ΑΝΑΛΥΣΗ ΠΕΡΙΕΧΟΜΕΝΟΥ

Το πρώτο βήμα αποτελεί την ανάλυση των αντικειμένων σε στοιχεία που τα περιγράφουν. Η διαδικασία αυτή είναι εφικτό να γίνει από τους σχεδιαστές του συστήματος, ωστόσο κάτι τέτοιο είναι ιδιαίτερα χρονοβόρο για πληθώρα αντικειμένων. Μια τέτοια προσέγγιση θα ήταν εφικτή π.χ. σε ένα e-shop, όπου τα αντικείμενα δεν ανανεώνονται από τους χρήστες αλλά από τους συντηρητές, και πέφτουν σε συγκεκριμένες κατηγορίες. Διαφορετικά, τα αντικείμενα που προέρχονται από την εκάστοτε πηγή πληροφοριών αναλύονται από κάποιον αλγόριθμο ο οποίος εξάγει στοιχεία.

3.2.1 TF-IDF

Θα παρουσιαστεί ένας ευρέως διαδεδομένος αλγόριθμος που χρησιμοποιείται για την κατηγοριοποίηση εγγράφων, ο TF-IDF. Τα αρχικά σημαίνουν Term Frequency - Inverse Document Frequency, και αντιστοιχούν στα δύο μέρη της συνάρτησης που θα παρουσιαστεί [1][26].

Έστω λοιπόν πως το σύστημα διαχειρίζεται έγγραφα. Κάθε αντικείμενο(έγγραφο d) μπορεί να περιγραφεί μέσω ενός διανύσματος $X(d)=(x_1,x_2,x_3,\dots,x_n)$, όπου το x_i είναι το βάρος το οποίο αντιστοιχεί στους διαφορετικούς όρους που εντοπίζονται μέσα στο έγγραφο.

Μια απλοϊκή προσέγγιση είναι πως μεγαλύτερο βάρος θα πρέπει να έχουν οι όροι που επαναλαμβάνονται περισσότερες φορές μέσα στο κείμενο. Αυτή η προσέγγιση ωστόσο, δεν

μπορεί να εφαρμοστεί επιτυχημένα για όρους όπως άρθρα, ανωνυμίες κλπ οι οποίοι εμφανίζονται συχνά χωρίς να αποτελούν ειδοποιό διαφορά ανάμεσα στα έγγραφα (δηλαδή, το γεγονός ότι ένα κείμενο περιέχει πολλές φορές την λέξη «το» δεν είναι σωστό να λογίζεται ως σημαντικό χαρακτηριστικό του).

Παρουσιάζεται λοιπόν η ανάγκη να υπάρχει ένας παράγοντας που θα προσδίδει λιγότερο βάρος σε όρους που επαναλαμβάνονται ούτως ή άλλως σε πολλά έγγραφα. Η τελική μορφή της συνάρτησης βάρους για τον κάθε όρο i είναι η εξής:

$$x_{d,i} = f_{d,i} \times \log \left(\frac{N}{n_i} \right) \text{ (Εξίσωση 15)}$$

Όπου

$f_{d,i}$: Η συχνότητα με την οποία εμφανίζεται ο όρος στο έγγραφο d

N : Ο συνολικός αριθμός των εγγράφων

n_i : Ο αριθμός των εγγράφων στα οποία παρουσιάζεται ο όρος έστω και μία φορά

Όταν το αντικείμενο αναλυθεί και αντιστοιχιστεί με ένα διάνυσμα που το περιγράφει (ή οποιαδήποτε άλλη μορφή αναπαράστασης χρησιμοποιεί ο εκάστοτε αλγόριθμος), τότε θεωρείται πως πλέον αναπαρίσταται, και το διάνυσμα αποθηκεύεται σε μια βάση δεδομένων (Πίνακα Αντικειμένων).

3.2.2 Naïve Bayes Classifier

Αντίστοιχα με τους CF αλγορίθμους, υπάρχουν model based υλοποιήσεις για τα CBF συστήματα, οι οποίοι χρησιμοποιούν μηχανική μάθηση, δηλαδή κατηγοριοποίηση σε κλάσεις, Τμήμα Μηχανολόγων και Αεροναυπηγών Μηχανικών – Τομέας Διοίκησης και Οργάνωσης 26

δέντρα αποφάσεων και νευρωνικά δίκτυα[27]. Με την κατηγοριοποίηση σε κλάσεις, γίνεται πιο εύκολη η ταξινόμηση των εγγράφων σε θεματικές κατηγορίες, και έπειτα το να ταιριάζουν με το προφίλ του κάθε χρήστη.

Στην περίπτωση του Naïve Bayes Classification, γίνεται μια απλή εφαρμογή του θεωρήματος Bayes[28] ώστε να κατηγοριοποιηθούν έγγραφα σε κλάσεις ανάλογα με τα περιεχόμενα τους. Το θεώρημα Bayes είναι το εξής:

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)} \text{ (Εξίσωση 19)}$$

Έστω πως η πιθανότητες $P(k_{1,s})$ για κάθε ενδεχόμενο που αντιστοιχεί σε κάποιο χαρακτηριστικό(λέξη κλειδί που αναφέρεται σε αυτό) του αντικειμένου που εξετάζεται να πραγματοποιηθεί(να υπάρχει στο έγγραφο s), είναι ανεξάρτητες μεταξύ τους. Εάν λοιπόν, $P(C_i)$ είναι η πιθανότητα του εγγράφου i να βρίσκεται στην κλάση C , τότε η πιθανότητα να ανήκει το έγγραφο στην κλάση αυτή δοσμένων των λέξεων κλειδιά που βρίσκονται σε αυτό εκφράζεται ως:

$$P(C_i|k_{1s} \& \dots \& k_{ns}) \text{ (Εξίσωση 20)}$$

Εάν υποθεθεί λοιπόν, πως υπάρχει πλήρης ανεξαρτησία ανάμεσα στην πιθανότητα ύπαρξης κάθε διαφορετικής λέξης κλειδί στο έγγραφο που εξετάζεται, τότε η παραπάνω πιθανότητα ισοδυναμεί με:

$$P(C_i) \sum_{x=1}^n (k_{xs}|C_i) \text{ (Εξίσωση 21)}$$

Οι πιθανότητες $P(k_{xs}), P(C_i)$ μπορούν σαφώς να προσδιοριστούν από τα δεδομένα που παρέχονται στον αλγόριθμο. Με αυτόν τον τρόπο, προκύπτει η πιθανότητα για κάθε κλάση, και έτσι το κάθε έγγραφο ταξινομείται σε εκείνη με την υψηλότερη πιθανότητα.

3.3 ΔΗΜΙΟΥΡΓΙΑ ΠΡΟΦΙΛ ΧΡΗΣΤΗ

Για να δημιουργηθεί το προφίλ(Διάνυσμα) που αναπαριστά τις προτιμήσεις του χρήστη c , οι αλληλεπιδράσεις του και οι αξιολογήσεις του με τα διάφορα αντικείμενα πρέπει να αποθηκεύονται στον Πίνακα Αλληλεπιδράσεων [6].

Σε αντιστοιχία με την περίπτωση των CF αλγορίθμων, στο σύστημα λαμβάνονται υπόψη τόσο οι άμεσες αξιολογήσεις του χρήστη(μηχανισμοί αξιολόγησης, like, dislike κλπ.), όσο και οι έμμεσες(επαφή με ένα αντικείμενο, κατέβασμα ενός εγγράφου, κλπ.). Στην περίπτωση της αξιοποίησης των αλληλεπιδράσεων του χρήστη, πρέπει να ορίζεται ένας διαφορετικός δείκτης σημαντικότητας για κάθε διαφορετικού τύπου αλληλεπίδραση χρήστη-αντικειμένου.

Το προφίλ $X(c)=(x_{c,1},x_{c,2},\dots,x_{c,n})$ του χρήστη c προκύπτει από την αξιοποίηση των δεδομένων αυτών, καθώς και από τις προτιμήσεις τις οποίες ενδεχομένως ο χρήστης έχει ήδη δηλώσει άμεσα (υπάρχουν πολλά συστήματα που ζητάνε από τον χρήστη να δηλώσει στοιχεία για τον ίδιο, όπως π.χ. τα αγαπημένα του είδη μουσικής), ή από λοιπά στοιχεία τα οποία έχει δηλώσει όπως ηλικία, φύλο, τοποθεσία (η τεχνική αυτή ονομάζεται Stereotyping και θα παρουσιαστεί σε επόμενη ενότητα)

3.4 ΔΗΜΙΟΥΡΓΙΑ ΣΥΣΤΑΣΕΩΝ

Σε έναν αλγόριθμο που χρησιμοποιεί Content Based Filtering (ή αλλιώς CBF), η εκτιμώμενη αξιολόγηση $r(c,s)$ ενός αντικειμένου s υπολογίζεται βάση της ομοιότητας του διανύσματος που

αναπαριστά το συγκεκριμένο αντικείμενο με το προφίλ του χρήστη, με χρήση εξισώσεων όπως οι 1 και 2, που χρησιμοποιούν ωστόσο σαν εισόδους τα δύο μονοδιάστατα διανύσματα:

$$\text{sim}(X(c), X(d)) = \frac{X(c) \cdot X(d)}{|X(c)| |X(d)|} \quad (\text{Εξίσωση 22})$$

Με κριτήριο την τιμή της συνάρτησης αυτής, προκύπτουν και ταξινομούνται οι προτάσεις που γίνονται στον χρήστη στην λίστα προτάσεων L_c . Υπάρχουν ωστόσο περιθώρια, τόσο το διάνυσμα που αναπαριστά τον χρήστη (Προφίλ Χρήστη), όσο και εκείνο που αναπαριστά το κάθε αντικείμενο να μεταβληθούν. Οι προτιμήσεις του χρήστη αλλάζουν, οπότε με τις νέες του αξιολογήσεις αλλάζει το προφίλ του και προστίθενται νέα αντικείμενα στη λίστα L_c [6]. Επίσης πολύ συχνά, ζητείται από τον χρήστη να αξιολογήσει τις προτάσεις που του γίνανε σε εκείνον, ώστε εάν ο ίδιος πιστεύει πως το αντικείμενο αξιολογήθηκε λάθος από τον αλγόριθμο, τότε το σύστημα να αλλάξει τα χαρακτηριστικά του[29]. Αυτές οι διαδικασίες επαναλαμβάνονται πολλές φορές, ώστε να προκύψουν βέλτιστα αποτελέσματα.

