



ΠΑΝΕΠΙΣΤΗΜΙΟ ΠΑΤΡΩΝ
ΠΟΛΥΤΕΧΝΙΚΗ ΣΧΟΛΗ
ΤΜΗΜΑ ΜΗΧΑΝΙΚΩΝ Η/Υ & ΠΛΗΡΟΦΟΡΙΚΗΣ

ΔΙΔΑΚΤΟΡΙΚΗ ΔΙΑΤΡΙΒΗ

**Τεχνικές Επεξεργασίας Φυσικής Γλώσσας, Εξαγωγής
Γνώσης και Αυτόματης Δημιουργίας Προσαρμοσμένης
στον Χρήστη Ανάδρασης για Ευφυή Συστήματα
Διδασκαλίας**

Ισίδωρος Περίκος

Επιβλέπων: Ιωάννης Χατζηλυγερούδης,
Αναπληρωτής Καθηγητής

Πάτρα, Σεπτέμβριος 2016



UNIVERSITY OF PATRAS
DEPARTMENT OF COMPUTER ENGINEERING AND INFOMATICS

DOCTORAL DISSERTATION

**Techniques for Natural Language Processing,
Knowledge Extraction and Automatic Creation
of User Adapted Feedback for Intelligent
Tutoring Systems**

Isidoros Perikos

Supervisor: Ioannis Hatzilygeroudis, Associate Professor

Patra, September 2016

ΠΑΝΕΠΙΣΤΗΜΙΟ ΠΑΤΡΩΝ
ΠΟΛΥΤΕΧΝΙΚΗ ΣΧΟΛΗ
ΤΜΗΜΑ ΜΗΧΑΝΙΚΩΝ Η/Υ & ΠΛΗΡΟΦΟΡΙΚΗΣ

ΔΙΔΑΚΤΟΡΙΚΗ ΔΙΑΤΡΙΒΗ

**Τεχνικές Επεξεργασίας Φυσικής Γλώσσας, Εξαγωγής
Γνώσης και Αυτόματης Δημιουργίας Προσαρμοσμένης
στον Χρήστη Ανάδρασης για Ευφυή Συστήματα
Διδασκαλίας**

Ισίδωρος Περίκος

Επταμελής Εξεταστική Επιτροπή:

Ιωάννης Χατζηλυγερουδης, Αναπληρωτής Καθηγητής, Τμήμα Μηχανικών Η/Υ & Πληροφορικής, Πανεπιστήμιου Πατρών (*Επιβλέπων*)

Σπυρίδων Λυκοθανάσης, Καθηγητής, Τμήμα Μηχανικών Η/Υ & Πληροφορικής, Πανεπιστήμιου Πατρών (*Μέλος της Τριμελούς Επιτροπής*)

Χρήστος Μακρής, Επίκουρος Καθηγητής, Τμήμα Μηχανικών Η/Υ & Πληροφορικής, Πανεπιστήμιου Πατρών (*Μέλος της Τριμελούς Επιτροπής*)

Νικόλαος Αβούρης, Καθηγητής, Τμήμα Ηλεκτρολόγων Μηχανικών & Τεχνολογίας Υπολογιστών, Πανεπιστήμιου Πατρών

Μαρία Βίρβου, Καθηγήτρια, Τμήμα Πληροφορικής, Πανεπιστήμιου Πειραιώς

Βασίλειος Μεγαλοοικονόμου, Καθηγητής, Τμήμα Μηχανικών Η/Υ & Πληροφορικής, Πανεπιστήμιου Πατρών

Νικόλαος Μπουρμπάκης, OBR Distinguished Professor, Department of Computer Science and Engineering, Wright State University, USA

Ευχαριστίες

Ολοκληρώνοντας την εκπόνηση της διδακτορικής διατριβής μου θα ήθελα να ευχαριστήσω όλους τους ανθρώπους με τους οποίους συνεργάστηκα καθ' όλη την διάρκεια της πορείας μου.

Αρχικά, ένα ιδιαίτερα θερμό ευχαριστώ οφείλω στον επιβλέποντα της διδακτορικής μου διατριβής, τον Αναπληρωτή Καθηγητή κ. Ιωάννη Χατζηλυγερούδη, που απετέλεσε τον καθοδηγητή μου σε όλα τα στάδια της διατριβής, για την πολύτιμη και πολύπλευρη βοήθεια του, την συνεχή καθοδήγηση του στην ερευνητική μου πορεία καθώς και την αμέριστη εμπιστοσύνη που μου έδειξε όλα αυτά τα χρόνια. Η προσφορά του στην επιστημονική μου εξέλιξη είναι πραγματικά ανεκτίμητη.

Θα ήθελα να ευχαριστήσω επίσης θερμά τον Καθηγητή κ. Σπυρίδωνα Λυκοθανάση και τον Επίκουρο Καθηγητή κ. Χρήστο Μακρή για τη συνεργασία και την συμβουλευτική υποστήριξη που παρείχαν κατά την εκπόνηση της διατριβής. Επίσης, ευχαριστώ θερμά τα μέλη της επταμελούς επιτροπής κ. Νικόλαο Αβούρη, κ. Μαρία Βίρβου, κ. Βασίλειο Μεγαλοοικονόμου και κ. Νικόλαο Μπουρμπάκη για την μεγάλη τιμή που μου έκαναν.

Επίσης, θα ήθελα να ευχαριστήσω τους συνεργάτες και φίλους μου στην Ομάδα Τεχνητής Νοημοσύνης, Φωτεινή Γριβοκωστοπούλου και Κωνσταντίνο Κόβα για την άριστη και παραγωγική συνεργασία μας όλα αυτά τα χρόνια.

Τέλος, οφείλω ένα μεγάλο ευχαριστώ στην οικογένεια μου για την υποστήριξη που μου παρείχε καθ' όλη την διάρκεια των σπουδών μου.

Ισίδωρος Περίκος

Πάτρα, Σεπτέμβριος 2016

Στην μνήμη του πατέρα μου

Περίληψη

Η παρούσα διδακτορική διατριβή αποσκοπεί στο σχεδιασμό, στην ανάπτυξη και στην αξιοποίηση μηχανισμών και καινοτόμων μεθοδολογιών από τις περιοχές της επεξεργασίας φυσικής γλώσσας, της μηχανικής μάθησης για την αλληλεπίδραση μεταξύ ανθρώπου και υπολογιστικών συστημάτων. Κύρια αντικείμενα του προβλήματος που πραγματεύεται η παρούσα διατριβή αφορούν την μοντελοποίηση του κύκλου παροχής ανάδρασης σε ευφυή εκπαιδευτικά συστήματα, την αναγνώριση της συναισθηματικής κατάστασης του χρήστη και την δυνατότητα κατανόησης κειμένου για επικοινωνία ανθρώπου-υπολογιστή σε φυσική γλώσσα.

Αρχικά, παρουσιάζεται μια νέα μεθοδολογία δημιουργίας ανάδρασης σε ευφυή εκπαιδευτικά συστήματα, με στόχο την παροχή κατάλληλης καθοδήγησης κατά την διάρκεια της αλληλεπίδρασης. Η μεθοδολογία που αναπτύχθηκε περιλαμβάνει την εισαγωγή μιας γενικής μοντελοποίησης της απάντησης των εκπαιδευόμενων, την εισαγωγή ενός γενικού ιεραρχικού μοντέλου οργάνωσης της ανάδρασης σε επίπεδα καθώς και το καθορισμό της κατηγοριοποίησης των λαθών με βάση το πεδίο της εφαρμογής. Επίσης, αναπτύχθηκε ένας μηχανισμός για την παραγωγή ανάδρασης, που έχει ως στόχο να παράγει και να παρέχει αυτόματα τους κατάλληλους τύπους ανάδρασης με βάση το ιεραρχικό μοντέλο. Τα αποτελέσματα της πειραματικής μελέτης έδειξαν ότι το πλαίσιο που εισάγεται είναι ιδιαίτερα αποτελεσματικό.

Για την κατανόηση κειμένου και την επικοινωνία σε φυσική γλώσσα, μια σημαντική συνεισφορά της διατριβής αποτελεί η εισαγωγή μιας νέας, καινοτόμου μεθοδολογίας για την αυτόματη μετατροπή φυσικής γλώσσας σε *κατηγορηματική λογική πρώτης τάξης* (*first-order predicate logic*), συντομογραφικά ΚΛΠΤ (FOPC). Η μεθοδολογία βασίζεται στη *βασισμένη σε περιπτώσεις συλλογιστική* (*case-based reasoning*) και στην αρχή ότι προτάσεις φυσικής γλώσσας που έχουν ίδια ή παρόμοια συντακτική δομή και *δέντρα εξαρτήσεων* (*dependency trees*) θα έχουν ίδια ή παρόμοια δομή έκφρασης σε κατηγορηματική λογική πρώτης τάξης. Για τον προσδιορισμό της ομοιότητας μεταξύ των προτάσεων φυσικής γλώσσας χρησιμοποιείται η *απόσταση μετασχηματισμού δέντρου* (*tree edit distance*), μετρική η οποία εφαρμόζεται στις δεντρικές δομές των δέντρων εξαρτήσεων. Επίσης, χρησιμοποιούνται κανόνες που έχουν ως στόχο να καθοριστεί η ακριβής έκφραση ΚΛΠΤ της νέας πρότασης. Τα πειραματικά αποτελέσματα έδειξαν ότι η προσέγγιση που εισάγεται είναι αποτελεσματική στις περιπτώσεις που η ομοιότητα είναι μεγάλη. Επίσης, εισάγεται μια προσέγγιση που σχεδιάστηκε για την ανάλυση των προτάσεων και τον προσδιορισμό του επιπέδου δυσκολίας της μετατροπής τους, η οποία βασίζεται σε χαρακτηριστικά της δομής των προτάσεων, στην ανάλυση των αντίστοιχων εκφράσεων ΚΛΠΤ, καθώς και σε χαρακτηριστικά που εξάγονται από την αντιστοίχισή τους. Τα αποτελέσματα της πειραματικής μελέτης έδειξαν ότι η προσέγγιση είναι ακριβής στον προσδιορισμό του κατάλληλου επιπέδου δυσκολίας, γεγονός που

επαληθεύεται και από τις επιδόσεις του μηχανισμού αυτόματης μετατροπής. Ακόμη, εισάγεται μια μεθοδολογία για την παραγωγή λογικά ισοδύναμων προτάσεων (παραφράσεων), που βασίζεται σε ένα σύνολο από *τεχνικές παράφρασης (paraphrasing techniques)*, οι οποίες μπορούν να τροποποιήσουν τη δομή προτάσεων και να μετατρέψουν δομικά μέρη του δέντρου εξάρτησης των προτάσεων σε ισοδύναμη μορφή. Τα πειραματικά αποτελέσματα της αξιολόγησης έδειξαν μια καλή απόδοση των τεχνικών και του εργαλείου που αναπτύχθηκε.

Ένα άλλο βασικό αντικείμενο που πραγματεύεται η διατριβή είναι η αναγνώριση της συναισθηματικής κατάστασης του χρήστη. Για τον σκοπό αυτό, σχεδιάστηκαν και αναπτύχθηκαν μέθοδοι για την ανάλυση κειμένου και τον προσδιορισμό του συναισθηματικού περιεχομένου του καθώς και μέθοδοι ανάλυσης εκφράσεων προσώπου και αναγνώρισης της συναισθηματικής τους κατάστασης.

Αρχικά, σχεδιάσαμε και αναπτύξαμε νέες προσεγγίσεις για την αυτόματη ανάλυση προτάσεων φυσικής γλώσσας και την εξαγωγή των συναισθηματικών πληροφοριών που φέρουν. Οι μέθοδοι που αναπτύχθηκαν βασίζονται σε σχήματα *ομαδοποιημένων ταξινομητών (ensemble classifiers)*, που συνδυάζουν μεθόδους επεξεργασίας φυσικής γλώσσας και ανάλυσης των δομών των προτάσεων με μεθόδους μηχανικής μάθησης. Πιο συγκεκριμένα, εισάγεται μια βασισμένη σε γνώση (knowledge based) προσέγγιση, που βασίζεται στην εις βάθος ανάλυση των δομών των προτάσεων, και γίνεται προσδιορισμός του συναισθηματικού περιεχόμενου τους με βάση τη συντακτική δομή τους και λεξικολογικούς πόρους. Επίσης, γίνεται χρήση μεθόδων μηχανικής μάθησης μέσω χρήσης βασικών στατιστικών ταξινομητών, όπως ο *απλοϊκός bayes (naïve Bayes)* και οι *μηχανές διανυσμάτων υποστήριξης (support vector machines)*, καθώς και συνδυασμός τους σε σχήμα ομαδοποιημένων ταξινομητών μαζί με τη βασισμένη στη γνώση προσέγγιση. Τα πειραματικά αποτελέσματα της αξιολόγησης έδειξαν ότι τα ομαδοποιημένα σχήματα παρουσίασαν μια αύξηση της απόδοσης κατά 4.8% έως 8.4% σε σχέση με τον καλύτερο ταξινομητή.