Στη συνέχεια θα παρουσιαστεί ένα παράδειγμα χρήσης της μεθόδου

Παράδειγμα 3

Έστω ο χρήστης c σε ένα σύστημα CBF. Ο χρήστης c αναπαρίσταται στο σύστημα με την βοήθεια του διανύσματος $X(c)$.

Στον παρακάτω πίνακα φαίνονται οι τιμές 5 χαρακτηριστικών A, B, Γ, Δ, E για αντικείμενα τα οποία έχει αξιολογήσει θετικά ο χρήστης c . Από τον μέσο όρο των τιμών των χαρακτηριστικών αυτών προκύπτει το προφίλ του χρήστη.

Πίνακας 5. Χαρακτηριστικά Αντικειμένων για το Παράδειγμα 3

	A	B	Γ	Δ	E
Αντικείμενο 1	6	8	2	5	9
Αντικείμενο 2	5	6	2	4	10
Αντικείμενο 3	8	9	3	5	4
Αντικείμενο 4	10	2	1	8	7
Αντικείμενο 5	5	7	4	7	8
X(c)	6.8	6.4	2.4	5.8	7.4

Προκύπτει άρα το διάνυσμα $X(c) = [6.8 \ 6.4 \ 2.4 \ 5.8 \ 7.4]$

Έστω τώρα τα αντικείμενα s_1, s_2 με τα διανύσματα

$$X(s_1) = [2.5 \ 6.9 \ 5.2 \ 6.7 \ 5.0],$$

$$X(s_2) = [5.4 \ 4.6 \ 3.0 \ 7.8 \ 9.1]$$

Με την βοήθεια της εξίσωσης 22, θα υπολογίσουμε την ομοιότητα των δύο αντικειμένων με το προφίλ του χρήστη.

Και προκύπτει:

$$\text{sim}(X(c), X(s_1)) = 0.84488$$

$$\text{sim}(X(c), X(s_2)) = 1.0973$$

Πιο κοντά στην τιμή 1 είναι η τιμή $\text{sim}(X(c), X(s_2))$, άρα το αντικείμενο s_2 είναι πιο πιθανό να αρέσει στον χρήστη c .

3.5 ΣΧΟΛΙΑ-ΑΔΥΝΑΜΙΕΣ-ΣΥΜΠΕΡΑΣΜΑΤΑ

Η εν λόγω μέθοδος είναι κατά μία έννοια η αντίστροφη της CF, καθώς εξετάζει την χρησιμότητα πολλών αντικειμένων αναφορικά με έναν μόνο χρήστη, ενώ η CF εξετάζει ένα αντικείμενο σε σχέση με πολλούς χρήστες.

Για τον λόγο αυτό, στην CBF δεν συναντάται το πρόβλημα *cold start* με τον ίδιο τρόπο που συναντάται στην CF, καθώς δεν είναι αναγκαίο να υπάρχουν πολλοί χρήστες (οι οποίοι έχουν μάλιστα αξιολογήσει ή έρθει σε επαφή με αντικείμενα) στο σύστημα ώστε αυτό να μπορεί να παράγει προτάσεις. Οι προτάσεις που παράγονται και η αξιοπιστία εκείνων, εξαρτώνται από την δραστηριότητα του χρήστη και μόνο. Το πρόβλημα *cold start* συνεχίζει βέβαια να υφίσταται, μέχρι ο ίδιος ο χρήστης να αξιολογήσει έναν κρίσιμο αριθμό αντικειμένων.

Σχετικά με την περίπτωση των εγγράφων, συμπεραίνει κανείς εύκολα πως το ποιοι όροι αναφέρονται σε ένα κείμενο δεν αποτελεί σε καμία περίπτωση ικανοποιητικό κριτήριο για την αξιολόγηση τους. Ένα άρθρο μπορεί να έχει το ίδιο ή πολύ όμοιο διάλυσμα με ένα άλλο, αλλά να είναι πολύ πιο κακογραμμένο από αυτό.

Ένα άλλο σημαντικό πρόβλημα της μεθόδου, είναι η μεγάλη πιθανότητα ο χρήστης να «εγκλωβιστεί» μεταξύ πολύ σχετικών μεταξύ τους αντικειμένων. Εάν το σύστημα του προτείνει μονάχα αντικείμενα τα οποία μοιάζουν μεταξύ τους, τότε μπορεί πιθανά ο χρήστης να βρίσκει με ακρίβεια εκείνα που ταιριάζουν στο προφίλ του, αλλά να μην ανακαλύπτει καινούργια αντικείμενα που θα του ήταν επίσης χρήσιμα, όσο ποιοτικά και εάν είναι αυτά, εάν δεν σχετίζονται με εκείνα που αυτός έχει ήδη έρθει σε επαφή (π.χ., εάν ένας χρήστης βλέπει μόνο ταινίες επιστημονικής φαντασίας, το σύστημα δε θα του πρότεινε ποτέ τον *Τιτανικό*). Επίσης, αν μιλάμε για εμπορικές εφαρμογές, εάν ο χρήστης έχει αγοράσει ένα συγκεκριμένο τύπο προϊόντος το οποίο δεν αγοράζει συχνά (π.χ. αυτοκίνητο, κινητό, καναπέ), δεν θα έχει κανένα όφελος εάν το σύστημα του προτείνει αντίστοιχα αντικείμενα. Το πρόβλημα αυτό μπορεί να λυθεί προσθέτοντας κάποια τυχαιότητα στις προτάσεις που παρουσιάζονται στον χρήστη, ή με την χρήση υβριδικών μεθόδων, προτείνοντας π.χ. αντικείμενα τα οποία βρίσκονται υψηλά στις προτιμήσεις όλων των χρηστών του συστήματος τη συγκεκριμένη χρονική στιγμή (μέθοδος Global Relevance).

Ένα πρακτικό σύστημα που προκύπτει, είναι η δυσκολία εφαρμογής της μεθόδου σε αντικείμενα των οποίων τα χαρακτηριστικά δεν μπορούν να αποτυπωθούν αυτόματα. Το παράδειγμα του e-shop είναι ένα τέτοιο. Ενώ ένα σύνολο από έγγραφα, μπορεί να μετατραπεί με σχετική αξιοπιστία σε διανύσματα μέσω του αλγορίθμου TFIDF, τα χαρακτηριστικά ενός συνόλου από ηλεκτρικές συσκευές θα πρέπει να περαστούν στο σύστημα ξεχωριστά για κάθε μία από αυτές.

4. ΣΥΜΠΛΗΡΩΜΑΤΙΚΟΙ ΑΛΓΟΡΙΘΜΟΙ

4.1 ΕΙΣΑΓΩΓΗ

Οι δύο κατηγορίες αλγορίθμων που παρουσιάστηκαν παραπάνω αποτελούν στην συντριπτική πλειοψηφία, μαζί με τους υβριδικούς συνδυασμούς τους, την βάση για κάθε ολοκληρωμένο σύστημα συστάσεων. Ωστόσο, υπάρχουν και κάποιοι επιπλέον αλγόριθμοι-κριτήρια, που χρησιμοποιούνται, συνήθως συμπληρωματικά, για την κατάταξη των αντικειμένων ή για την δημιουργία επιπλέον προτάσεων.

Χάρη σε αυτούς τους αλγορίθμους, μπορούν να υπερνικούνται προβλήματα και να καλύπτονται κενά που προκύπτουν στο κυρίως σύστημα, αλλά συνήθως δεν παρέχουν υψηλό επίπεδο εξατομίκευσης.

4.2 STEREOTYPING

Πρόκειται για μία από τις πιο πρώιμες μεθόδους, η οποία όμως είναι ακόμα αποτελεσματική, σε συνδυασμό συνήθως με άλλες μεθόδους (Hybrid Methods) [30]

Η μέθοδος χρησιμοποιεί στερεότυπα όπως ακριβώς αυτά ορίζονται από την επιστήμη της ψυχολογίας. Όπως δηλαδή οι άνθρωποι συχνά κρίνουν κάποιον και σχηματίζουν μια εικόνα για

εκείνον βασιζόμενοι σε πολύ λίγα πράγματα που γνωρίζουνε για αυτόν, ο αλγόριθμος χρησιμοποιεί ορισμένα (ενδεικτικά) χαρακτηριστικά του χρήστη, για να τον εντάξει σε μία κλάση χρηστών που τους αρέσουν τα ίδια πράγματα.

Τα χαρακτηριστικά αυτά μπορούν ενίοτε να μην σχετίζονται καν με την επαφή που έχει ο χρήστης με τα αντικείμενα (χωρίς αυτό, φυσικά, να αποκλείεται) . Μερικά από αυτά μπορούν να είναι η ηλικία, το φύλο, η γεωγραφική τοποθεσία ,και άλλα.