Ως προς την αυτόματη ανάλυση εκφράσεων προσώπου και την αναγνώριση του συναισθηματικού τους περιεχομένου, αρχικά, εισάγεται ένα πλαίσιο μοντελοποίησης των εκφράσεων του προσώπου, το οποίο ακολουθεί μια αναλυτική, τοπικής εμβέλειας προσέγγιση και αναπαριστά μια έκφραση με 25 χαρακτηριστικά που εξάγει από σημεία όπως, τα μάτια, το στόμα και τα φρύδια. Στην συνέχεια, μελετάται η απόδοση νευρο-ασαφούς προσέγγισης και μεθόδων μηχανικής μάθησης, όπως οι random forest και support vector machines, στον προσδιορισμό της ύπαρξης ή όχι συναισθημάτων σε μια έκφραση καθώς και στην αναγνώριση της ακριβούς συναισθηματικής της κατάστασης. Έπειτα, με βάση την απόδοση τους, παρουσιάζεται ομαδοποιημένο σχήμα ταξινόμησης, το οποίο συνδυάζει ταξινομητές που βασίζονται στις τρεις παραπάνω μεθόδους σε ένα σχήμα πλειοψηφίας. Στόχος του ομαδοποιημένου σχήματος αποτελεί η αύξηση της απόδοσης και η αποτελεσματική γενίκευση σε νέα δεδομένα εκφράσεων προσώπου. Τα αποτελέσματα της πειραματικής μελέτης έδειξαν ότι το ομαδοποιημένο σχήμα παρουσίασε μια αύξηση της απόδοσης της τάξης του 3.5% σε σχέση με τον καλύτερο ταξινομητή.

Συμπερασματικά, οι τεχνικές που προτείνονται στην παρούσα διδακτορική διατριβή επεκτείνουν εργασίες της διεθνούς βιβλιογραφίας, προσθέτοντας και εισάγοντας νέες μεθόδους αντιμετώπισης στα θεματικά που πραγματεύεται.

Abstract

The main aim of this thesis concerns the design, development and use of innovative mechanisms and methodologies from the areas of natural language processing and machine learning in the field of human computer interaction. The main subjects of the problem addressed by this thesis concern, the modeling of the feedback delivery cycle in intelligent educational systems (ITSs), the recognition of the emotional state of users and the ability of text understanding for human-computer communication in natural language.

Initially, we introduce a new feedback generation methodology for intelligent educational systems, in order to provide appropriate guidance during the interaction with them. The developed methodology includes, the introduction of a general model of learners' responses, the introduction of a general hierarchical model for the organization of feedback in levels and the creation of a classification scheme of the domain errors. Furthermore, a mechanism for the production of feedback is presented, designed to generate and automatically provide the appropriate types of feedback based on the hierarchical model. The experimental results showed that the introduced framework is quite effective.

Regarding text understanding and communication in natural language, an important contribution of this thesis is the introduction of an innovative methodology to automatically convert natural language into first order predicate logic (FOPL). The methodology is based on a case-based reasoning cycle and on the principle that natural language sentences that have the same or similar syntactic structures and dependency trees will have the same or similar expressions in FOPL. The *tree edit distance* metric is applied to the tree structures of the sentences' dependency trees, in order to determine the similarity. Also, rules are used to determine the FOPL expression of a new sentence. The experimental results showed that the introduced approach is effective, when the similarity between new sentences and the formalized ones stored in the knowledge based is high. Furthermore, an approach was designed to analyze sentences and determine the difficulty level of their formalization process. The approach is based on characteristics of sentence structure, characteristics of their corresponding FOL expressions as well as features extracted by comparing with their NL and FOL representation. The results of the experimental study show that the approach is accurate in determining the appropriate level of difficulty. Moreover, a method for producing logically equivalent sentences (paraphrases) is introduced. It is based on a set of developed paraphrase techniques, which can modify the sentences structure in order to convert parts of a sentence's dependency tree into equivalent forms. The experimental results of the evaluation show good behavior of the techniques and tools designed.

Another main aspect of the thesis concerns the recognition of users' emotional states. Methods for text analysis and recognition of the emotional content and also methods for analyzing facial expressions and recognizing their emotional state are presented. Initially, we present new approaches for the automatic analysis of natural language sentences and the extraction of emotional information. The developed methods are based on an ensemble classifier schema, which combine deep natural language processing methods with machine learning methods. More specifically, a knowledge based approach is introduced,

based on the in-depth analysis of the syntactical structures of sentences and the estimation of the emotional content via the syntactic structure and lexical resources. Machine learning methods are also used via using basic statistical classifiers, such as naive Bayes and support vector machines (SVMs), combined into an ensemble classifier schema including the knowledge-based approach too. The experimental results of the evaluation indicate that the ensemble classifier schema shows a performance increase by 4.8% to 8.4% relative to the best individual classifier.

With regards to the methods that automatically analyze facial expressions and recognize their emotional content, first, a modeling framework of facial expressions is presented, which follows an analytical, local-based approach. It represents an expression with a set of twenty-five features extracted from points such as the eyes, the mouth and the eyebrows. Then, the performance of a neuro-fuzzy approach as well as machine learning methods, such as random forest and support vector machines, are studied in determining the existence of emotions in an expression and recognizing the exact emotional state. After that, an ensemble classifier schema is presented, that combines the above classifiers in a majority schema. The goal of the ensemble schema is to increase the efficiency and the scalability in new facial expression data. The experimental study and the results show that the ensemble schema showed a performance increase of about 3.5% compared to the best individual classifier.

In conclusion, the techniques proposed in this thesis extend the existing approaches in international literature, adding and introducing new methods for coping with the problems they are dealing with.

Περιεχόμενα

1	ΕΙΣΑΓΩΓΗ	25
1.1	ΑΝΤΙΚΕΙΜΕΝΑ ΚΑΙ ΣΥΝΕΙΣΦΟΡΑ ΤΗΣ ΕΡΓΑΣΙΑΣ	26
1.2	ΔΙΑΡΘΡΩΣΗ ΤΗΣ ΔΙΑΤΡΙΒΗΣ	28
2	ΜΟΝΤΕΛΟΠΟΙΗΣΗ ΑΝΑΔΡΑΣΗΣ ΣΕ ΕΥΦΥΗ ΕΚΠΑΙΔΕΥΤΙΚΑ ΣΥΣΤΗΜΑΤΑ	31
2.1	ΕΙΣΑΓΩΓΗ	31
2.1.1	Γενικό Πλαίσιο Ανάδρασης	33
2.1.2	Μοντελοποίηση Απάντησης Εκπαιδευόμενων	33
2.1.3	Σχεδιασμός Ιεραρχίας Ανάδρασης	35
2.1.3.1	Ανάδραση πριν την υποβολή απάντησης	36
2.1.3.2	Ανάδραση μετά την υποβολή λανθασμένης απάντησης	37
2.1.3.3	Μετα-γνωστική (Meta-cognitive) Ανάδραση	39
2.1.4	Πλαίσιο Παραγωγής Ανάδρασης	41
2.1.4.1	Μηχανισμός Αναγνώρισης λαθών	42
2.1.4.2	Μηχανισμός Παραγωγής Ανάδρασης	42
2.2	ΕΦΑΡΜΟΓΗ ΤΟΥ ΠΛΑΙΣΙΟΥ ΑΝΑΔΡΑΣΗΣ ΣΤΟ ΕΥΦΥΕΣ ΕΚΠΑΙΔΕΥΤΙΚΟ ΣΥΣΤΗΜΑ NLtoFOL	43
2.2.1	Διαδικασία Μετατροπής Φυσικής Γλώσσας σε ΚΛΠΤ.....	43
2.2.2	Κατηγοριοποίηση Λαθών στο Σύστημα NLtoFOL.....	47
2.2.3	Παραγωγή Κειμενικών Μηνυμάτων Ανάδρασης	51
2.2.4	Πειραματική Αξιολόγηση	53
2.2.4.1	Πρώτη πειραματική μελέτη	53
2.2.4.2	Ανάλυση των μαθησιακών ενεργειών των εκπαιδευόμενων	59
2.2.4.3	Δεύτερη Πειραματική Μελέτη	64
3	ΑΥΤΟΜΑΤΗ ΜΕΤΑΤΡΟΠΗ ΠΡΟΤΑΣΕΩΝ ΦΥΣΙΚΗΣ ΓΛΩΣΣΑΣ ΣΕ ΛΟΓΙΚΗ	67
3.1	ΕΙΣΑΓΩΓΗ	67
3.2	ΣΧΕΤΙΚΕΣ ΕΡΓΑΣΙΕΣ	69
3.3	ΒΑΣΙΚΕΣ ΈΝΝΟΙΕΣ.....	71
3.3.1	Σύνταξη και Σημασιολογία Κατηγορηματικής Λογικής	71
3.3.2	Οι Δυσκολίες Αυτοματοποίησης του Φορμαλισμού Φυσικής Γλώσσας	73
3.4	ΠΡΟΣΔΙΟΡΙΣΜΟΣ ΠΟΛΥΠΛΟΚΟΤΗΤΑΣ ΦΟΡΜΑΛΙΣΜΟΥ ΠΡΟΤΑΣΕΩΝ ΦΥΣΙΚΗΣ ΓΛΩΣΣΑΣ	74
3.5	ΜΕΘΟΔΟΛΟΓΙΑ ΑΥΤΟΜΑΤΗΣ ΜΕΤΑΤΡΟΠΗΣ ΦΥΣΙΚΗΣ ΓΛΩΣΣΑΣ ΣΕ ΚΛΠΤ	81
3.5.1	Συλλογιστική Βασισμένη σε Περιπτώσεις.....	81
3.5.2	Αυτόματη Μετατροπή ΦΓ σε ΚΛΠΤ Βασισμένη σε Περιπτώσεις	83
3.6	ΠΕΙΡΑΜΑΤΙΚΑ ΑΠΟΤΕΛΕΣΜΑΤΑ.....	90

3.6.1	Δημιουργία Συνόλου Δεδομένων.....	91
3.6.2	Πειραματική Αξιολόγηση Προσδιορισμού Δυσκολίας.....	91
3.6.3	Πειραματική Αξιολόγηση του Μηχανισμού Αυτόματης Μετατροπής ΦΓ σε ΚΛΠΤ....	92
4	ΑΝΑΛΥΣΗ ΚΕΙΜΕΝΟΥ ΚΑΙ ΔΗΜΙΟΥΡΓΙΑ ΠΑΡΑΦΡΑΣΕΩΝ	99
4.1	ΕΙΣΑΓΩΓΗ	99
4.2	ΣΧΕΤΙΚΕΣ ΕΡΓΑΣΙΕΣ	101
4.3	ΜΗΧΑΝΙΣΜΟΣ ΠΑΡΑΓΩΓΗΣ ΠΑΡΑΦΡΑΣΕΩΝ.....	103
4.3.1	Ανάλυση της Φυσικής Γλώσσας.....	104
4.3.2	Δημιουργία Τεχνικών Παράφρασης.....	105
4.4	ΠΕΙΡΑΜΑΤΙΚΗ ΑΞΙΟΛΟΓΗΣΗ	111
5	ΕΠΕΞΕΡΓΑΣΙΑ ΚΕΙΜΕΝΟΥ ΚΑΙ ΠΡΟΣΔΙΟΡΙΣΜΟΣ ΣΥΝΑΙΣΘΗΜΑΤΙΚΟΥ ΠΕΡΙΕΧΟΜΕΝΟΥ	115
5.1	ΕΙΣΑΓΩΓΗ	115
5.2	ΜΕΘΟΔΟΙ, ΛΕΞΙΚΟΛΟΓΙΚΟΙ ΠΟΡΟΙ ΚΑΙ ΣΧΕΤΙΚΕΣ ΕΡΓΑΣΙΕΣ.....	117
5.2.1	Λεξικόλογικοί Πόροι.....	118
5.2.2	Σχετικές Εργασίες.....	119
5.3	ΜΟΝΤΕΛΟΠΟΙΗΣΗ ΣΥΝΑΙΣΘΗΜΑΤΩΝ.....	123
5.4	ΣΧΕΔΙΑΣΗ ΚΑΙ ΑΝΑΠΤΥΞΗ ΟΜΑΔΟΠΟΙΗΜΕΝΟΥ ΤΑΞΙΝΟΜΗΤΗ ΑΝΑΓΝΩΡΙΣΗΣ ΣΥΝΑΙΣΘΗΜΑΤΙΚΟΥ ΠΕΡΙΕΧΟΜΕΝΟΥ ΚΕΙΜΕΝΟΥ	126
5.4.1	Βασικές αρχές σχεδιασμού	126
5.4.2	Ομαδοποιημένος Ταξινομητής Αναγνώρισης Συναισθημάτων.....	128
5.4.3	Εξαγωγή Χαρακτηριστικών από Κείμενα	130
5.4.4	Ταξινομητής Naïve Bayes	130
5.4.5	Μηχανές Διανυσμάτων Υποστήριξης – SVM ταξινομητής	131
5.4.6	Μηχανισμός Βασισμένος στην Γνώση (Knowledge-Based Mechanism)	132
5.4.6.1	Μορφοσυντακτική ανάλυση προτάσεων.....	134
5.4.6.2	Δημιουργία συναισθηματικών δομών	135
5.4.7	Ομαδοποίηση Ταξινομητών	137
5.4.7.1	Αρχή Πλειοψηφίας.....	137
5.4.7.2	Bagging	138
5.4.7.3	Boosting.....	139
5.5	ΠΡΟΣΔΙΟΡΙΣΜΟΣ ΣΥΝΑΙΣΘΗΜΑΤΙΚΗΣ ΠΟΛΙΚΟΤΗΤΑΣ.....	141
5.6	ΠΕΙΡΑΜΑΤΙΚΗ ΑΞΙΟΛΟΓΗΣΗ	142
5.6.1	Δημιουργία Συνόλου Δεδομένων.....	143
5.6.2	Αξιολόγηση Προσδιορισμού Ύπαρξης Συναισθηματικού Περιεχόμενου	143
5.6.3	Αξιολόγηση Προσδιορισμού Συναισθηματικής Πολικότητας	151

5.6.4	Αξιολόγηση Αναγνώρισης Συναισθηματικών Καταστάσεων	156
5.7	ΠΛΑΙΣΙΟ ΣΥΝΑΙΣΘΗΜΑΤΙΚΗΣ ΜΟΝΤΕΛΟΠΟΙΗΣΗΣ ΚΕΙΜΕΝΙΚΗΣ ΣΥΖΗΤΗΣΗΣ	160
5.7.1	Ορισμός του Πλαισίου	160
5.7.2	Εφαρμογή του Πλαισίου	161
6	ΑΝΑΓΝΩΡΙΣΗ ΣΥΝΑΙΣΘΗΜΑΤΙΚΗΣ ΚΑΤΑΣΤΑΣΗΣ ΕΚΦΡΑΣΕΩΝ ΠΡΟΣΩΠΟΥ	165
6.1	ΕΙΣΑΓΩΓΗ	165
6.2	ΒΑΣΙΚΕΣ ΈΝΝΟΙΕΣ ΚΑΙ ΣΧΕΤΙΚΕΣ ΕΡΓΑΣΙΕΣ	167
6.2.1	Βασικές έννοιες	167
6.2.1.1	Facial Action Coding System (FACS)	167
6.2.1.2	Παράμετροι απόδοσης κίνησης προσώπου (Facial Animation Parameters -FAPs)	168
6.2.2	Σχετικές Εργασίες	169
6.3	ΣΥΣΤΗΜΑ ΣΥΝΑΙΣΘΗΜΑΤΙΚΗΣ ΑΝΑΓΝΩΡΙΣΗΣ ΕΚΦΡΑΣΕΩΝ ΠΡΟΣΩΠΟΥ	171
6.3.1	Ανάλυση Προσώπου	172
6.3.2	Εξαγωγή Χαρακτηριστικών	173
6.3.3	Ταξινομητές Αναγνώρισης Συναισθηματικών Καταστάσεων	176
6.3.3.1	Νευρο-ασαφής (Neuro-Fuzzy) Ταξινομητής	176
6.3.3.2	Ταξινομητής Μηχανής Διανουσμάτων Υποστήριξης (Support Vector Machine)	179
6.3.3.3	Ταξινομητής Τυχαίου Δάσους (Random Forest)	179
6.4	ΠΕΙΡΑΜΑΤΙΚΗ ΑΞΙΟΛΟΓΗΣΗ	181
6.4.1	Δημιουργία Συνόλου Δεδομένων	181
6.4.2	Αξιολόγηση προσδιορισμού ύπαρξης Συναισθηματικού Περιεχόμενου	181
6.4.3	Αξιολόγηση Αναγνώρισης Συναισθηματικών Καταστάσεων	184
7	ΣΥΜΠΕΡΑΣΜΑΤΑ ΚΑΙ ΜΕΛΛΟΝΤΙΚΗ ΈΡΕΥΝΑ	189
7.1	ΣΥΜΠΕΡΑΣΜΑΤΑ	189
7.2	ΜΕΛΛΟΝΤΙΚΗ ΕΡΕΥΝΑ	192

Κατάλογος Εικόνων

<i>Εικόνα 2.1 Αρχιτεκτονική Παραγωγής Ανάδρασης</i>	<i>41</i>
<i>Εικόνα 2.2 Διαδικασία Πειράματος</i>	<i>54</i>
<i>Εικόνα 2.3 Τα κέρδη μάθησης για τις δύο ομάδες</i>	<i>56</i>
<i>Εικόνα 2.4 Τα αποτελέσματα του ερωτηματολογίου της ομάδας GroupA</i>	<i>58</i>
<i>Εικόνα 2.5 Τα αποτελέσματα του ερωτηματολογίου της ομάδας GroupB</i>	<i>59</i>
<i>Εικόνα 2.6 Μέσος όρος των λαθών ανά εκπαιδευόμενο για τις δυο ομάδες</i>	<i>64</i>
<i>Εικόνα 2.7 Διαδικασία Πειράματος</i>	<i>65</i>
<i>Εικόνα 2.8 Κέρδος μάθησης για τις δύο ομάδες.</i>	<i>66</i>
<i>Εικόνα 3.1. Αρχιτεκτονική Συστήματος Προσδιορισμού Δυσκολίας Φορμαλισμού</i>	<i>75</i>
<i>Εικόνα 3.2 Κύκλος λειτουργίας μετατροπής ΦΓ σε ΚΛΠΤ</i>	<i>83</i>
<i>Εικόνα 3.3 Σχηματική αναπαράσταση του προσδιορισμού όμοιων προτάσεων</i>	<i>86</i>
<i>Εικόνα 3.4 Τα δέντρα Εξαρτήσεων των προτάσεων.....</i>	<i>89</i>
<i>Εικόνα 3.5 Η απόδοση ως προς το επίπεδο δυσκολίας των προτάσεων.</i>	<i>94</i>
<i>Εικόνα 3.6 Η αξιολόγηση ως προς το επίπεδο δυσκολίας των προτάσεων.....</i>	<i>98</i>
<i>Εικόνα 4.1 Αρχιτεκτονική Συστήματος.....</i>	<i>103</i>
<i>Εικόνα 4.2 Το δέντρο εξαρτήσεων της πρότασης Π1</i>	<i>106</i>
<i>Εικόνα 4.3 Οι μετατροπές στο δέντρο εξαρτήσεων της Π1.....</i>	<i>106</i>
<i>Εικόνα 4.4 Το δέντρο εξαρτήσεων της πρότασης Π2</i>	<i>107</i>
<i>Εικόνα 4.5 Το δέντρο εξαρτήσεων της πρότασης Π3</i>	<i>108</i>
<i>Εικόνα 4.6 Η διεπαφή του εργαλείου παραγωγής παραφράσεων.</i>	<i>110</i>
<i>Εικόνα 4.7 Προσδιορισμός καταλληλότητας των παραγόμενων παραφράσεων.</i>	<i>110</i>
<i>Εικόνα 4.8 Απόδοση του μηχανισμού σε δείγμα 15 προτάσεων.....</i>	<i>112</i>
<i>Εικόνα 4.9 Η απόδοση για κάθε είδος τεχνικής</i>	<i>113</i>
<i>Εικόνα 5.1 Μέθοδοι και τεχνικές ανάλυσης και συναισθηματικής ταξινόμησης κειμένου</i>	<i>118</i>
<i>Εικόνα 5.2 Επίπεδα και συναισθηματικές καταστάσεις του μοντέλου Parrot.....</i>	<i>124</i>
<i>Εικόνα 5.3 Το συναισθηματικό μοντέλο του Plutchik.</i>	<i>125</i>
<i>Εικόνα 5.4 Το Circumflex μοντέλο του Russell.</i>	<i>126</i>
<i>Εικόνα 5.5 Βασικές σκοπιές ομαδοποίησης ταξινομητών.</i>	<i>128</i>
<i>Εικόνα 5.6 Βασική αρχιτεκτονική του ομαδοποιημένου ταξινομητή.</i>	<i>129</i>

<i>Εικόνα 5.7. Η αρχιτεκτονική του βασισμένου σε γνώση μηχανισμού</i>	<i>132</i>
<i>Εικόνα 5.8 Το δέντρο εξαρτήσεων της πρότασης “She kissed her aunt with great happiness”</i>	<i>134</i>
<i>Εικόνα 5.9 Η μέθοδος Bagging.</i>	<i>139</i>
<i>Εικόνα 5.10 Η μέθοδος Boosting.</i>	<i>141</i>
<i>Εικόνα 5.11. Αποτύπωση των συναισθημάτων στον χώρο του μοντέλου του Russell</i>	<i>142</i>
<i>Εικόνα 5.12 Η ορθότητα των ταξινομητών σε όλο το σύνολο δεδομένων</i>	<i>145</i>
<i>Εικόνα 5.13 Η ορθότητα των βασικών ταξινομητών σε κάθε είδος δεδομένων κειμένου.</i>	<i>145</i>
<i>Εικόνα 5.14 Η ορθότητα των ομαδοποιημένων ταξινομητών σε κάθε είδος δεδομένων κειμένου.....</i>	<i>148</i>
<i>Εικόνα 5.15 Η ορθότητα των ταξινομητών σε όλο το σύνολο δεδομένων</i>	<i>152</i>
<i>Εικόνα 5.16 Η ορθότητα των βασικών ταξινομητών σε κάθε είδος δεδομένων κειμένου</i>	<i>153</i>
<i>Εικόνα 5.17 Η ορθότητα των ομαδοποιημένων ταξινομητών σε κάθε είδος δεδομένων κειμένου.....</i>	<i>155</i>
<i>Εικόνα 5.18 Συναισθηματικό γράφημα της θεματικής περιοχής “Syrian civil war”</i>	<i>162</i>
<i>Εικόνα 5.19 Συναισθηματικό γράφημα της θεματικής περιοχής “Valentine’s day”</i>	<i>163</i>
<i>Εικόνα 6.1. Μερικά AUs του προσώπου και συνδυασμοί τους.....</i>	<i>168</i>
<i>Εικόνα 6.2 Διάγραμμα Ροής του Συστήματος.....</i>	<i>172</i>
<i>Εικόνα 6.3 Τα στάδια της ανάλυσης των εικόνων προσώπων.....</i>	<i>173</i>
<i>Εικόνα 6.4. Απομόνωση περιοχών ενδιαφέροντος</i>	<i>174</i>
<i>Εικόνα 6.5 Εξαγωγή Χαρακτηριστικών από κάθε περιοχή ενδιαφέροντος και η μοντελοποίηση της έκφρασης</i>	<i>174</i>
<i>Εικόνα 6.6 Βασική αρχιτεκτονική του ομαδοποιημένου ταξινομητή.</i>	<i>176</i>
<i>Εικόνα 6.7 Αρχιτεκτονική ANFIS.....</i>	<i>177</i>
<i>Εικόνα 6.8 Λειτουργία τυχαίων δασών.....</i>	<i>180</i>