Η μέθοδος έχει κάποια αξιοσημείωτα πλεονεκτήματα. Αρχικά, υπερνικά το πρόβλημα *cold start*, το οποίο είχε παρουσιαστεί από κοινού στις προηγούμενες δύο μεθόδους που εξετάσαμε. Με πολύ λίγα δεδομένα, το σύστημα μπορεί να προτείνει σχετικά αξιόπιστες προτάσεις, πράγμα ιδιαίτερα χρήσιμο σε ένα σύστημα που περιέχει πολλά αντικείμενα. Για παράδειγμα, σε ένα κατάστημα με ρούχα, είναι συνήθως πολύ χρήσιμο και αυτονόητο σε έναν άντρα να προτείνονται αντρικά ρούχα, καθώς στην πλειοψηφία των περιπτώσεων, αυτά είναι τα αντικείμενα που αυτός θα αναζητάει.

Η μέθοδος γενικά δεν απαιτεί υψηλή συμμετοχή από τους χρήστες, οπότε κρίνεται κατάλληλη για εφαρμογές που ο χρήστης δε θα χρησιμοποιήσει μακροπρόθεσμα, και πρέπει να του προταθούν αντικείμενα γρήγορα και αποτελεσματικά. Επίσης, η μέθοδος δεν έχει υψηλές απαιτήσεις σε υπολογιστική ισχύ.

Η μέθοδος έχει χρησιμότητα σε μηχανές αναζήτησης, καθώς διάφοροι παράγοντες όπως η τοποθεσία και η ηλικία καθορίζουν τα αποτελέσματα που θα είναι πιο χρήσιμα για τον κάθε χρήστη (πχ, όταν ένας φοιτητής που βρίσκεται στην Πάτρα αναζητήσει τους όρους «eclass», «Μηχανολόγοι μηχανικοί πρόγραμμα εξεταστικής» , και «εφημερεύοντα φαρμακεία», τα

αποτελέσματα που θα σχετίζονται με την πλατφόρμα του ΕΜΠ, και τα φαρμακεία που βρίσκονται στην Αθήνα, αν και θα έχουν περισσότερη επισκεψιμότητα, θα του είναι παντελώς άχρηστα).

Τα προβλήματα είναι επίσης αρκετά και καθιστούν την μέθοδο ανεπαρκή για τις περισσότερες εφαρμογές. Ουσιαστικά, δεν παρέχεται μεγάλος βαθμός εξατομίκευσης στις προτάσεις που παρέχει, και υπάρχει μεγάλη πιθανότητα σφάλματος (πχ, τα αθλητικά είδη αρέσουν σε πολλούς άνδρες ηλικίας 16-25, αλλά όχι σε όλους. Για τους υπόλοιπους, η πρόταση είναι άχρηστη). Το σφάλμα αυτό για να αποφευχθεί, πρέπει να υπάρχουν περισσότερα δεδομένα σχετικά με τον χρήστη, αλλά αυτό καταλήγει στο να αυξάνει πολύ την πολυπλοκότητα και την δυσκολία σχεδιασμού του συστήματος, καθώς πρέπει να σχηματιστούν πολλά στερεότυπα και υποκατηγορίες. Είναι επίσης προφανές πως κάθε νέο αντικείμενο που εισάγεται στο σύστημα, πρέπει να εντάσσεται ξεχωριστά σε ένα ή παραπάνω στερεότυπα.

Γενικά, όπως και οι άλλοι συμπληρωματικοί αλγόριθμοι, είναι χρήσιμος για να καλυφθούν κάποια κενά, όπως η παροχή προτάσεων όταν υπάρχουν ανεπαρκή δεδομένα από τον χρήστη, ή παροχή «τυχαίων» προτάσεων οι οποίες δεν είναι όμως άχρηστες, οπότε και πιο κατάλληλος να χρησιμοποιηθεί σε κάποια υβριδική υλοποίηση, όπως αναφέρθηκε παραπάνω.

4.3 GLOBAL RELEVANCE

Μια ακόμα συμπληρωματική μέθοδος.

Η συγκεκριμένη μέθοδος, δεν παράγει εξατομικευμένες προτάσεις για τον κάθε χρήστη, αλλά υιοθετεί μια διαφορετική προσέγγιση. Με ορισμένα κριτήρια, όπως η δημοτικότητα ή οι θετικές αξιολογήσεις ενός αντικειμένου, παράγει προτάσεις που αφορούν όλους τους χρήστες.

Επί της ουσίας, η μέθοδος μπορεί να αποτελέσει ένα δευτερεύον κριτήριο ώστε να ταξινομηθούν προτάσεις που παράγονται από έναν αλγόριθμο CF ή CBF. Σε πολλές περιπτώσεις, αυτό είναι πολύ χρήσιμο καθώς έτσι φτάνουν στον χρήστη νέα αντικείμενα, όπως μια νέα μουσική κυκλοφορία, η οποία όμως θα ήταν ούτως ή άλλως κοντά στις προτιμήσεις του χρήστη.

4.4 CO-OCCURRENCE

Η μέθοδος αυτή, αντί να εστιάζει στην ομοιότητα δύο αντικειμένων, εστιάζει στο πόσο συχνά έρχεται κάποιος σε επαφή και με τα δύο. Αν μιλάμε για έγγραφα, μπορεί το κριτήριο να είναι το πόσο συχνά γίνεται αναφορά και στα δύο έγγραφα σε μια δημοσίευση. Εάν μιλάμε για προϊόντα, το κριτήριο συνήθως είναι το πόσο συχνά αυτά αγοράζονται μαζί.

Το ότι η μέθοδος συσχετίζει αντικείμενα αντί να τα συγκρίνει με βάση τα χαρακτηριστικά τους αποτελεί συχνά πλεονέκτημα. Η CBF για παράδειγμα, προτείνει παρόμοια αντικείμενα στον χρήστη, και αυτό έχουμε ήδη εξηγήσει πως μπορεί να είναι πρόβλημα. Αυτή η προσέγγιση είναι επίσης λιγότερο απαιτητική σε πόρους, αφού δεν απαιτείται η ανάλυση του ίδιου του αντικειμένου.

Από την άλλη πλευρά, αυτό σημαίνει πως ένα αντικείμενο δεν μπορεί να προταθεί αν δεν αγοραστεί μαζί με ένα άλλο τουλάχιστον μία φορά, και αυτό καθιστά συχνά την μέθοδο ανεπαρκή, ιδιαίτερα σε ένα σύστημα που εισάγονται συχνά καινούρια αντικείμενα.

5. ΥΒΡΙΔΙΚΕΣ ΜΕΘΟΔΟΙ

5.1 ΕΙΣΑΓΩΓΗ

Ουσιαστικά, υβριδικές είναι σχεδόν όλες οι μέθοδοι που εφαρμόζονται στην πράξη. Σπάνια κάποια από τις προηγούμενες κατηγορίες να υλοποιείται «καθαρά». Αυτό συμβαίνει επειδή στην πράξη, μέσω της σύνθεσης διαφορετικών μεθόδων καλύπτονται οι αδυναμίες που έχουν η κάθε μια ξεχωριστά, και επειδή τα πραγματικά προβλήματα είναι πολύπλοκα και διαφορετικά μεταξύ τους οπότε χρειάζεται η κάθε λύση να προσαρμόζεται σε αυτά

Είναι δύσκολο να περιγραφεί επακριβώς το πως δομείται ένας υβριδικός αλγόριθμος, ακριβώς επειδή δεν υπάρχει κάποια «συνταγή», αλλά πολλοί διαφορετικοί τρόποι υλοποίησης. Σε δύο από τις μελέτες που περιλαμβάνονται στην βιβλιογραφία αυτής της εργασίας, έχει γίνει μια κατηγοριοποίηση των υβριδικών αλγορίθμων[31] [1] [11]. Θα παρουσιαστούν ορισμένες αντιπροσωπευτικές κατηγορίες.

5.2 ΣΥΝΔΥΑΣΤΙΚΑ ΜΟΝΤΕΛΑ – ΣΥΓΚΕΡΑΣΜΟΣ ΠΡΟΤΑΣΕΩΝ

Σε αυτήν την περίπτωση, δύο οι παραπάνω αλγόριθμοι εκτελούνται ανεξάρτητα και μετά υφίσταται κάποιας μορφής συγκερασμός των αποτελεσμάτων τους. Μια λύση είναι αυτά να συνδυαστούν γραμμικά, και η τελική πρόταση να είναι ο μέσος όρος των αποτελεσμάτων των

αλγορίθμων. Είναι επίσης εφικτό να παρουσιάζονται κάθε φορά τα αποτελέσματα από έναν μονάχα αλγόριθμο, ανάλογα με το ποια αποτελέσματα αξιολογούνται καλύτερα με βάση κάποια τρίτα κριτήρια.

Θα παρουσιαστούν μερικές προσεγγίσεις που πέφτουν σε αυτήν την κατηγορία, όπως αυτές παρουσιάζονται από τον Robin Burke [32]

5.2.1 Συνδυασμός Προτάσεων με συντελεστές βαρύτητας

Ένα υβρίδιο που χρησιμοποιεί αυτήν την τεχνική, θα κατατάξει τις προτάσεις του συμψηφίζοντας την εκτιμώμενη αξιολόγηση που δίνουν διαφορετικοί αλγόριθμοι ξεχωριστά, δίνοντας διαφορετικό συντελεστή βαρύτητας στην αξιολόγηση κάθε διαφορετικού αλγορίθμου.

Στην πραγματικότητα, υφίσταται ένα σύστημα «ψηφοφορίας», όπου οι διαφορετικοί αλγόριθμοι «ψηφίζουν» ποια πρόταση θα παρουσιαστεί στον χρήστη, και οι συντελεστές βαρύτητας αντιπροσωπεύουν το πόσες «ψήφοι» αναλογούν στον κάθε αλγόριθμο. Οι συντελεστές αυτοί στην αρχή θα είναι κατά κανόνα ίδιοι για όλους τους αλγόριθμους, αλλά έπειτα αλλάζουν ανάλογα με το πόσο χρήσιμη αξιολογεί την κάθε πρόταση ο χρήστης.