Κατάλογος Πινάκων

Πίνακας 2.1 Οργάνωση ανάδρασης πριν την υποβολή απάντησης.....	36
Πίνακας 2.2 Οργάνωση ανάδρασης μετά την υποβολή απάντησης	37
Πίνακας 2.3 Παραδείγματα ανάδρασης με βάση την συμπεριφορά του εκπαιδευόμενου	39
Πίνακας 2.4 Ανάλυση του βήματος 2 της διαδικασίας NLtoFOL	45
Πίνακας 2.5 Κατηγοριοποίηση λαθών της διαδικασίας φορμαλισμού.....	47
Πίνακας 2.6 Απάντηση Εκπαιδευόμενου και σωστή απάντηση για το δεύτερο βήμα της διαδικασίας.	49
Πίνακας 2.7 Παραδείγματα προτύπων μηνυμάτων	51
Πίνακας 2.8 Πειραματικά Αποτελέσματα ANCOVA.....	56
Πίνακας 2.9 Ερωτήσεις ερωτηματολογίου και για τις δυο ομάδες	57
Πίνακας 2.10 Ανάλυση συμπεριφοράς εκπαιδευόμενων.....	60
Πίνακας 2.11 Ποσοστό σωστών απαντήσεων ανά επίπεδο παρεχόμενης βοήθειας ...	62
Πίνακας 2.12 Αποτελέσματα ANCOVA.....	66
Πίνακας 3.1 Κανόνες για τον προσδιορισμό της δυσκολίας με βάση την ΚΛΠΤ έκφραση.....	76
Πίνακας 3.2 Κανόνες για τον προσδιορισμό του επιπέδου δυσκολίας με βάση την έκφραση ΚΛΠΤ και τη σημασιολογία της Φυσικής Γλώσσας	81
Πίνακας 3.3 Κανόνες δημιουργίας εκφράσεων ΚΛΠΤ	86
Πίνακας 3.4 Λέξεις ποσοτικοποίησης μεταβλητών ατομικών εκφράσεων.....	88
Πίνακας 3.5 Πειραματικά αποτελέσματα μηχανισμού εκτίμησης δυσκολίας προ	91
Πίνακας 3.6 Επίπεδα ομοιότητας μεταξύ προτάσεων φυσικής γλώσσας	93
Πίνακας 3.7 Απόδοση με βάση την μετρική ομοιότητας	93
Πίνακας 3.8 Σχήματα Βαθμολόγησης Ασκήσεων Μετατροπής ΦΓ σε ΚΛΠΤ.....	95
Πίνακας 3.9 Συνολική βαθμολογία εκφράσεων ΚΛΠΤ	96
Πίνακας 3.10 Συνολική βαθμολογία Εκφράσεων ΚΛΠΤ.....	96
Πίνακας 5.1 Λέξεις ποσοτικοποίησης και βαθμός επίδρασης τους.....	135
Πίνακας 5.2 Συνολική απόδοση των ταξινομητών	144
Πίνακας 5.3 Ανάλυση της απόδοσης των βασικών ταξινομητών σε κάθε είδος κειμενικών δεδομένων	145
Πίνακας 5.4 Ανάλυση της απόδοσης των ομαδοποιημένων ταξινομητών.....	147

<i>Πίνακας 5.5 Υπολογισμός Cohen's Kappa.....</i>	<i>149</i>
<i>Πίνακας 5.6. Αποτελέσματα Cohen's kappa μετρικής.....</i>	<i>150</i>
<i>Πίνακας 5.7 Η απόδοση των ταξινομητών ως προς τη συναισθηματική πολικότητα.....</i>	<i>151</i>
<i>Πίνακας 5.8 Ανάλυση της απόδοσης των βασικών ταξινομητών σε κάθε είδος κειμενικών δεδομένων</i>	<i>152</i>
<i>Πίνακας 5.9 Ανάλυση της απόδοσης των ομαδοποιημένων ταξινομητών σε κάθε είδος κειμενικών δεδομένων</i>	<i>154</i>
<i>Πίνακας 5.10. Cohen's kappa μετρική.....</i>	<i>155</i>
<i>Πίνακας 5.11 Συνολική Απόδοση των ταξινομητών.....</i>	<i>156</i>
<i>Πίνακας 5.12 Ανάλυση της απόδοσης των βασικών ταξινομητών σε κάθε είδος κειμενικών δεδομένων.</i>	<i>157</i>
<i>Πίνακας 5.13 Ανάλυση της απόδοσης των ταξινομητών σε κάθε μια συναισθηματική κατηγορία.</i>	<i>157</i>
<i>Πίνακας 5.14. Cohen's kappa μετρική.....</i>	<i>158</i>

1

Εισαγωγή

Ένα από τα βασικότερα χαρακτηριστικά των υπολογιστικών συστημάτων είναι η δυνατότητα αλληλεπίδρασης που προσφέρουν στον χρήστη. Για να γίνει η επικοινωνία ανθρώπου - υπολογιστή ποιοτικότερη και πιο φυσική, είναι απαραίτητο τα υπολογιστικά συστήματα να έχουν την ικανότητα να αναγνωρίζουν και να προσαρμόζονται στις καταστάσεις του χρήστη και να διαμορφώνουν κατάλληλα τις ενέργειές τους.

Τα τελευταία χρόνια η ραγδαία ανάπτυξη της τεχνολογίας και η εξάπλωση του διαδικτύου έχουν αλλάξει τον τρόπο με τον οποίο το εκπαιδευτικό περιεχόμενο και η διαδικασία της μάθησης παρέχεται στους εκπαιδευόμενους και έχουν αναδείξει νέα μέσα που μπορούν να συμβάλουν σε αυτήν. Τα Ευφυή Συστήματα Διδασκαλίας (ΕΣΔ) (Intelligent Tutoring Systems - ITSs) αποτελούν την πιο δημοφιλή κατηγορία εκπαιδευτικών συστημάτων και έχουν ως στόχο να αυξήσουν την αποτελεσματικότητα της μάθησης.

Μια βασική πτυχή των ΕΣΔ αφορά την αλληλεπίδρασή τους με τους εκπαιδευόμενους και τη δυνατότητα να παρέχουν αποτελεσματική καθοδήγηση μέσω ανάδρασης κατά τη διάρκεια της μαθησιακής διαδικασίας. Η ανάδραση (feedback) στα ΕΣΔ ορίζεται ως οποιοδήποτε μήνυμα εμφανίζει το σύστημα στον εκπαιδευόμενο, είτε μετά από αίτημά του για βοήθεια είτε μετά από μια ενέργειά του, και αποτελεί μια βασική πτυχή των ΕΣΔ, που συμβάλει στο να γίνει πιο αποτελεσματική η διδασκαλία (Mory, 2004). Η ανάδραση εξαρτάται από πολλές παραμέτρους και μπορεί να διαφοροποιείται ως προς το χρόνο παράδοσης, την ποσότητα της παρεχόμενης πληροφορίας καθώς και τον τρόπο που παρέχεται (Shute, 2008; VanLehn, 2011). Η παροχή σημασιολογικά πλούσιας και εξατομικευμένης ανάδρασης είναι μια ιδιαίτερα σύνθετη και πολύπλοκη διαδικασία, όπου

απαιτείται η εις βάθος ανάλυση των ενεργειών των εκπαιδευόμενων καθώς και η αναλυτική καταγραφή της γνώσης του πεδίου.

Μια άλλη βασική παράμετρος αφορά την προσαρμογή των ΕΣΔ, αλλά και εν γένει των υπολογιστικών συστημάτων, στη συναισθηματική κατάσταση του χρήστη, κάτι που απαιτεί την αναγνώριση της συναισθηματικής του κατάστασης. Η αναγνώριση συναισθημάτων από ένα υπολογιστικό σύστημα είναι μια σύνθετη διαδικασία και αφορά κυρίως την προσπάθεια να αναλύσει και να προσδιορίσει τη συναισθηματική κατάσταση ενός χρήστη που έρχεται σε επαφή με αυτό. Το πεδίο της συναισθηματικής υπολογιστικής σχετίζεται με την μελέτη και ανάπτυξη μεθόδων και συστημάτων τα οποία μπορούν να αναγνωρίσουν, να επεξεργαστούν και να εκφράσουν τα ανθρώπινα συναισθήματα. Δεδομένου ότι η ανθρώπινη επικοινωνία δεν γίνεται μόνο με τον λόγο, αλλά και έμμεσα, μέσω εκφράσεων του προσώπου, μέθοδοι και συστήματα που θα μπορούσαν να εντοπίσουν, να επεξεργαστούν και να αναγνωρίσουν το συναισθηματικό περιεχόμενο των εκφράσεων του προσώπου θα ήταν ιδιαίτερα χρήσιμα και θα μπορούσαν να συμβάλουν αποτελεσματικά όχι μόνο στα εκπαιδευτικά συστήματα, αλλά και σε διάφορα άλλα πεδία και εφαρμογές.

Βασικό στόχο της παρούσας διατριβής αποτελεί η ανάπτυξη μηχανισμών που μπορούν να συμβάλλουν στην καλύτερη αλληλεπίδραση μεταξύ ανθρώπου και υπολογιστικών συστημάτων καθώς και η βελτίωση της ανάδρασης σε ευφυή εκπαιδευτικά συστήματα με στόχο να γίνει πιο αποτελεσματική η διαδικασία της μάθησης. Στην συνέχεια παρουσιάζεται αναλυτικά η συνεισφορά της διατριβής σε αυτά τα πεδία.

1.1 Αντικείμενα και συνεισφορά της Εργασίας

Αρχικά, ένα αντικείμενο της παρούσας εργασίας αποτελεί η αλληλεπίδραση του χρήστη και η παροχή ανάδρασης (feedback) και καθοδήγησης (guidance) στα πλαίσια των ευφών εκπαιδευτικών συστημάτων (intelligent tutoring systems). Αξίζει να σημειωθεί ότι στη διεθνή βιβλιογραφία δεν υπάρχουν αναφορές σε γενικές μεθοδολογίες που να μοντελοποιούν τη διαδικασία της δημιουργίας ανάδρασης. Για τον σκοπό αυτό, εισάγεται μια γενική μεθοδολογία που περιλαμβάνει μια συστηματική και αναλυτική μοντελοποίηση των τύπων των ενεργειών και των απαντήσεων των εκπαιδευόμενων στα πλαίσια της ενασχόλησής τους με διάφορες εκπαιδευτικές δραστηριότητες. Επίσης, εισάγεται ένα

ιεραρχικό μοντέλο οργάνωσης της ανάδρασης σε επίπεδα, προχωρώντας από γενικότερη σε ειδικότερη ανάδραση-βοήθεια στο χρήστη. Στα πλαίσια εφαρμογής της μεθόδου για την παροχή ανάδρασης, σχεδιάστηκε και αναπτύχθηκε ένα ευφυές έμπειρο σύστημα βασισμένο σε κανόνες, το οποίο αναλύει την συμπεριφορά του εκπαιδευόμενου και καθορίζει το κατάλληλο επίπεδο ανάδρασης προς παροχή από το σύστημα. Είναι η πρώτη φορά στη διεθνή βιβλιογραφία που εισάγεται μια τέτοια μεθοδολογία.

Ένα δεύτερο αντικείμενο της διατριβής αποτελεί η ανάπτυξη τεχνικών ανάλυσης κειμένου, που έχουν ως στόχο την εξαγωγή πληροφορίας σχετικά με το συναισθηματικό περιεχόμενο που φέρει (sentiment analysis). Για το σκοπό αυτό, μια πρώτη συνεισφορά της διατριβής αποτελεί ο σχεδιασμός μιας προσέγγισης και η ανάπτυξη αντίστοιχου εργαλείου, που εκτελεί εις βάθος ανάλυση του κειμένου και προσδιορίζει το συναισθηματικό περιεχόμενό του με βάση τη συντακτική δομή του και λεξικολογικούς πόρους. Το εργαλείο που αναπτύχθηκε δεν απαιτεί εκπαίδευση από δεδομένα και μπορεί να χρησιμοποιηθεί αποτελεσματικά και άμεσα σε μεγάλο πλήθος κειμενικών δεδομένων. Στη συνέχεια, μια ακόμη συνεισφορά της διατριβής αποτελεί ο σχεδιασμός και η ανάπτυξη ομαδοποιημένων σχημάτων ταξινόμησης (ensemble classifiers), που συνδυάζουν το προηγούμενο εργαλείο με μεθόδους μηχανικής μάθησης, καθώς και η μελέτη της απόδοσής τους σε διάφορα είδη κειμενικών δεδομένων. Στόχος των ομαδοποιημένων σχημάτων αποτελεί η αποτελεσματική σύνθεση διαφορετικών ειδών ταξινομητών με στόχο τόσο την αύξηση της απόδοσης της διεργασίας της ταξινόμησης όσο και την αποτελεσματική γενίκευση σε μεγάλο πλήθος διαφορετικών ειδών κειμενικών δεδομένων. Τα ομαδοποιημένα σχήματα παρουσίασαν μια αύξηση της απόδοσης της τάξης του 5-10% σε σχέση με τον καλύτερο ταξινομητή. Επίσης, μια ακόμη συνεισφορά της διατριβής αποτελεί η εισαγωγή ενός γενικού πλαισίου αναπαράστασης του συναισθηματικού περιεχομένου της συζήτησης σε θεματική περιοχή. Για τον σκοπό αυτό, παρουσιάζεται μια προσέγγιση που δημιουργεί το συναισθηματικό γράφημα μιας συζήτησης και το οποίο παρέχει ένα δομημένο τρόπο αναπαράστασης των συναισθηματικών καταστάσεων με βάση τις επιμέρους αναλύσεις του συναισθηματικού περιεχομένου κάθε μέρους της συζήτησης.

Ένα τρίτο βασικό αντικείμενο της διατριβής αποτελεί η κατανόηση κειμένου και η αυτόματη μετατροπή προτάσεων φυσικής γλώσσας (ΦΓ) σε κατηγορηματική λογική πρώτης τάξης (ΚΛΠΤ). Για τον σκοπό αυτό, μια σημαντική συνεισφορά της διατριβής αποτελεί η εισαγωγή μιας νέας, καινοτόμου μεθοδολογίας αυτόματης μετατροπής φυσικής

γλώσσας σε ΚΛΠΤ, η οποία βασίζεται στην συλλογιστική βασισμένη στις περιπτώσεις (case-based reasoning) και στην αρχή ότι προτάσεις που έχουν ίδια ή παρόμοια συντακτική δομή και δέντρα εξαρτήσεων (dependency trees) θα έχουν ίδια ή παρόμοια δομή έκφρασης σε ΚΛΠΤ.

Μία ακόμη συνεισφορά της διατριβής αποτελεί ο σχεδιασμός και η ανάπτυξη μεθοδολογίας για τον προσδιορισμό της δυσκολίας μετατροπής προτάσεων ΦΓ σε ΚΛΠΤ, που βασίζεται σε χαρακτηριστικά της πρότασης ΦΓ καθώς και σε χαρακτηριστικά της αντίστοιχης πρότασης ΚΛΠΤ. Στην συνέχεια, μια ακόμη συνεισφορά αποτελεί ο σχεδιασμός προσέγγισης και η ανάπτυξη αντίστοιχου εργαλείου που έχει ως στόχο την αυτόματη ανάλυση προτάσεων φυσικής γλώσσας και τη δημιουργία λογικά ισοδύναμων παραφράσεων τους, οι οποίες διατηρούν το σημασιολογικό τους περιεχόμενο. Για το σκοπό αυτό, σχεδιάστηκαν και αναπτύχθηκαν τεχνικές παράφρασης, οι οποίες μπορούν να μετατρέπουν δομικά μέρη των δέντρων εξάρτησης των προτάσεων σε ισοδύναμες μορφές, καθώς και να τροποποιούν κατάλληλα συγκεκριμένες λέξεις των προτάσεων.

Ένα τελευταίο αντικείμενο της διατριβής αποτελεί η ανάλυση εκφράσεων προσώπου και η αυτόματη αναγνώριση του συναισθηματικού τους περιεχομένου. Για το σκοπό αυτό, μια συνεισφορά της διατριβής αποτελεί ο προσδιορισμός ενός πλαισίου μοντελοποίησης των εκφράσεων του προσώπου με την εξαγωγή κατάλληλων χαρακτηριστικών που αναπαριστούν την κατάσταση κάθε μέρους του προσώπου. Επίσης, μελετάται η απόδοση νεύρο-ασαφών προσεγγίσεων καθώς και μεθόδων μηχανικής μάθησης στην αναγνώριση συναισθημάτων. Με βάση την απόδοση τους, σχεδιάστηκε ένα ομαδοποιημένο σχήμα ταξινόμησης. Στόχος του ομαδοποιημένου σχήματος αποτελεί τόσο η αύξηση της απόδοσης όσο και η αποτελεσματική γενίκευση σε νέα δεδομένα εκφράσεων προσώπου. Το ομαδοποιημένο σχήμα παρουσίασε μια αύξηση της απόδοσης της τάξης του 3.5% σε σχέση με την απόδοση του καλύτερου ταξινομητή και επίσης η μελέτη σε γνωστές βάσεις δεδομένων ανέδειξε πολύ καλά αποτελέσματα.

1.2 Διάρθρωση της Διατριβής

Η διδακτορική διατριβή αποτελείται συνολικά από 7 κεφάλαια. Στο παρόν κεφάλαιο γίνεται μια σύντομη παρουσίαση των αντικειμένων και της συνεισφοράς της διατριβής.

Στο κεφάλαιο 2, γίνεται παρουσίαση της μεθοδολογίας για την μοντελοποίηση της ανάδρασης σε ευφυή εκπαιδευτικά συστήματα. Πιο συγκεκριμένα, γίνεται αναλυτική παρουσίαση του κύκλου λειτουργίας της και ανάλυση κάθε βασικού μηχανισμού της. Επίσης, παρουσιάζεται η εφαρμογή του πλαισίου ανάδρασης στο ευφυές εκπαιδευτικό σύστημα NLtoFOL, το οποίο διδάσκει τη διαδικασία μετατροπής προτάσεων φυσικής γλώσσας σε κατηγορηματική λογική στο πεδίο της λογικής και παρουσιάζονται τα αποτελέσματα που προέκυψαν στα πλαίσια της πειραματικής μελέτης.

Στο κεφάλαιο 3, γίνεται παρουσίαση της προσέγγισης για τον φορμαλισμό προτάσεων φυσικής γλώσσας (ΦΓ) και την αυτόματη μετατροπή τους σε κατηγορηματική λογική πρώτης τάξης (ΚΛΠΤ). Πιο συγκεκριμένα, γίνεται ανάλυση της μεθοδολογίας, η οποία βασίζεται στον συλλογισμό με βάση περιπτώσεις (case-based reasoning), και γίνεται αναλυτική παρουσίαση του τρόπου λειτουργίας της. Επίσης, παρουσιάζεται η προσέγγιση για τον προσδιορισμό της δυσκολίας μετατροπής προτάσεων (ΦΓ) σε ΚΛΠΤ και αναλύεται ο τρόπος λειτουργίας της. Τέλος, παρουσιάζεται η πειραματική μελέτη και τα αποτελέσματα της αποτίμησης της απόδοσης της μεθοδολογίας φορμαλισμού και της προσέγγισης προσδιορισμού της δυσκολίας φορμαλισμού. Τα πειραματικά αποτελέσματα δείχνουν ότι η μέθοδος είναι αποτελεσματική και έχει πολύ καλή απόδοση στον φορμαλισμό προτάσεων και ιδίως σε ορθά συντακτικά δομημένες προτάσεις.

Στο κεφάλαιο 4, παρουσιάζεται η μεθοδολογία για την αναγνώριση και την παραγωγή παραφράσεων προτάσεων φυσικής γλώσσας. Η μεθοδολογία βασίζεται στην μορφοσυντακτική ανάλυση του κειμένου και χρησιμοποιεί ένα σύνολο τεχνικών παράφρασης οι οποίες μπορούν να τροποποιήσουν τη δομή της πρότασης και να αλλάξουν συγκεκριμένες λέξεις με στόχο να δημιουργήσουν σημασιολογικά ισοδύναμες προτάσεις (παραφράσεις). Επίσης, παρουσιάζεται η πειραματική μελέτη αξιολόγησης της απόδοσης της προσέγγισης στη δημιουργία συντακτικά ορθών παραφράσεων, που διατηρούν το νόημα των προτάσεων, από όπου προκύπτει ότι η μεθοδολογία είναι ακριβής και λειτουργεί αποτελεσματικά ακόμη και σε πολύπλοκες προτάσεις.

Στο κεφάλαιο 5, παρουσιάζεται η μεθοδολογία για την αυτόματη ανάλυση κειμένου και την αναγνώριση του συναισθηματικού περιεχομένου που φέρει. Αρχικά, παρουσιάζεται η προσέγγιση που αναπτύχθηκε, και βασίζονται στην εις βάθος ανάλυση της φυσικής γλώσσας. Στην συνέχεια, γίνεται παρουσίαση των σχημάτων ομαδοποιημένων ταξινομητών (ensemble classifiers), που συνδυάζουν το εργαλείο με μεθόδους μηχανικής μάθησης και

της μελέτης της απόδοσής τους στην αναγνώριση συναισθηματικού περιεχομένου διαφόρων ειδών κειμένων. Επίσης, παρουσιάζεται ένα γενικό πλαίσιο μοντελοποίησης και αναπαράστασης του συναισθηματικού περιεχομένου συζητήσεων και αναλύεται η λειτουργία του σε διάφορες θεματικές περιοχές

Στο κεφάλαιο 6, παρουσιάζεται η μεθοδολογία για την ανάλυση εκφράσεων προσώπου και τον προσδιορισμό της συναισθηματικής τους κατάστασης. Στα πλαίσια της μεθοδολογίας που αναπτύχθηκε, αρχικά γίνεται μια αναλυτική (analytical local-based) προσέγγιση εξαγωγής χαρακτηριστικών του ανθρώπινου προσώπου από περιοχές ενδιαφέροντος. Οι περιοχές αυτές είναι τα μάτια, τα φρύδια και το στόμα, τα οποία αναλύονται και εξάγονται χαρακτηριστικά τα οποία μοντελοποιούν το σχήμα τους. Στην συνέχεια, με βάση τα χαρακτηριστικά, ταξινομητές που βασίζονται σε νευρο-ασαφείς μεθόδους και μεθόδους μηχανικής μάθησης εκπαιδεύονται και χρησιμοποιούνται για την αναγνώριση των εκφράσεων του προσώπου, καθώς επίσης και ομαδοποιημένα σχήματα ταξινομητών. Το κεφάλαιο ολοκληρώνεται με την παρουσίαση των πειραματικών αποτελεσμάτων σε γνωστές βάσεις προσώπων, που χρησιμοποιούνται για την αποτίμηση της απόδοσης συστημάτων.

Τέλος, στο κεφάλαιο 7 παρουσιάζονται τα συμπεράσματα που προέκυψαν και γίνεται παρουσίαση των κατευθύνσεων που θα μπορούσε να εστιάσει η μελλοντική έρευνα πάνω στα θέματα που διερευνώνται στην παρούσα διδακτορική διατριβή.

2

Μοντελοποίηση Ανάδρασης σε Ευφυή Εκπαιδευτικά Συστήματα

2.1 Εισαγωγή

Η ραγδαία ανάπτυξη της εκπαιδευτικής τεχνολογίας και η εξάπλωση του διαδικτύου έχουν αλλάξει τον τρόπο με τον οποίο το εκπαιδευτικό περιεχόμενο και η διαδικασία της μάθησης παρέχεται στους εκπαιδευόμενους και έχουν αναδείξει νέα μέσα που μπορούν να συμβάλουν στην εκπαιδευτική διαδικασία. Τα ευφυή Συστήματα Διδασκαλίας (ΕΣΔ) (Intelligent Tutoring Systems - ITSs) αποτελούν την πιο δημοφιλή και ευρέως χρησιμοποιούμενη κατηγορία εκπαιδευτικών συστημάτων, τα οποία έχουν ως στόχο να αυξήσουν την αποτελεσματικότητά της μάθησης και της διδασκαλίας. Τα ευφυή συστήματα διδασκαλίας έχουν χρησιμοποιηθεί με μεγάλη επιτυχία σε πολλές περιοχές για να προσφέρουν εξατομικευμένη μάθηση και να παρέχουν τη δυνατότητα στους εκπαιδευόμενους να μελετήσουν και να μάθουν αποτελεσματικά με τον δικό τους τρόπο και ρυθμό (Mitrovic et al., 2007; Mitrovic et al., 2009; VanLehn et al., 2005).