Εάν δηλαδή, ο χρήστης επικυρώνει πως οι προτάσεις που προέκυψαν από έναν CF αλγόριθμο ήταν οι πιο χρήσιμες, τότε το μοντέλο του συγκεκριμένου χρήστη θα δίνει στο μέλλον μεγαλύτερη βαρύτητα σε αυτές.

Η συγκεκριμένη μέθοδος είναι ιδιαίτερα απλή και κατανοητή στην δομή της, και μας επιτρέπει την κατανόηση και την τροποποίηση σε μεγάλο βαθμό του πως παράγονται οι

προτάσεις. Μια μεγάλη της αδυναμία ωστόσο, είναι το γεγονός πως υπονοεί ότι για κάθε χρήστη, όλα τα αντικείμενα αντιμετωπίζονται παρόμοια και τα κριτήρια διαμορφώνονται με τον ίδιο τρόπο, ενώ για διαφορετικά αντικείμενα κάτι τέτοιο είναι πολύ πιθανό να μην ισχύει.

5.2.2 Συνδυασμός Προτάσεων με Εναλλαγή Αλγορίθμων

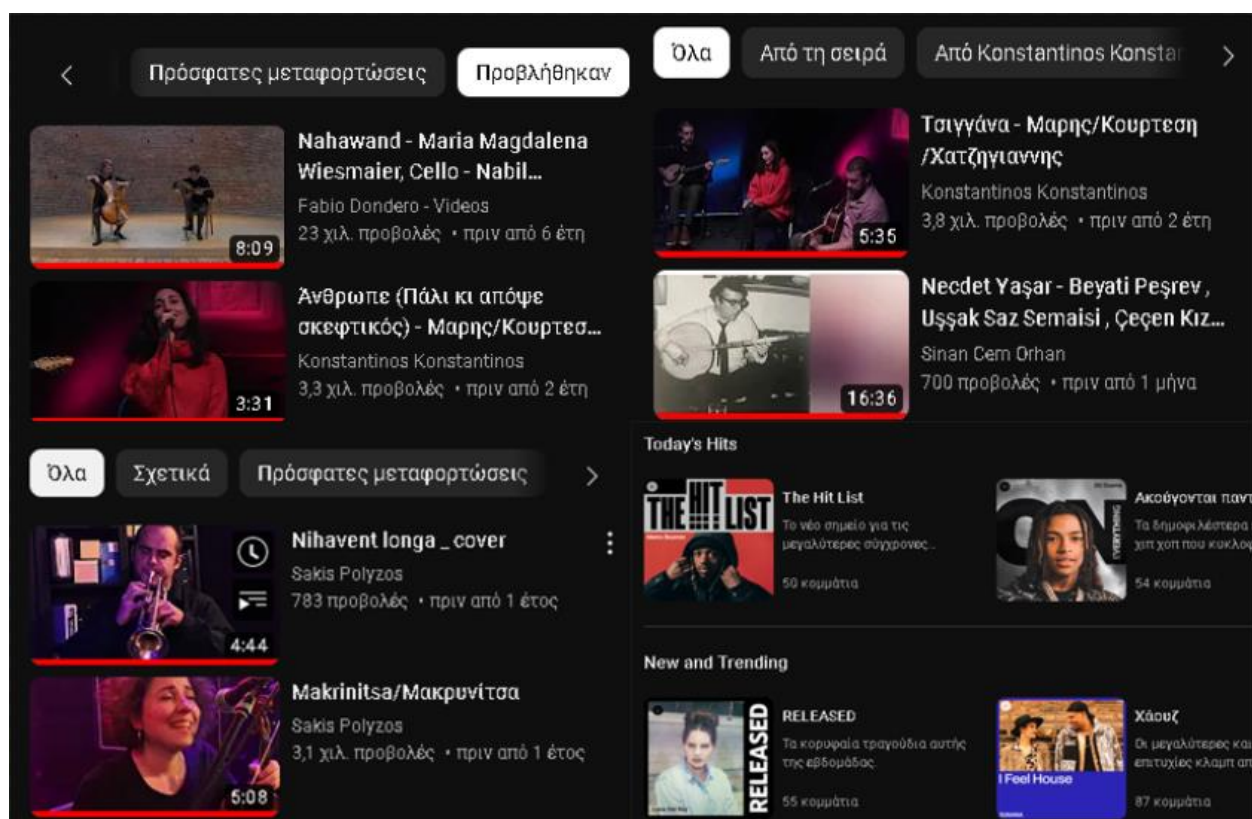
Το πρόβλημα της ακριβώς προηγούμενης προσέγγισης, το γεγονός δηλαδή πως εφόσον υπάρχουν στο σύστημα διαφορετικές κατηγορίες αντικειμένων, για τον κάθε χρήστη δεν υφίσταται ένας αλγόριθμος ο οποίος θα παράγει πάντοτε την καλύτερη πρόταση, αλλά ο αλγόριθμος αυτός αλλάζει ανάλογα με την κατηγορία αντικειμένων, αντιμετωπίζεται με την εισαγωγή κάποιου κριτηρίου εναλλαγής του αλγορίθμου που χρησιμοποιείται.

Στην περίπτωση του συστήματος του DailyLearner [33], το οποίο χρησιμοποιεί έναν CF και έναν CBF αλγόριθμο, οι αλγόριθμοι εναλλάσσονται μεταξύ τους όταν εκείνος που χρησιμοποιείται αποτυγχάνει να δημιουργήσει μια χρήσιμη για την χρήστη πρόταση. Με αυτήν την μέθοδο, παράγονται προτάσεις οι οποίες μπορεί να μην σχετίζονται με τις προηγούμενες, αλλά εξακολουθούν να είναι σχετικές με τον χρήστη. Η τεχνική αυτή γίνεται να συνδυαστεί και με την προηγούμενη, ώστε η συνοχή των αξιολογήσεων των προτάσεων από τον χρήστη για τον κάθε αλγόριθμο να καθορίζει το ποιος αλγόριθμος χρησιμοποιείται κάθε φορά.

Η μέθοδος αυτή εισάγει έναν νέο βαθμό πολυπλοκότητας στο σύστημα, καθώς τα κριτήρια εναλλαγής πρέπει να καθοριστούν, ωστόσο είναι πολύ πιο ευέλικτο και προσαρμόζεται στα πλεονεκτήματα και τα μειονεκτήματα του κάθε αλγορίθμου.

5.2.3 Παρουσίαση Μικτών Προτάσεων

Η μέθοδος αυτή χρησιμοποιείται σε περιπτώσεις όπου πρέπει να γίνουν στον χρήστη πάρα πολλές προτάσεις ταυτόχρονα. Τέτοιες είναι συνήθως οι πλατφόρμες πολυμέσων, όπου προτείνεται περιεχόμενο από πολλές πηγές ταυτόχρονα, και συχνά μάλιστα οι προτάσεις συνοδεύονται από μια επεξήγηση για το πως σχηματίστηκαν.



Εικόνα 4. Παράδειγμα Μικτών Προτάσεων στην Πλατφόρμα "Youtube"

Στην παραπάνω εικόνα, φαίνεται για παράδειγμα πως δίνεται μάλιστα στον χρήστη η επιλογή να εξετάσει ορισμένες από τις προτάσεις ανάλογα με το πως εκείνες παράχθηκαν.

Ένα πλεονέκτημα της μεθόδου αυτής, είναι το ζήτημα της συνεχούς προσθήκης νέων αντικειμένων σε ένα σύστημα το οποίο βολεύει να χρησιμοποιεί CF μεθόδους, καθώς μπορούνε

να υπάρχουν προτάσεις για αντικείμενα τα οποία δεν έχουν αξιολογηθεί. Όπως και η προηγούμενες μέθοδοι, προσφέρει μια ευελιξία, που η χρήση πιο «αγνών» αλγορίθμων δεν προσφέρει.

5.3 ΕΝΣΩΜΑΤΩΣΗ ΧΑΡΑΚΤΗΡΙΣΤΙΚΩΝ – ΠΡΑΓΜΑΤΙΚΑ ΥΒΡΙΔΙΑ

Μια κατηγορία προσεγγίσεων η οποία συνδυάζει χαρακτηριστικά CF και CBF.

Το κοινό στοιχείο ανάμεσα σε όσες υλοποιήσεις εμπίπτουν σε αυτήν την κατηγορία είναι το ότι τα χαρακτηριστικά λήψης της απόφασης είναι παρμένα και από τις δύο μεθόδους , ή και από περισσότερες.

Τα συστήματα αυτά συχνά υποστηρίζονται από βοηθητικά συστήματα που βασίζονται σε υπαρκτές γνώσεις και πληροφορίες, τα οποία βοηθάνε στο να ξεπεραστούν διάφορες δυσκολίες.

5.3.1 Ενσωμάτωση CF χαρακτηριστικών σε CBF Υλοποιήσεις

Στην περίπτωση αυτή, τα αποτελέσματα του CF αλγορίθμου λαμβάνονται απλά σαν ένα επιπλέον χαρακτηριστικό στο προφίλ $X(c)$ του Χρήστη c . Έπειτα, τα δεδομένα φιλτράρονται και παράγονται προτάσεις για τον χρήστη.

Εναλλακτικά, μπορεί να χρησιμοποιηθεί κάποιος αλγόριθμος SVD (Εξίσωση 17) πάνω σε ομάδες από προφίλ χρηστών, ώστε να υπάρχει μια κατηγοριοποίηση των χρηστών σε κλάσεις βάσει των προφίλ τους.