Ένα από τα κύρια χαρακτηριστικά των ΕΣΔ είναι ότι μπορούν να προσαρμόζουν την εκπαιδευτική διαδικασία στις ανάγκες του κάθε εκπαιδευόμενου με σκοπό να μεγιστοποιήσουν τα μαθησιακά αποτελέσματα. Αυτό επιτυγχάνεται κυρίως με τη χρήση μεθόδων τεχνητής νοημοσύνης για να αναπαραστήσουν το παιδαγωγικό μοντέλο, τη γνώση του πεδίου διδασκαλίας, το μοντέλο του μαθητή και τις διαδικασίες αξιολόγησης (Shute & Zapata-Rivera, 2010). Μια βασική πτυχή των ευφύων συστημάτων διδασκαλίας αφορά την αλληλεπίδρασή τους με τους εκπαιδευόμενους και τη δυνατότητα να παρέχουν

αποτελεσματική βοήθεια και καθοδήγηση στους εκπαιδευόμενους κατά τη διάρκεια της μαθησιακής διαδικασίας.

Η ανάδραση (feedback) στα ευφυή συστήματα διδασκαλίας ορίζεται ως οποιοδήποτε μήνυμα εμφανίζει το σύστημα στον εκπαιδευόμενο μετά από αίτημά του για βοήθεια ή μετά από μια ενέργειά του και αποτελεί μια βασική πτυχή των ΕΣΔ, που συμβάλει στο να γίνει πιο αποτελεσματική η μάθηση και η διδασκαλία (Wager & Wager, 1985; Mory, 2004). Η ανάδραση εξαρτάται από πολλές παραμέτρους και μπορεί να διαφοροποιείται ως προς το χρόνο παράδοσης, την ποσότητα της παρεχόμενης πληροφορίας καθώς και τον τρόπο που παρέχεται (McKendree, 1990). Η παροχή σημασιολογικά πλούσιας και εξατομικευμένης ανάδρασης είναι μια ιδιαίτερα σύνθετη και πολύπλοκη διαδικασία όπου απαιτείται η εις βάθος ανάλυση των ενεργειών των εκπαιδευόμενων καθώς και η αναλυτική καταγραφή της γνώσης του πεδίου.

Στα πλαίσια της ενότητας αυτής παρουσιάζουμε μια νέα γενική μεθοδολογία για την μοντελοποίηση της ανάδρασης σε ευφυή εκπαιδευτικά συστήματα. Αξίζει να σημειωθεί ότι στη διεθνή βιβλιογραφία δεν υπάρχουν αναφορές σε γενικές μεθοδολογίες που να μοντελοποιούν τη διαδικασία της δημιουργίας ανάδρασης. Η μεθοδολογία που αναπτύχθηκε βασίζεται στη μοντελοποίηση της απάντησης του εκπαιδευόμενου, στη δημιουργία και εισαγωγή ενός ιεραρχικού μοντέλου οργάνωσης της ανάδρασης καθώς επίσης και στην κατηγοριοποίηση των λαθών του πεδίου εφαρμογής. Πιο συγκεκριμένα, η μεθοδολογία που εισάγεται περιλαμβάνει μια γενική, συστηματική και αναλυτική μοντελοποίηση των τύπων των απαντήσεων των εκπαιδευόμενων στα πλαίσια της ενασχόλησης με διαφόρων ειδών εκπαιδευτικές δραστηριότητες και διαδικασίες. Επίσης, εισάγει ένα ιεραρχικό μοντέλο, που δομεί την ανάδραση σε επίπεδα και όπου κάθε επίπεδο παρέχει συγκεκριμένα είδη ανάδρασης. Επίσης, περιλαμβάνει ένα μηχανισμό για την κατηγοριοποίηση των λαθών ανάλογα με το πεδίο της εφαρμογής, με στόχο την αναγνώριση των τύπων των λαθών που έγιναν από τους εκπαιδευόμενους. Τέλος, παρουσιάζουμε την εφαρμογή του πλαισίου ανάδρασης στο εκπαιδευτικό σύστημα NLtoFOL, καθώς και τα πειραματικά αποτελέσματα που προέκυψαν από την αξιολόγηση του πλαισίου σχετικά με την αποτελεσματικότητά του στην διαδικασία της μάθησης των εκπαιδευόμενων.

2.1.1 Γενικό Πλαίσιο Ανάδρασης

Η ανάδραση αποτελεί μια βασική πτυχή της εκπαιδευτικής διαδικασίας και αφορά κάθε μήνυμα και πληροφορία που παρουσιάζει το σύστημα στον εκπαιδευόμενο μετά από αίτημα του ή μετά από απόφαση του συστήματος σχετικά με την βελτίωση της απόδοσης του εκπαιδευόμενου (Mogy, 2004). Το γενικό πλαίσιο παροχής ανάδρασης που σχεδιάστηκε και αναπτύχθηκε έχει ως στόχο να μοντελοποιήσει την ανάδραση που προσφέρεται στα πλαίσια των ευφύων εκπαιδευτικών συστημάτων και να εισάγει ένα συστηματικό τρόπο εφαρμογής της δημιουργίας ανάδρασης. Αξίζει να τονίσουμε ότι στη διεθνή βιβλιογραφία δεν υπάρχουν εργασίες και συγκεκριμένες μεθοδολογίες για την μοντελοποίηση της διαδικασίας δημιουργίας ανάδρασης. Επίσης, όσες εργασίες υπάρχουν, πραγματεύονται μόνο την διαδικασία της ακολουθία παροχής ανάδρασης αναλύοντας την σειρά με την οποία θα παρέχεται το κάθε είδος ανάδρασης (Vanlehn, 2006, Vanlehn, 2011). Στα πλαίσια της παρούσας διατριβής, γίνεται μια λεπτομερής ανάλυση της λειτουργίας της ανάδρασης, που περιλαμβάνει την εισαγωγή μιας γενικής μεθόδου μοντελοποίησης της απάντησης των εκπαιδευόμενων, την εισαγωγή ενός γενικού ιεραρχικού μοντέλου οργάνωσης της ανάδρασης σε επίπεδα καθώς και του καθορισμού μιας συγκεκριμένης κατηγοριοποίησης των λαθών του πεδίου εφαρμογής.

2.1.2 Μοντελοποίηση Απάντησης Εκπαιδευόμενων

Η ανάλυση της απάντησης του εκπαιδευόμενου αποτελεί μια πολύ σημαντική διεργασία που συμβάλλει σημαντικά στην παραγωγή της κατάλληλης ανάδρασης. Πιο συγκεκριμένα, για να γίνει ακριβής η ανάλυση της απάντησης των εκπαιδευόμενων, σχεδιάστηκε ένα σχήμα μοντελοποίησης της απάντησης τους. Στα πλαίσια του σχήματος, ο όρος “στοιχείο” χρησιμοποιείται για να αναπαραστήσει ένα τμήμα πληροφορίας που πρέπει να προσδιορίσει ο εκπαιδευόμενος στην απάντηση του. Για τον σκοπό της καλύτερης και λεπτομερέστερης ανάλυσης της απάντησης του εκπαιδευόμενου, το σχήμα μοντελοποίησης της απάντησης του εκπαιδευόμενου χαρακτηρίζει την απάντηση ως προς την *πληρότητα* (completeness) και την *ακρίβεια* (accuracy) (Fiedler & Tsovaltzi, 2003). Μια απάντηση χαρακτηρίζεται ως *πλήρης* (complete) αν όλα τα στοιχεία της σωστής απάντησης εμφανίζονται στην απάντηση του εκπαιδευόμενου. Αντίθετα, όταν ο εκπαιδευόμενος δεν προσδιορίζει όλα τα στοιχεία της σωστής απάντησης, η απάντησή του χαρακτηρίζεται ως *ελλιπής* (incomplete). Η απάντηση του εκπαιδευόμενου χαρακτηρίζεται ως *ακριβής* (accurate), όταν όλα τα στοιχεία της απάντησης είναι σωστά, ενώ χαρακτηρίζεται ως *ανακριβής* (inaccurate) όταν ένα ή περισσότερα στοιχεία από την απάντηση του εκπαιδευόμενου δεν έχουν οριστεί σωστά.

Επίσης, το σχήμα μοντελοποίησης της απάντησης που σχεδιάστηκε περιέχει και τον χαρακτηρισμό της απάντησης ως *πλεονασματικής* (*superfluous*) (Gouli et al., 2006). Πιο συγκεκριμένα, ως *πλεονασματική* (*superfluous*) χαρακτηρίζεται η απάντηση του εκπαιδευόμενου όταν περιέχει όλα τα απαιτούμενα στοιχεία καθώς και ένα ή περισσότερα επιπλέον στοιχεία τα οποία δεν είναι απαραίτητα. Επομένως, η απάντηση ενός εκπαιδευόμενου μπορεί να μοντελοποιηθεί στις ακόλουθες κατηγορίες:

- *Incomplete – Accurate*: Μια απάντηση όπου δεν έχουν προσδιοριστεί ορισμένα στοιχεία από τον εκπαιδευόμενο, ωστόσο, τα στοιχεία που προσδιορίστηκαν είναι σωστά.
- *Incomplete – Inaccurate*: Μια απάντηση όπου δεν έχουν προσδιοριστεί ορισμένα στοιχεία από τον εκπαιδευόμενο και επίσης ένα ή περισσότερα από τα στοιχεία που προσδιορίστηκαν είναι λανθασμένα.
- *Complete – Inaccurate*: Μια απάντηση όπου όλα τα απαραίτητα στοιχεία της σωστής απάντησης εμφανίζονται, ωστόσο, ένα ή περισσότερα από τα στοιχεία που προσδιορίστηκαν είναι λανθασμένα.
- *Complete- Accurate*: Αποτελεί τη σωστή απάντηση, όπου έχουν καθοριστεί ορθά όλα τα στοιχεία της σωστής απάντησης.
- *Superfluous*: Ο εκπαιδευόμενος έχει προσδιορίσει ένα ή περισσότερα στοιχεία επί πλέον από όσα είναι απαραίτητα και επίσης η απάντηση είναι λανθασμένη. Η απάντηση μπορεί να αναλυθεί περαιτέρω ως εξής:
 - *Superfluous – Accurate*: τα απαιτούμενα στοιχεία προσδιορίστηκαν σωστά, όμως τα επιπλέον στοιχεία καθιστούν την απάντηση λανθασμένη.
 - *Superfluous–Inaccurate*: ένα ή περισσότερα από τα απαιτούμενα στοιχεία έχουν προσδιοριστεί λανθασμένα και επίσης έχουν προσδιοριστεί ένα ή περισσότερα επί πλέον στοιχεία από τα απαιτούμενα.

Σε περίπτωση που η απάντηση χαρακτηρίζεται ως *complete - accurate*, ο εκπαιδευόμενος έχει προσδιορίσει τη σωστή απάντηση ή μια σωστή απάντηση σε περίπτωση που περισσότερες από μία απαντήσεις είναι σωστές. Σε όλες τις άλλες περιπτώσεις, η απάντηση του εκπαιδευόμενου θεωρείται ως λανθασμένη και αναλύεται περαιτέρω από τη μονάδα αναγνώρισης λαθών προκειμένου να προσδιοριστεί το είδος των λαθών που έγιναν και να παραχθεί η κατάλληλη ανάδραση με βάση την ιεραρχία ανάδρασης.

2.1.3 Σχεδιασμός Ιεραρχίας Ανάδρασης

Η παροχή ανάδρασης και βοήθειας σε κάθε στάδιο της εκπαιδευτικής διαδικασίας αποτελεί μια βασική πτυχή ενός ευφυούς συστήματος διδασκαλίας (Gheorghiu & Vanlehn, 2008). Διαφορετικοί τύποι ανάδρασης μπορούν να έχουν διαφορετική επίδραση στη διαδικασία μάθησης των εκπαιδευόμενων (Koedinger & Aleven, 2007; Mory, 2004; Vasilyeva et al., 2008). Η παροχή κατάλληλης ανάδρασης όσον αφορά τη σημασιολογία και την ποσότητα της παρεχόμενης πληροφορίας αποτελεί μια σημαντική λειτουργία ενός ευφυούς συστήματος διδασκαλίας και οι κύριοι τύποι ανατροφοδότησης που χρησιμοποιούνται στο πλαίσιο παροχής ανάδρασης είναι οι εξής:

Minimal feedback: Αποτελεί το πιο απλό είδος ανάδρασης, το οποίο ενημερώνει τον εκπαιδευόμενο αν η απάντηση του είναι σωστή ή όχι.