Οι αλγόριθμοι αυτοί έχουν σαν πλεονέκτημα το να διαφαίνεται η ομοιότητα μεταξύ αντικειμένων, κάτι που δεν συμβαίνει στις καθαρά CBF υλοποιήσεις.

5.3.2 Ενσωμάτωση CBF χαρακτηριστικών σε CF Υλοποιήσεις

Πολλοί αλγόριθμοι λειτουργούν με αυτόν τον τρόπο, κρατώντας δηλαδή τον CF χαρακτήρα τους, αλλά χρησιμοποιώντας την έννοια του προφίλ ($X(c)=(x_{c,1},x_{c,2},\dots,x_{c,n})$) του χρήστη ώστε να συγκριθούν δύο χρήστες και να θεωρηθούν παρόμοιοι. Αντί δηλαδή οι χρήστες να συγκρίνονται μόνο βάση των κοινών αντικειμένων που έχουν αξιολογήσει ($S_{cc'}$), συγκρίνονται σε σχέση με το πόσο μοιάζουν τα προφίλ (διανύσματα) τους, μέσω μια συνάρτησης sim σαν αυτές που έχουμε ήδη συναντήσει.

Αυτή η τροποποίηση βοηθάει να ξεπεραστεί ένα πρόβλημα το οποίο συναντάται εύκολα σε συστήματα που διαθέτουν πολλά αντικείμενα, δηλαδή την έλλειψη ή τον ανεπαρκή αριθμό αντικειμένων που ανήκουν στο σύνολο $S_{cc'}$.

5.3.3 Αλληλουχίες Αλγορίθμων

Σε αυτήν την τεχνική, χρησιμοποιείται αρχικά ένας αλγόριθμος ώστε να «φιλτράρει» τα αρχικά δεδομένα σε ένα πρωταρχικό επίπεδο, και ύστερα άλλος ένας ο οποίος θα επιλέξει μεταξύ των υποψηφίων για το ποιες θα είναι οι τελικές προτάσεις. Δε αποκλείεται να υπάρχουν και επιπλέον βήματα διαλογής σε τέτοιου τύπου υλοποιήσεις.

Η χρήση αλληλουχιών βοηθάει το σύστημα να αποφεύγει να χρησιμοποιεί αλγορίθμους που απαιτούν υπολογιστική ισχύ ώστε να αξιολογήσει αντικείμενα που είναι ήδη είτε πολύ καλά αξιολογημένα από τον πρώτο αλγόριθμο, είτε τόσο χαμηλά που δεν θα μπορούσαν να χρησιμοποιηθούν σαν προτάσεις, και να εξαλείφει τον θόρυβο που μπορεί να επηρεάσει τα αποτελέσματα του δεύτερου αλγορίθμου.

5.4 ΣΧΟΛΙΑ – ΑΔΥΝΑΜΙΕΣ – ΣΥΜΠΕΡΑΣΜΑΤΑ

Σε όλη την βιβλιογραφία η οποία μελετήθηκε για την εκπόνηση της παρούσας εργασίας, ήταν εμφανές πως οι υβριδικές υλοποιήσεις είναι εκείνες που εν τέλει χρησιμοποιούνται στην πράξη. Κάτι τέτοιο είναι αναμενόμενο, καθώς σε όλους τους αλγόριθμους που περιγράψαμε, προκύπτουν διάφορα προβλήματα. Τα προβλήματα αυτά δεν σχετίζονται βέβαια μονάχα με τους αλγορίθμους, αλλά και με το κάθε μοναδικό σύστημα και τα εγγενή προβλήματα του.

Οι υβριδικές υλοποιήσεις, οι οποίες είδαμε πως μπορεί να διαφέρουν πάρα πολύ μεταξύ τους και σίγουρα δεν μπορούν να συνοψιστούν επαρκώς σε μια οποιαδήποτε εργασία, δύνανται

να καλύπτουν τα συγκεκριμένα κενά και να επιλύουν τα προβλήματα που προκύπτουν. Επί της ουσίας, το στήσιμο ενός συστήματος συστάσεων για τον μηχανικό είναι σχεδόν σίγουρο πως θα περιλαμβάνει τον συνδυασμό τεχνικών ή έστω την χρήση συμπληρωματικών συστημάτων σε κάποιο βαθμό.

Οι αλγόριθμοι αυτοί, είναι βεβαίως αναμενόμενο να είναι πιο απαιτητικοί στον σχεδιασμό τους, και να απαιτούν σε υψηλό βαθμό παραμετροποίηση για το κάθε ξεχωριστό σύστημα.

6. ΣΥΣΤΗΜΑ ΣΥΣΤΑΣΕΩΝ ΜΟΥΣΙΚΗΣ ΜΕ ΤΗΝ ΧΡΗΣΗ ΕΠΑΝΑΛΑΜΒΑΝΟΜΕΝΟΥ ΝΕΥΡΩΝΙΚΟΥ ΔΙΚΤΥΟΥ ΣΥΝΕΛΙΞΗΣ

Η Μηχανική Μάθηση βρίσκεται στις μέρες μας εφαρμογή σε πολλούς διαφορετικούς τομείς της τεχνολογίας και της έρευνας. Τα συστήματα συστάσεων αποτελούν συστήματα τα οποία επεξεργάζονται μεγάλο όγκο δεδομένων, και είναι αναμενόμενο να χρησιμοποιούνται αλγόριθμοι που παράγουν μοντέλα μέσω μηχανικής μάθησης, καθώς αυτή αποτελεί συνήθως τη βέλτιστη λύση. [34]

Στα προηγούμενα κεφάλαια αναλύθηκαν οι διαφορετικοί αλγόριθμοι των συστημάτων συστάσεων. Εισαγωγικά στοιχεία και πληροφορίες για τις έννοιες της μηχανικής μάθησης, της τεχνητής νοημοσύνης και της βαθιάς μηχανικής μάθησης υπάρχουνε στην εισαγωγή της εργασίας.

Σε αυτήν και την επόμενη ενότητα παρουσιάζονται κάποιες εφαρμογές της μηχανικής μάθησης εμφωλευμένες σε συστήματα συστάσεων. Ουσιαστικά, στην εργασία υπάρχει ήδη αναφορά στην χρήση μηχανικής μάθησης σε συστήματα συστάσεων, στην ενότητα 1.1.3, όμως καθώς οι model based υλοποιήσεις αναφέρονταν στην βιβλιογραφία ως μια υποκατηγορία των CF συστημάτων, θεωρήθηκε πως πρέπει για τους σκοπούς της εργασίας να αφιερωθεί μια ξεχωριστή ενότητα στην παρουσίαση εφαρμογών μηχανικής μάθησης σε συστήματα συστάσεων. Στην πράξη, υπάρχουν πάρα πολλές τέτοιες περιπτώσεις, καθώς τα συστήματα μηχανικής μάθησης είναι πολύ χρήσιμα στην επεξεργασία δεδομένων. Για την παρούσα εργασία επιλέχθηκαν κάποια χαρακτηριστικά παραδείγματα[10], [35]–[38].

Η μελέτη των Adiyansjah, Alexander A S Gunawan, και Derwin Suhartono αποσκοπεί στην ανάπτυξη ενός συστήματος συστάσεων το οποίο θα μπορεί να παράγει συστάσεις μουσικής που θα βασίζονται στην ομοιότητα χαρακτηριστικών ηχητικών σημάτων [39].

6.1 ΕΠΑΝΑΛΑΜΒΑΝΟΜΕΝΑ ΝΕΥΡΩΝΙΚΑ ΔΙΚΤΥΑ ΣΥΝΕΛΙΞΗΣ

Τα Επαναλαμβανόμενα Νευρωνικά Δίκτυα Συνέλιξης (Convolutional Recurrent Neural Networks – CRNNs) αποτελούν έναν συνδυασμό των Νευρωνικών Δικτύων Συνέλιξης (Convolutional Neural Networks – CNNs) και των Επαναλαμβανόμενων Νευρωνικών Δικτύων (Recurrent Neural Networks – RNNs)[40]. Τα CNNs είναι δίκτυα που αποτελούνται από πολλά επίπεδα, με τα βασικά να είναι τα επίπεδα συνέλιξης, όπου απομονώνονται και αναγνωρίζονται μοτίβα και χαρακτηριστικά, συνήθως σε εικόνες. Στην περίπτωση που το σύστημα αναλύει ηχητικά κύματα, είναι απαραίτητη η μετατροπή των αρχείων ήχου σε Διαγράμματα Συχνοτήτων μέσω κάποιου μετασχηματισμού Fourier (στην συγκεκριμένη περίπτωση χρησιμοποιείται ο μετασχηματισμός Short Time Fourier Transform – STFT), οπότε αναλύονται τα συνολικά συχνотικά χαρακτηριστικά κάθε αρχείου. Τα RNNs έχουν σχεδιαστεί για να λειτουργούν με δεδομένα χρονοσειρών, και είναι ιδιαίτερα χρήσιμα στην επίλυση προβλημάτων πρόβλεψης κάποιας ακολουθίας, οπότε στην περίπτωση που αναλύουν αρχεία ήχου ανιχνεύουν διάφορα μοτίβα που βρίσκονται μέσα στο αρχείο. Τα CRNNs αποτελούν συνδυασμό των δύο προαναφερθέντων αρχιτεκτονικών, και επιτυγχάνουν υψηλής ακρίβειας αποτελέσματα σε προβλήματα κατάταξης αρχείων ήχου σε συγκεκριμένες κατηγορίες (μουσικά είδη).