Flag feedback: Ενημερώνει τον εκπαιδευόμενο ποια μέρη της απάντησης του είναι σωστά και ποια λάθος χρωματίζοντάς τα με πράσινο και κόκκινο χρώμα αντίστοιχα.

Knowledge about Concept: Ενημερώνει τον εκπαιδευόμενο σχετικά με την θεωρία των εννοιών και των διαδικασιών, καθώς και παρέχει υποδείξεις και επεξηγήσεις στις αντίστοιχες έννοιες.

Positive feedback: Ενημερώνει τον εκπαιδευόμενο για τα σωστά μέρη της απάντησής του και παρέχει αντίστοιχη αιτιολόγησή τους.

Procedural feedback: Παρέχει υποδείξεις στους εκπαιδευόμενους σχετικά με το πώς να προχωρήσουν προς την κατεύθυνση της λύσης.

Error-specific feedback: Όταν η απάντηση ενός εκπαιδευόμενου είναι λανθασμένη, ο μηχανισμός αναγνωρίζει τα λάθη που έγιναν και παρέχει την κατάλληλη ανάδραση με βάση τα λάθη.

Bottom-out hints: Ενημερώνει τον εκπαιδευόμενο σχετικά με την σωστή απάντηση ή ενός μέρους της σωστής απάντησης. Αυτό μπορεί να γίνει μετά από αίτημα του εκπαιδευόμενου ή μετά από συνεχείς αποτυχίες ανάλογα με τις συνθήκες της μάθησης.

Knowledge on meta-cognition: Το σύστημα αναλύει τις ενέργειες και την κατάσταση των εκπαιδευόμενων κατά τη διάρκεια της διαδικασίας της μάθησης και μπορεί να παρέχει μετά-γνωστική (meta-cognitive) καθοδήγηση και υποδείξεις.

Το πλαίσιο κατηγοριοποιεί την ανάδραση σε επίπεδα (levels) και δημιουργείται ένα ιεραρχικό μοντέλο, όπου κάθε επίπεδο παρέχει στους εκπαιδευόμενους συγκεκριμένους

τύπους ανάδρασης. Το ιεραρχικό μοντέλο έχει σχεδιαστεί να ακολουθεί μια κλιμακωτή πολιτική παροχής βοήθειας και ανάδρασης. Αρχικά, παρέχεται ανάδραση που σχετίζεται με το πρώτο επίπεδο και σταδιακά φτάνει στο τελικό επίπεδο, όπου παρέχεται η σωστή απάντηση. Πιο συγκεκριμένα, το ιεραρχικό μοντέλο αναλύει και καθορίζει την ανάδραση που θα παραχθεί και προσπαθεί να προσομοιώσει τον τρόπο που θα παρείχε ανάδραση ο διδάσκοντας. Ο διδάσκοντας θεωρείται ότι είναι εξαιρετικά αποτελεσματικός για την παροχή βοήθειας και καθοδήγησης στους εκπαιδευόμενους κατά την διάρκεια των εκπαιδευτικών δραστηριοτήτων και μια βασική πτυχή της συμπεριφοράς του είναι ότι δεν αποκαλύπτει τη σωστή απάντηση άμεσα αλλά προσπαθεί να εμπλέξει τους εκπαιδευόμενους στην διαδικασία της μάθησης συμβάλλοντας στο να γίνουν πιο ενεργοί. Πιο συγκεκριμένα, η βοήθεια του συστήματος παρέχεται σε δύο κύρια στάδια: (i) πριν την υποβολή μιας απάντησης από τον εκπαιδευόμενο σε μια εκπαιδευτική διαδικασία και (ii) μετά την υποβολή λανθασμένης απάντησης.

2.1.3.1 Ανάδραση πριν την υποβολή απάντησης

Αρχικά, ο εκπαιδευόμενος μπορεί να ζητήσει την βοήθεια του συστήματος κατά την διάρκεια της εκπαιδευτικής δραστηριότητας, πριν υποβάλλει την απάντησή του. Το σχήμα βοήθειας για την παραγωγή των κατάλληλων υποδείξεων πριν υποβάλλει μια απάντηση ο εκπαιδευόμενος αποτελείται από δύο επίπεδα. Το κάθε επίπεδο συσχετίζεται με συγκεκριμένους τύπους ανάδρασης και ο κάθε τύπος συνδέεται με τους κατάλληλους τύπους υποδείξεων. Τα επίπεδα βοήθειας, οι τύποι ανάδρασης και οι υποδείξεις παρουσιάζονται στον Πίνακα 2.1

Πίνακας 2.1 Οργάνωση ανάδρασης πριν την υποβολή απάντησης

Επίπεδο Βοήθειας	Είδος Ανάδρασης	Επεξηγήσεις
Πρώτο Επίπεδο	Knowledge about Concept	Υποδείξεις σχετικά με την θεωρία του θεματικού. Υποδείξεις & εξηγήσεις για ορισμούς εννοιών.
Δεύτερο Επίπεδο	Procedural feedback	Υποδείξεις για το τι θα πρέπει να γίνει προς σχηματισμό της απάντησης. Παρουσίαση παρόμοιων σχετικών ασκήσεων που έχουν ολοκληρωθεί από τον εκπαιδευόμενο.

Στο πρώτο επίπεδο ανάδρασης, παρέχονται υποδείξεις και ανάδραση τύπου *knowledge about concept* (KC). Οι υποδείξεις αυτές σχετίζονται με χαρακτηριστικά των θεμάτων της θεωρίας με στόχο να ενημερώσουν τους εκπαιδευόμενους σχετικά με το θεωρητικό πλαίσιο (ορισμοί, σύνταξη κ.λπ.) της εκπαιδευτικής δραστηριότητας.

Στο δεύτερο επίπεδο, παρέχονται στον εκπαιδευόμενο υποδείξεις και ανάδραση τύπου *procedural feedback*. Οι υποδείξεις αυτές έχουν ως στόχο να υπενθυμίσουν στους εκπαιδευόμενους τους αντίστοιχους στόχους και επίσης να δώσουν συγκεκριμένη καθοδήγηση για το πώς θα επιτευχθούν αυτοί. Ο κύριος στόχος είναι να εστιάσει την προσοχή των εκπαιδευόμενων σε συγκεκριμένα σημεία της εκπαιδευτικής διαδικασίας που είναι δύσκολα και μπορούν να τον οδηγήσουν σε μεγάλο πλήθος λαθών και μεγάλο ποσοστό αποτυχημένων προσπαθειών. Επιπλέον, αυτή η ιεραρχία ανάδρασης μπορεί να παρέχει στους εκπαιδευόμενους παρόμοια παραδείγματα ανάλογα με την εκπαιδευτική διαδικασία. Οι υποδείξεις και η καθοδήγηση των επιπέδων της ανάδρασης παρέχονται μετά από αίτημα του εκπαιδευόμενου για βοήθεια.

2.1.3.2 Ανάδραση μετά την υποβολή λανθασμένης απάντησης

Το δεύτερο στάδιο της ανάδρασης παρέχεται μετά από μια ενέργεια του εκπαιδευόμενου η οποία δεν είναι σωστή. Το σχήμα της ανάδρασης αποτελείται από τέσσερα επίπεδα, όπως παρουσιάζονται στον Πίνακα 2.2.

Πίνακας 2.2 Οργάνωση ανάδρασης μετά την υποβολή απάντησης

Επίπεδο Βοήθειας	Είδος ανάδρασης	Επεξηγήσεις
Πρώτο Επίπεδο	Flag feedback	Σταδιακή ενημέρωση για την ορθότητα κάθε ενός στοιχείου της απάντησης του εκπαιδευόμενου.
		Ενημερώσεις για την ορθότητα όλων των στοιχείων της απάντησης του εκπαιδευόμενου.
Δεύτερο Επίπεδο	Positive feedback	Αναλύσεις και επεξηγήσεις για τα σωστά στοιχεία.
	Error-specific feedback	Αναλύσεις και επεξηγήσεις για τα λανθασμένα στοιχεία.
Τρίτο Επίπεδο	Procedural feedback	Υποδείξεις πώς να εργαστεί για την σωστή απάντηση.
		Παρουσίαση παρόμοιων εκπαιδευτικών διαδικασιών που ολοκλήρωσε με επιτυχία στο παρελθόν.

Τέταρτο Επίπεδο	Bottom-out	Σταδιακή παρουσίαση ενός στοιχείου της σωστής απάντησης κάθε φορά. Παρουσίαση της σωστής απάντησης.
--------------------	------------	--

Το πρώτο επίπεδο ενημερώνει τους εκπαιδευόμενους για την ορθότητα της απάντησής τους χρωματίζοντας με πράσινο χρώμα τα σωστά μέρη της απάντησης και με κόκκινο τα λανθασμένα στοιχεία. Αυτό ο τύπος της ανάδρασης (*flag feedback*) μπορεί να προσφέρεται στους εκπαιδευόμενους με δύο διαφορετικούς τρόπους. Αρχικά, ο πρώτος τρόπος αφορά τον χρωματισμό ενός μόνο στοιχείου της απάντησης με το κατάλληλο χρώμα και στην συνέχεια ο εκπαιδευόμενος μπορεί να ζητήσει επιπρόσθετες υποδείξεις και για τα υπόλοιπα στοιχεία. Ο δεύτερος τρόπος αφορά τον χρωματισμό όλων των στοιχείων της απάντησης με το κατάλληλο χρώμα και με τον τρόπο αυτό ο εκπαιδευόμενος λαμβάνει την πλήρη λειτουργικότητα της ανάδραση τύπου *flag feedback* σε όλα τα στοιχεία της απάντησής του.

Στο δεύτερο επίπεδο παρέχονται αιτιολογήσεις για τα σωστά στοιχεία της απάντησης και οι κατάλληλες εξηγήσεις για τα λάθη που έγιναν από τον εκπαιδευόμενο. Ο εκπαιδευόμενος έχει ήδη ενημερωθεί από την ανάδραση του πρώτου επιπέδου σχετικά με την ορθότητα του κάθε στοιχείου της απάντησής του και στο επίπεδο αυτό παρέχονται υποδείξεις αιτιολογώντας τα σωστά στοιχεία και εξηγώντας τα λάθη που έγιναν.

Στο τρίτο επίπεδο της βοήθειας παρέχονται στον εκπαιδευόμενο υποδείξεις τύπου *procedural hints* και πιο συγκεκριμένα, παρέχονται υποδείξεις σχετικά με την αντιμετώπιση συγκεκριμένων σφαλμάτων που έγιναν και το πώς να διορθώσει την απάντησή του. Επιπλέον, με βάση το ιστορικό των προηγούμενων απαντήσεων του εκπαιδευόμενου παρέχονται παρόμοιες ασκήσεις που έχει ολοκληρώσει επιτυχώς.

Το τέταρτο επίπεδο βοήθειας παρέχει στον εκπαιδευόμενο τη σωστή απάντηση ή ένα μέρος αυτής. Η σωστή απάντηση μπορεί να προσφέρεται με δύο τρόπους. Αρχικά, το σύστημα μπορεί να αποκαλύψει ένα στοιχείο της σωστής απάντησης και με βάση αυτό ο εκπαιδευόμενος να προσπαθήσει να επιτύχει τη σωστή απάντηση. Ο δεύτερος τρόπος αφορά την παρουσίαση όλων των στοιχείων της σωστής απάντησης, οπότε ο εκπαιδευόμενος ενημερώνεται για όλα τα στοιχεία της σωστής απάντησης στην τρέχουσα εκπαιδευτική δραστηριότητα.