6.2 ΠΕΡΙΓΡΑΦΗ ΣΥΣΤΗΜΑΤΟΣ

Στην μελέτη που αναλύεται, χρησιμοποιήθηκε ένα ήδη υπάρχον σύνολο δεδομένων το οποίο αρχικά περιείχε 25.000 ηχητικά δείγματα διάρκειας 30 δευτερολέπτων έκαστο, τα οποία αντιστοιχούσαν σε ένα ή περισσότερα είδη μουσικής. Η τελική διαλογή περιείχε εν τέλει 7000 δείγματα, ισομοιρασμένα σε 7 διαφορετικά μουσικά είδη, με κάθε τραγούδι να αντιστοιχεί σε ένα μόνο είδος και διασφάλιση υψηλής ποιότητας για κάθε δείγμα. Όλα τα αρχεία ήχου μετατράπηκαν έπειτα σε διαγράμματα συχνотήτων στην κλίμακα Mel. Τα ίδια δεδομένα επεξεργάστηκαν με αλγόριθμο CNN και αλγόριθμο CRNN ξεχωριστά.

Οι αλγόριθμοι παρήγαγαν εν τέλει για κάθε αρχείο ένα διάνυσμα χαρακτηριστικών, δίνοντας την δυνατότητα να γίνονται συγκρίσεις ανάμεσα σε τραγούδια μέσω συναρτήσεων ομοιότητας παρόμοιες με αυτές που περιγράψαμε στην ενότητα 1.

Για να παραχθούν συστάσεις χρησιμοποιήθηκαν έπειτα δύο διαφορετικοί αλγόριθμοι. Ο πρώτος λάμβανε υπόψιν του μόνο την ομοιότητα μεταξύ των αρχείων που είχε επιλέξει ο χρήστης, ενώ ο δεύτερος λάμβανε επίσης υπόψιν και το είδος μουσικής στο οποίο ήταν καταχωρισμένο. Τα αποτελέσματα αυτών των αλγορίθμων αξιολογήθηκαν ξεχωριστά.

6.3 ΑΞΙΟΛΟΓΗΣΗ ΑΠΟΤΕΛΕΣΜΑΤΩΝ

Τα αποτελέσματα της αξιολόγησης ήταν θετικά υπέρ των αλγορίθμων που λάμβαναν υπόψιν και το μουσικό είδος στο οποίο κατατασσότανε κάθε τραγούδι. Αναφορικά με τους αλγόριθμους εξαγωγής χαρακτηριστικών, ο αλγόριθμος CRNN που εξήγαγε τα χαρακτηριστικά

από τα τραγούδια, λαμβάνοντας υπόψιν του τόσο τα χαρακτηριστικά των συχνοτήτων που εμπεριεχόταν στο κομμάτι όσο και τα μοτίβα που εμπεριέχονταν στην χρονική ακολουθία του τραγουδιού, ήτανε πιο αποτελεσματικός από τον αλγόριθμο CNN που λάμβανε υπόψιν του μόνο τα χαρακτηριστικά των συχνοτήτων.

7. ΣΥΣΤΗΜΑ ΣΥΣΤΑΣΕΩΝ ΕΙΔΗΣΕΩΝ ΒΑΣΙΣΜΕΝΟ ΣΕ ΜΕΜΟΝΩΜΕΝΕΣ ΣΥΝΕΔΡΙΕΣ ΜΕ ΧΡΗΣΗ ΒΑΘΙΩΝ ΝΕΥΡΩΝΙΚΩΝ ΔΙΚΤΥΩΝ

Σε αυτήν την μελέτη, προτείνεται από τους Gabriel de Souza Pereira Moreira και Felipe Ferreira μια υλοποίηση της αρχιτεκτονικής CHAMELEON για το πρόβλημα της δημιουργίας προτάσεων σε χρήστες ειδησεογραφικών ιστοσελίδων κατά την επίσκεψη τους σε αυτές [41].

7.1 ΣΥΝΑΡΤΗΣΕΙΣ ΠΟΥ ΧΡΗΣΙΜΟΠΟΙΗΘΗΚΑΝ

Σε αυτήν την ενότητα θα παρουσιαστούν κάποιες συναρτήσεις που αποτελούν τμήματα του αλγορίθμου.

7.1.1 Η συνάρτηση softmax

Η συνάρτηση softmax αποτελεί μια συνάρτηση κανονικοποίησης των τιμών των κόμβων εξόδου ενός νευρωνικού δικτύου, ώστε εκείνες να αναπαρίστανται σε μορφή πιθανότητας. Υπάρχουν διάφορες συναρτήσεις που μετατρέπουν το αποτέλεσμα των κόμβων αυτών, όπως η argmax η οποία δίνει επί της ουσίας το αποτέλεσμα του αλγορίθμου, ωστόσο οι τιμές που δίνει η softmax μπορούν να χρησιμοποιηθούν πιο εύκολα για την αξιολόγηση της ακρίβειας των αποτελεσμάτων του δικτύου.

$$\sigma(x_j) = \frac{e^{x_j}}{\sum_i e^{x_i}} \text{ (Εξίσωση 23)}$$

Όπου x_i η αρχική τιμή του κάθε κόμβου i , και j ο εξεταζόμενος κόμβος.

Θα ισχύει πάντοτε πως $\sum_i \sigma(x_i) = 1$.

7.1.2 Η λογαριθμική απώλεια διασταυρούμενης εντροπίας (cross-entropy log loss)

Η συνάρτηση αυτή αξιοποιεί τα αποτελέσματα της συνάρτησης softmax ώστε να βελτιστοποιηθούν οι παράμετροι του δικτύου.

$$l(i) = -a_i \log(\sigma(x_i)) \text{ (Εξίσωση 23)}$$

Όπου η τιμή a_i αντιπροσωπεύει την πραγματική τιμή του κόμβου την οποία έπρεπε να προβλέψει το μοντέλο, δηλαδή στην περίπτωση των δικτύων κατηγοριοποίησης που εξετάζουμε, 1 εάν το αρχικό διάνυσμα ανήκει στην κατηγορία που αντιπροσωπεύει ο συγκεκριμένος κόμβος εξόδου και 0 εάν αυτό δεν συμβαίνει, και το $l(i)$ αντιπροσωπεύει την λογαριθμική απώλεια.

7.3 ΠΕΡΙΓΡΑΦΗ ΣΥΣΤΗΜΑΤΟΣ

Σε αυτήν την ενότητα θα παρουσιαστεί συνολικά η λειτουργία του συστήματος. Το σύστημα χωρίζεται σε δύο κυρίως υποσυστήματα, το σύστημα αναπαράστασης άρθρου (Article Content Representation - ACR) και το σύστημα πρόβλεψης επόμενου άρθρου (Next-Article Recommendation - NAR).

7.3.1 Σύστημα αναπαράστασης άρθρου (ACR)

Η λειτουργία του υποσυστήματος είναι η ανάλυση του περιεχομένου των άρθρων και η κατηγοριοποίηση τους, με τις εξής μεταβλητές εκπαίδευσης:

1. Τα στοιχεία του άρθρου (εκδότης, ανατιθέμενη κατηγορία, κλπ.)
2. Το περιεχόμενο του άρθρου

Χρησιμοποιείται ένα νευρωνικό δίκτυο διαχωριστικής αρχιτεκτονικής, το οποίο κατηγοριοποιεί το κάθε άρθρο. Στο τελευταίο επίπεδο του δικτύου, χρησιμοποιείται η συνάρτηση softmax και έπειτα η συνάρτηση λογαριθμικής απώλειας, η οποία συγκρίνει το αποτέλεσμα του δικτύου με την ανατιθέμενη κατηγορία ώστε να εκπαιδευτεί το δίκτυο.

7.3.2 Σύστημα πρόβλεψης επόμενου άρθρου (NAR)

Η λειτουργία του υποσυστήματος είναι οι παροχή προτάσεων στον χρήστη, με τις εξής μεταβλητές για το κάθε άρθρο που εξετάζεται:

1. Η αναπαράσταση του άρθρου από το ήδη εκπαιδευμένο υποσύστημα ACR
2. Η δημοτικότητα και η ημερομηνία δημοσίευσης του άρθρου
3. Τα στοιχεία που παρέχονται από την δραστηριότητα του χρήστη (τοποθεσία, ώρα εισόδου στο σύστημα, συσκευή κλπ.)

Με τα στοιχεία αυτά, το σύστημα παράγει προτάσεις (*προσωποποιημένα άρθρα*) για τον χρήστη. Το σύστημα χρησιμοποιεί επίσης ένα δίκτυο RNN το οποίο μοντελοποιεί την ακολουθία των άρθρων που διαβάζουν οι χρήστες στις συνεδρίες τους, και για κάθε άρθρο που διαβάζει ο χρήστης εκπαιδεύεται να παράγει μια πρόταση για το επόμενο (*πρόβλεψη επόμενου άρθρου*).

Για να εκπαιδευτεί το σύστημα αυτό, τα προσωποποιημένα άρθρα και οι προβλέψεις επόμενου άρθρου συγκρίνονται με μια συνάρτηση ομοιότητας με τα άρθρα που διαβάζει όντως ο χρήστης (θετικά δείγματα), καθώς και με εκείνα που δεν διαβάζει (αρνητικά δείγματα). Στην πρώτη περίπτωση στόχος είναι η ομοιότητα να μεγιστοποιηθεί, ενώ στη δεύτερη να ελαχιστοποιηθεί.

Με την χρήση της στρατηγικής αυτής, ένα άρθρο μπορεί να προταθεί κατευθείαν στον χρήστη εφόσον περάσει από το υποσύστημα ACR, ενώ αν το σύστημα λειτουργούσε με Collaborative Filtering, θα υπήρχε πρόβλημα με τα νέα αντικείμενα.

7.4 ΑΞΙΟΛΟΓΗΣΗ ΑΠΟΤΕΛΕΣΜΑΤΩΝ

.Για την συνολική αξιολόγηση του συστήματος εκτελέστηκαν 2 πειράματα, τα οποία αξιολογήθηκαν με μετρήσεις HR@5 και MRR@5

1. Συνεχή εκπαίδευση και αξιολόγηση κάθε πέντε ώρες, για διάστημα δεκαπέντε ημερών.
2. Συνεχή εκπαίδευση και αξιολόγηση κάθε μία ώρα, για διάστημα μίας ημέρας.

Και στα δύο πειράματα, η υλοποίηση του συστήματος CHAMELEON πέτυχε μεγαλύτερη ακρίβεια και στις δύο μετρήσεις από τις μεθόδους που χρησιμοποιήθηκαν ως σημεία Τμήμα Μηχανολόγων και Αεροναυπηγών Μηχανικών – Τομέας Διοίκησης και Οργάνωσης 53

αναφοράς.. Το υποσύστημα ACR, κατάφερε πράγματι να κατηγοριοποιήσει σε μεγάλο ποσοστό τα άρθρα στις ανατιθέμενες κατηγορίες τους.

8. ΣΥΝΟΨΗ – ΣΥΜΠΕΡΑΣΜΑΤΑ

Σε αυτήν την Σπουδαστική Εργασία έγινε μια παρουσίαση των πιο βασικών κατηγοριών Συστημάτων Συστάσεων, και κάποιων περιπτώσεων χρήσης Μηχανικής Μάθησης μέσα σε αυτά. Ο συγκεκριμένος τομέας έχει πληθώρα εφαρμογών, και οι αλγόριθμοι που παρουσιάστηκαν μελετώνται από πριν ακόμα περάσουμε σαν ανθρωπότητα ολοκληρωτικά στην ψηφιακή εποχή. Είναι λοιπόν σαφές πως η εργασία αυτή δεν αποτελεί μια εξαντλητική παρουσίαση όλων των υποπεριπτώσεων ή του πλήρους ιστορικού ανάπτυξης αυτών των αλγορίθμων, καθώς κάτι τέτοιο θα ξέφευγε κατά πολύ από τα πλαίσια της σπουδαστικής εργασίας.

Η βιβλιογραφία ήταν πλούσια, και μέσα από την ανάλυση των κειμένων που ενσωματώθηκαν στην εργασία σαν βιβλιογραφικές αναφορές, αλλά και άλλων που εν τέλει δεν επιλέχθηκαν για αυτόν τον σκοπό, προέκυψε μια σκιαγράφηση των χαρακτηριστικών του συγκεκριμένου τομέα της πληροφορικής, του πλαισίου στο οποίο τίθενται τα σχετικά προβλήματα τα οποία επιλύονται με την χρήση των συστημάτων συστάσεων, και τα συνήθη προβλήματα που προκύπτουν με την χρήση πραγματικών δεδομένων.

Στο πρώτο μέρος της εργασίας, παρουσιάζονται οι πιο ευρέως χρησιμοποιημένες κατηγορίες αλγορίθμων συστάσεων (recommendation algorithms). Οι δύο σημαντικότερες κατηγορίες είναι το συλλογικό φιλτράρισμα (collaborative filtering), και το φιλτράρισμα βάση των χαρακτηριστικών των αντικειμένων (content based filtering).

- Στο συλλογικό φιλτράρισμα, οι συστάσεις παράγονται μέσω της εξέτασης της ομοιότητας του χρήστη με άλλους χρήστες και τις προτιμήσεις αυτών. Υπάρχουν διαφορετικές προσεγγίσεις αυτών των αλγορίθμων. Άλλες εξετάζουν την εύρεση των χρηστών που είναι παρόμοιοι με τον χρήστη για τον οποίον παράγονται οι συστάσεις (user based), ενώ άλλες την εύρεση των αντικειμένων που είναι παρόμοια μεταξύ του βάσης των αξιολογήσεων των χρηστών (item based). Υπάρχουν επίσης οι model based προσεγγίσεις, όπου διαμορφώνονται μοντέλα μέσω μηχανικής μάθησης για την πρόβλεψη των αξιολογήσεων. Στην συγκεκριμένη εργασία, παρουσιάσαμε την περίπτωση της χρήσης κλάσεων κατηγοριοποίησης των χρηστών ανάλογα με τις αξιολογήσεις τους, ή την εκπαίδευση δικτύων που θα προβλέπουν την αξιολόγηση των χρηστών για αντικείμενα.
- Στο φιλτράρισμα βάσης των χαρακτηριστικών των αντικειμένων, η εκτιμώμενη αξιολόγηση ενός αντικειμένου εξάγεται κατά κανόνα βάσης των αξιολογήσεων του χρήστη για αντικείμενα που θεωρούνται παρόμοια με αυτό. Σε αυτές τις περιπτώσεις, υφίστανται διάφοροι τρόποι εξαγωγής των χαρακτηριστικών του αντικειμένου. Στην εργασία αναλύεται συνολικά η διαδικασία της παραγωγής των συστάσεων, η οποία αποτελείται από την ανάλυση του περιεχομένου των αντικειμένων, την δημιουργία του προφίλ του κάθε χρήστη, βάσης των προτιμήσεων του και άλλων πληροφοριών που εξάγονται έμμεσα, και την δημιουργία των συστάσεων μέσα από την σύγκριση του προφίλ με τα αντικείμενα. Σε αυτήν την κατηγορία, όπως και στην προηγούμενη, υπάρχουν επίσης model based προσεγγίσεις.

Σε κάθε μία από αυτές τις κατηγορίες, εξετάζονται τόσο οι έμμεσες όσο και οι άμεσες αξιολογήσεις του χρήστη για τα αντικείμενα. Αρκετά σημαντική είναι επίσης η χρήση των συναρτήσεων ομοιότητας. Αρκετές από αυτές παρουσιάζονται στην παρούσα εργασία.

Εκτός από τις κύριες αυτές κατηγορίες, υφίστανται και συμπληρωματικές κατηγορίες αλγορίθμων. Αυτές που παρουσιάζονται είναι:

- Η χρήση στερεοτύπων, δηλαδή η κατηγοριοποίηση των χρηστών σε ήδη διαμορφωμένες κλάσεις μέσω κάποιων πολύ βασικών χαρακτηριστικών τους που εξάγονται έμμεσα (πχ ηλικία, τόπος διαμονής)
- Η χρήση καθολικών κριτηρίων για την παρούσα δημοτικότητα (trend) ενός αντικειμένου
- Η δημιουργία συστάσεων βάση του πόσο συχνά δύο αντικείμενα αξιολογούνται θετικά από έναν χρήστη σε αλληλουχία. Αυτή η μέθοδος (Co-Occurrence) συνήθως εφαρμόζεται σε e-shop όταν δύο αντικείμενα αγοράζονται συχνά μαζί

Τέλος, μια ενότητα αφιερώνεται στην περιγραφή των υβριδικών μεθόδων, που είναι και η συχνότερη περίπτωση που συναντάται στην πράξη. Οι υβριδικές υλοποιήσεις είναι πάρα πολλές και δεν είναι εφικτό να περιγράψουν συνολικά σε μια ακαδημαϊκή μελέτη, ωστόσο στην παρούσα εργασία γίνεται μια κατηγοριοποίηση αυτών. Οι κατηγορίες που αναλύονται είναι η χρήση συνδυαστικών μοντέλων και ο συγκερασμός των προτάσεων, με την χρήση συντελεστών βαρύτητας ή χωρίς, ή με την εναλλαγή των αλγορίθμων που χρησιμοποιούνται, η παρουσίαση προτάσεων από διαφορετικούς αλγορίθμους ταυτόχρονα, καθώς και η ενσωμάτωση χαρακτηριστικών της μιας μεθόδου στην άλλη.

Οι υβριδικές προσεγγίσεις είναι εκείνες οι οποίες επιτυγχάνουν στην πραγματικότητα να προσαρμόζονται σε κάθε ξεχωριστό σύστημα το οποίο έχει τα δικά του ιδιαίτερα προβλήματα και συνθήκες λειτουργίας.

Τα κύρια προβλήματα που παρουσιάζονται στα συστήματα συστάσεων έχουν συνήθως να κάνουν με την αραιότητα δεδομένων, την προσθήκη νέων αντικειμένων ή χρηστών στο σύστημα, και την έλλειψη αξιολογήσεων. Η χρήση μηχανικής μάθησης και μοντέλων, καθώς και η χρήση συμπληρωματικών κριτηρίων και μεθόδων τα οποία επίσης εντάχθηκαν στην εργασία αποτελούν τις πιο συνήθεις λύσεις των προβλημάτων αυτών.

Στο δεύτερο κεφάλαιο, δίνεται βαρύτητα στην παρουσίαση της έννοιας της Μηχανικής Μάθησης (Machine Learning), και της χρήσης αυτής στα συστήματα συστάσεων. Η Μηχανική Μάθηση είναι μια ευρεία έννοια, η οποία περιλαμβάνει διαφορετικές κατηγορίες συστημάτων (πχ επιβλεπόμενα, μη επιβλεπόμενα, ημί-επιβλεπόμενα), και έχει πληθώρα εφαρμογών σε πολλούς τεχνολογικούς τομείς. Στον τομέα των συστάσεων, χρησιμοποιούνται κυρίως νευρωνικά δίκτυα βαθιάς ή μη μηχανικής μάθησης (Recurrent Neural Networks, Convolutional Neural Networks, κ.α.).

Κατά την διάρκεια της ανασκόπησης της βιβλιογραφίας, μελετήθηκαν διάφορα συστήματα τα οποία βρίσκονται σε εκείνη την κατηγορία, και πολλές πληροφορίες που βρίσκονται στην παρούσα εργασία εξάχθηκαν από αυτά. Επιλέχθηκε δύο από αυτά να περιγραφούνε στο εν λόγω κεφάλαιο.

- Ένα σύστημα συστάσεων μουσικής το οποίο χρησιμοποιεί επαναλαμβανόμενα νευρωνικά δίκτυα συνέλιξης, και

- Ένα σύστημα συστάσεων ειδήσεων το οποίο εξετάζει μεμονωμένες συνεδρίες χρηστών σε ειδησεογραφικές ιστοσελίδες, και χρησιμοποιεί βαθιά νευρωνικά δίκτυα.

Οι τεχνολογίες που παρουσιάζονται στην παρούσα εργασία βρίσκονται ακόμη σε ανάπτυξη και έχουν ευρεία χρήση σε πληθώρα από πλατφόρμες και υπηρεσίες. Όπως έχει ήδη αναφερθεί, η βιβλιογραφία ήτανε πλούσια και υπήρχαν αρκετές πηγές που ανέλυαν τις έννοιες που παρουσιάζονται στην εργασία. Ο τομέας των συστάσεων είναι λογικό να αναπτύσσεται καθώς υπάρχει ανταγωνισμός που προκύπτει λόγω του πλεονεκτήματος που παρέχει η ύπαρξη καλών συστάσεων για έναν χρήστη σε πολλές διαφορετικές περιπτώσεις. Η συμβολή των μηχανικών στα συστήματα συστάσεων ήταν εμφανής κατά την ανασκόπηση της βιβλιογραφίας, καθώς κοινό ζητούμενο είναι η βελτιστοποίηση των αλγορίθμων αυτών.

ΒΙΒΛΙΟΓΡΑΦΙΑ

- [1] G. Adomavicius and A. Tuzhilin, “Toward the Next Generation of Recommender Systems: A Survey of the State-of-the-Art and Possible Extensions.”
- [2] S. Greengard, “Advertising Gets Personal,” *Commun ACM*, vol. 55, no. 8, pp. 18–20, Aug. 2012, doi: 10.1145/2240236.2240243.
- [3] “Recommender Systems.”
- [4] W. Carrer-Neto, M. L. Hernández-Alcaraz, R. Valencia-García, and F. García-Sánchez, “Social knowledge-based recommender system. Application to the movies domain,” *Expert Syst Appl*, vol. 39, no. 12, pp. 10990–11000, Sep. 2012, doi: 10.1016/j.eswa.2012.03.025.
- [5] *Recommender Systems Handbook*. Springer US, 2011. doi: 10.1007/978-0-387-85820-3.
- [6] I. Portugal, P. Alencar, and D. Cowan, “The use of machine learning algorithms in recommender systems: A systematic review,” *Expert Systems with Applications*, vol. 97. Elsevier Ltd, pp. 205–227, May 01, 2018. doi: 10.1016/j.eswa.2017.12.020.
- [7] A. Da’u and N. Salim, “Recommendation system based on deep learning methods: a systematic review and new directions,” *Artif Intell Rev*, vol. 53, no. 4, pp. 2709–2748, Apr. 2020, doi: 10.1007/s10462-019-09744-1.
- [8] A. Singhal, P. Sinha, and R. Pant, “Use of Deep Learning in Modern Recommendation System: A Summary of Recent Works,” 2017.
- [9] H. Wang, N. Wang, and D. Y. Yeung, “Collaborative deep learning for recommender systems,” in *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, Association for Computing Machinery, Aug. 2015, pp. 1235–1244. doi: 10.1145/2783258.2783273.

- [10] L. Zheng, V. Noroozi, and P. S. Yu, “Joint deep modeling of users and items using reviews for recommendation,” in *WSDM 2017 - Proceedings of the 10th ACM International Conference on Web Search and Data Mining*, Association for Computing Machinery, Inc, Feb. 2017, pp. 425–433. doi: 10.1145/3018661.3018665.
- [11] M. Khalaji, C. Dadkhah, and J. Gharibshah, “Hybrid Movie Recommender System based on Resource Allocation.”
- [12] D. Billsus and M. J. Pazzani, “Learning Collaborative Information Filters.” [Online]. Available: www.firefly.com
- [13] B. Marlin, “Modeling User Rating Profiles For Collaborative Filtering.”
- [14] L. Chen, G. Chen, and F. Wang, “Recommender systems based on user reviews: the state of the art,” *User Model User-adapt Interact*, vol. 25, no. 2, pp. 99–154, Apr. 2015, doi: 10.1007/s11257-015-9155-5.
- [15] J. S. Breese, D. Heckerman, and C. Kadie, “Empirical Analysis of Predictive Algorithms for Collaborative Filtering.”
- [16] M. Pazzani, D. Billsus, R. S. Michalski, and J. Wnek, “Learning and Revising User Profiles: The Identification of Interesting Web Sites,” Kluwer Academic Publishers, 1997. [Online]. Available: <http://ai.iit.nrc.ca/subjects/ai>
- [17] T. Hofmann, “Collaborative Filtering via Gaussian Probabilistic Latent Semantic Analysis,” 2003. [Online]. Available: <http://doi.acm.org>
- [18] Y. Cao, R. Klamka, and M. C. Pham, “Clustering Technique for Collaborative Filtering and the Application to Venue Recommendation MILKI-PSY-Multimodal Immersive Learning with AI for Psychomotoric Skills View project Clustering Technique for

- Collaborative Filtering and the Application to Venue Recommendation,” 2010. [Online]. Available: <http://www.cs.wisc.edu/dbworld/>
- [19] D. Billsus and M. J. Pazzani, “Learning Collaborative Information Filters.” [Online]. Available: www.firefly.com
- [20] H.-J. Xue, X.-Y. Dai, J. Zhang, S. Huang, and J. Chen, “Deep Matrix Factorization Models for Recommender Systems *,” 2017.
- [21] L. Wan, F. Xia, X. Kong, C. H. Hsu, R. Huang, and J. Ma, “Deep Matrix Factorization for Trust-Aware Recommendation in Social Networks,” *IEEE Trans Netw Sci Eng*, vol. 8, no. 1, pp. 511–528, Jan. 2021, doi: 10.1109/TNSE.2020.3044035.
- [22] D. Billsus and M. J. Pazzani, “Learning Collaborative Information Filters.” [Online]. Available: www.firefly.com
- [23] W. H. Hsu, P. Boddhireddy, and R. Joehanes, “Using Probabilistic Relational Models for Collaborative Filtering.” [Online]. Available: <http://www.kddresearch.org>
- [24] *Recommender Systems Handbook*. Springer US, 2011. doi: 10.1007/978-0-387-85820-3.
- [25] A. Agostinelli *et al.*, “MusicLM: Generating Music From Text,” Jan. 2023, [Online]. Available: <http://arxiv.org/abs/2301.11325>
- [26] J. Beel, B. Gipp, S. Langer, and C. Breiting, “Research-paper recommender systems: a literature survey,” *International Journal on Digital Libraries*, vol. 17, no. 4, pp. 305–338, Nov. 2016, doi: 10.1007/s00799-015-0156-0.
- [27] L. Shah, R. Scholar, H. Gaudani, and P. Balani, “Survey on Recommendation System,” 2016.
- [28] M. F. Triola, “Bayes’ Theorem.”

- [29] R. J. Mooney, “Content-Based Book Recommending Using Learning for Text Categorization,” 2000.
- [30] E. Rich, “User Modeling via Stereotypes*,” 1979.
- [31] X. Wang and Y. Wang, “Improving content-based and hybrid music recommendation using deep learning,” in *MM 2014 - Proceedings of the 2014 ACM Conference on Multimedia*, Association for Computing Machinery, Nov. 2014, pp. 627–636. doi: 10.1145/2647868.2654940.
- [32] R. Burke, “Hybrid Recommender Systems: Survey and Experiments 1.” [Online]. Available: <http://www.google.com>
- [33] D. Billsus and M. J. Pazzani, “User Modeling for Adaptive News Access.”
- [34] A. Da’u and N. Salim, “Recommendation system based on deep learning methods: a systematic review and new directions,” *Artif Intell Rev*, vol. 53, no. 4, pp. 2709–2748, Apr. 2020, doi: 10.1007/s10462-019-09744-1.
- [35] A. Singhal, P. Sinha, and R. Pant, “Use of Deep Learning in Modern Recommendation System: A Summary of Recent Works,” 2017.
- [36] T. Hofmann, “Collaborative Filtering via Gaussian Probabilistic Latent Semantic Analysis,” 2003. [Online]. Available: <http://doi.acm.org>
- [37] Y. Cao, R. Klamka, and M. C. Pham, “Clustering Technique for Collaborative Filtering and the Application to Venue Recommendation MILKI-PSY-Multimodal Immersive Learning with AI for Psychomotoric Skills View project Clustering Technique for Collaborative Filtering and the Application to Venue Recommendation,” 2010. [Online]. Available: <http://www.cs.wisc.edu/dbworld/>

- [38] D. Billsus and M. J. Pazzani, “User Modeling for Adaptive News Access.”
- [39] Adiyansjah, A. A. S. Gunawan, and D. Suhartono, “Music recommender system based on genre using convolutional recurrent neural networks,” in *Procedia Computer Science*, Elsevier B.V., 2019, pp. 99–109. doi: 10.1016/j.procs.2019.08.146.
- [40] R. Devooght and H. Bersini, “Collaborative Filtering with Recurrent Neural Networks,” Aug. 2016, [Online]. Available: <http://arxiv.org/abs/1608.07400>
- [41] G. De Souza Pereira Moreira, F. Ferreira, and A. M. Da Cunha, “News Session-Based Recommendations using Deep Neural Networks,” in *ACM International Conference Proceeding Series*, Association for Computing Machinery, Oct. 2018, pp. 15–23. doi: 10.1145/3270323.3270328.