2.1.3.3 Μετα-γνωστική (Meta-cognitive) Ανάδραση

Στα περισσότερα ευφυή συστήματα διδασκαλίας, ο εκπαιδευόμενος έχει τον έλεγχο για την παροχή βοήθειας. Επομένως, ο εκπαιδευόμενος θα πρέπει να αποφασίσει πώς να χρησιμοποιήσει την υποστήριξη του συστήματος, γεγονός που προϋποθέτει ότι θα πρέπει να είναι σε θέση να έχει καλή γνώση σχετικά με το επίπεδο γνώσης του (Aleven & Koedinger, 2000). Ωστόσο, σε πολλές περιπτώσεις, ένας εκπαιδευόμενος δεν μπορεί να χρησιμοποιήσει τη βοήθεια των συστημάτων διδασκαλίας με τον κατάλληλο τρόπο, κάτι που μπορεί να μειώσει τις μαθησιακές δυνατότητες της διαδικασίας διδασκαλίας (Aleven et al., 2003; Aleven et al., 2006). Για το σκοπό αυτό, ένα ειδικό τμήμα της αλληλεπίδρασης του συστήματος με το εκπαιδευόμενο βασίζεται στη μεταγνωστική μαθησιακή συμπεριφορά (metacognitive learning behavior) του. Η μοντελοποίηση των μεταγνωστικών δεξιοτήτων (metacognitive skills) και η χρήση των μοντέλων αυτών στην υποστήριξη των εκπαιδευόμενων μπορεί να συμβάλει σε πιο αποτελεσματική και αποδοτική μάθηση (Mathan & Koedinger, 2005). Για το σκοπό αυτό, κατά τη διάρκεια της εκπαιδευτικής διαδικασίας, το πλαίσιο ανάδρασης καταγράφει όλες τις ενέργειες των εκπαιδευόμενων, τις αναλύει και προσδιορίζει τη μεταγνωστική (metacognitive) συμπεριφορά του κάθε εκπαιδευόμενου. Η μεταγνωστική συμπεριφορά αφορά τη μοντελοποίηση των μεταγνωστικών ενεργειών του εκπαιδευόμενου και μπορεί να βοηθήσει ένα ευφύες σύστημα διδασκαλίας να παρέχει κατάλληλες υποδείξεις και καθοδήγηση με βάση την συμπεριφορά του εκπαιδευόμενου. Οι βασικές περιπτώσεις μεταγνωστικής συμπεριφοράς του εκπαιδευόμενου παρουσιάζονται στον Πίνακα 2.3

Πίνακας 2.3 Παραδείγματα ανάδρασης με βάση την συμπεριφορά του εκπαιδευόμενου

Συμπεριφορά Εκπαιδευόμενου	Ανάδραση
Ο εκπαιδευόμενος εργάζεται για πάρα πολύ ώρα πάνω στην εκπαιδευτική διαδικασία χωρίς να έχει στείλει κάποια απάντηση.	Υπόδειξη να ζητήσει καθοδήγηση από το σύστημα.
Ο εκπαιδευόμενος ζητάει συνεχώς την βοήθεια του συστήματος χωρίς να έχει αποστείλει κάποια απάντηση (help abuse)	Υπόδειξη να προσπαθήσει να αποστείλει μια απάντηση στην τρέχουσα εκπαιδευτική δραστηριότητα.
Ο εκπαιδευόμενος δεν έχει ζητήσει κάποια βοήθεια από το σύστημα και δυσκολεύεται να προσδιορίσει	Υπόδειξη να ζητήσει καθοδήγηση από το σύστημα.

την σωστή απάντηση στην εκπαιδευτική
δραστηριότητα (help refusal)

Το σύστημα σε πολλές περιπτώσεις μπορεί να κάνει συστάσεις προς τον εκπαιδευόμενο να ζητήσει βοήθεια ή και να παρέχει στον εκπαιδευόμενο τη κατάλληλη υπόδειξη. Για παράδειγμα, όταν ο εκπαιδευόμενος επανειλημμένα αποτυγχάνει να προσδιορίσει ένα συγκεκριμένο στοιχείο της απάντησης, το σύστημα μπορεί να κάνει μια πρόταση για να ζητήσει την βοήθεια του συστήματος. Όλες οι προτάσεις που κάνει το πλαίσιο ανάδρασης και ο τρόπος που τις προσφέρει στον εκπαιδευόμενο καθορίζονται από ένα *έμπειρο σύστημα* βασισμένο σε κανόνες. Στα πλαίσια των εκπαιδευτικών δραστηριοτήτων των εκπαιδευόμενων όλες οι ενέργειες τους καταγράφονται και αποθηκεύονται στο σύστημα όπως για παράδειγμα ο αριθμός των απαντήσεων που υποβλήθηκαν, ο χρόνος που χρειάστηκε για την υποβολή κάθε απάντησης, τα αιτήματα για βοήθεια και παράμετροι σχετικά με την ανάλυση των λαθών που έγιναν. Το έμπειρο σύστημα παροχής υποδείξεων αποτελείται από δύο βασικές ομάδες κανόνων οι οποίες είναι, οι *κανόνες επικαιροποίησης προφίλ του χρήστη (profiling rules)* και οι *κανόνες συστάσεων (recommendation rules)*. Οι κανόνες δημιουργίας προφίλ προσδιορίζουν με βάση τις ενέργειες του εκπαιδευόμενου στοιχεία της συμπεριφοράς όπως για παράδειγμα αν ο εκπαιδευόμενος είναι ενεργός ή ανενεργός κατά την αλληλεπίδραση με το σύστημα και τις εκπαιδευτικές δραστηριότητες. Από την άλλη πλευρά, οι κανόνες συστάσεων (recommendation rules) χρησιμοποιούνται για να παρέχουν υποδείξεις και βοήθεια στους εκπαιδευόμενους με βάση την συμπεριφορά τους. Στην συνέχεια, παρουσιάζονται μερικά ενδεικτικά παραδείγματα κανόνων οι οποίοι για λόγους αναγνωσιμότητας, έχουν εκφραστεί σε φυσική γλώσσα. Ο πρώτος κανόνας R1 είναι ένας κανόνας επικαιροποίησης προφίλ, ενώ οι κανόνες R2 και R3 είναι δύο κανόνες παροχής συστάσεων.

R1: If student_creates_events is high

then student_is_working_on_the_exercise is yes

R2: If time_between_submission is high and answer_score is low and
hints_delivered are low and student_is_working_on_the_exercise

then take_hint_recommendation.

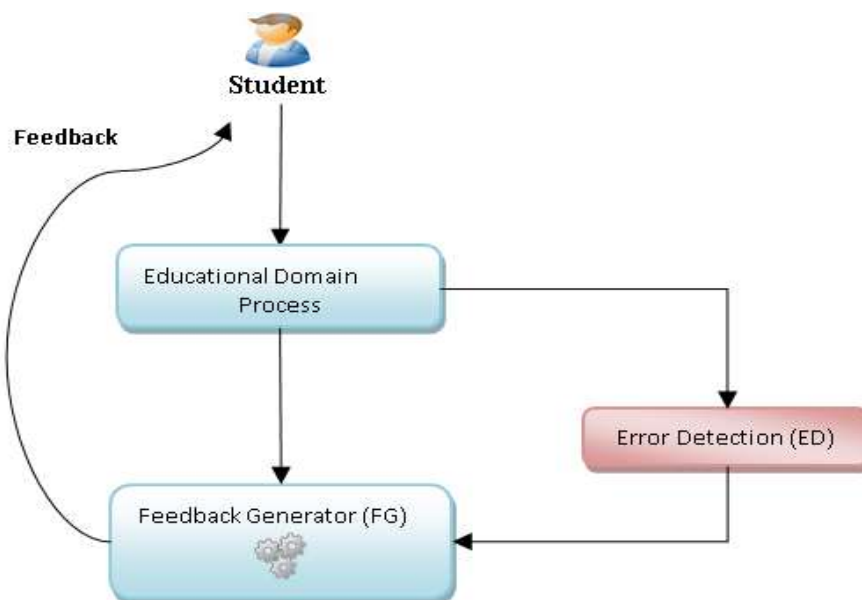
R3: If submitted_answers is low and answer_score is low and hints_delivered are
many

then do_not_take_hint_recommendation and submit_answer_recommendation.

Η ανάπτυξη των κανόνων βασίστηκε σε ιστορικά δεδομένα των εκπαιδευόμενων σε συνδυασμό με την εμπειρία και την καθοδήγηση ειδικών-εμπειρογνομόνων. Τα χαρακτηριστικά που σχετίζονται με την συμπεριφορά των εκπαιδευόμενων ενημερώνονται και επικαιροποιούνται συνεχώς κατά τη διάρκεια των εκπαιδευτικών δραστηριοτήτων με βάση την αλληλεπίδρασή τους με το σύστημα.

2.1.4 Πλαίσιο Παραγωγής Ανάδρασης

Το πλαίσιο για την παραγωγή ανάδρασης (feedback generation framework) αποτελείται από τρεις βασικές μονάδες: τον *μηχανισμό αναγνώρισης λαθών* (*Error Detection Mechanism*), τον *μηχανισμό παραγωγής ανάδρασης* (*Feedback Generation Mechanism*) και το *πεδίο της εκπαιδευτικής διαδικασίας* (*Educational Domain Process*) όπως παρουσιάζεται στην Εικόνα 2.1. Ο *μηχανισμός αναγνώρισης λαθών* είναι η μονάδα που είναι υπεύθυνη να αναγνωρίζει και να κατηγοριοποιεί τα λάθη των εκπαιδευόμενων. Ο *μηχανισμός παραγωγής ανάδρασης* παράγει κατάλληλα μηνύματα ανάδρασης ανάλογα με τα λάθη των εκπαιδευόμενων. Το *πεδίο της εκπαιδευτικής διαδικασίας* αφορά το ευφύες εκπαιδευτικό σύστημα στο οποίο μπορεί να ενσωματωθεί το πλαίσιο παραγωγής ανάδρασης.



Εικόνα 2.1 Αρχιτεκτονική Παραγωγής Ανάδρασης

2.1.4.1 Μηχανισμός Αναγνώρισης Λαθών

Η διαδικασία αναγνώρισης λαθών είναι μια θεμελιώδης διεργασία στα ευφυή συστήματα διδασκαλίας και συμβάλει ώστε να γίνει η κατανόηση και η μάθηση πιο αποτελεσματική (Hirashima, 1998). Η αναγνώριση των λαθών και ο σαφής προσδιορισμός των τύπων των λαθών είναι ιδιαίτερα σημαντικός παράγοντας για την παραγωγή της κατάλληλης ανάδρασης. Ο μηχανισμός αναγνώρισης λαθών έχει ως στόχο να αναλύει τις απαντήσεις των εκπαιδευόμενων και να αναγνωρίζει τα είδη των λαθών που έγιναν. Ο μηχανισμός της αναγνώρισης λαθών επιτελεί δύο βασικές λειτουργίες. Αρχικά, χαρακτηρίζει την απάντηση του εκπαιδευόμενου όσον αφορά την πληρότητα (completeness) και την ακρίβεια (accuracy) και την κατηγοριοποιεί στον κατάλληλο τύπο απάντησης. Στη συνέχεια, σε περίπτωση που η απάντηση που υπέβαλε ο εκπαιδευόμενος δεν είναι σωστή, γίνεται εις βάθος ανάλυση των λαθών προσδιορίζοντας τον αριθμό και το είδος των λαθών που έγιναν. Για τον προσδιορισμό του τύπου των λαθών χρησιμοποιείται ένα σχήμα κατηγοριοποίησης λαθών του πεδίου της εφαρμογής που προσδιορίζει τα είδη των λαθών και τους περιορισμούς που ισχύουν.

2.1.4.2 Μηχανισμός Παραγωγής Ανάδρασης

Ο μηχανισμός παραγωγής ανάδρασης (feedback generation mechanism) χρησιμοποιείται για να παράγει και να παρέχει στους εκπαιδευόμενους τους κατάλληλους τύπους μηνυμάτων ανάδρασης, με βάση το μοντέλο της ιεραρχίας ανάδρασης.

Ο γενικός μηχανισμός παραγωγής ανάδρασης που ακολουθείται έχει ως εξής:

1. Όταν υποβληθεί μια απάντηση του εκπαιδευόμενου, τότε προωθείται στη μονάδα αναγνώριση λαθών για να γίνει η ανάλυση της και η κατηγοριοποίηση στην αντίστοιχη κατηγορία απαντήσεων.
 - 1.1 Αν η απάντηση είναι σωστή (complete- accurate), προχωράει στο επόμενο στάδιο της εκπαιδευτικής διαδικασίας.
 - 1.2 Διαφορετικά, η απάντηση του εκπαιδευόμενου αναλύεται εις βάθος από το μηχανισμό αναγνώρισης λαθών για να αναγνωριστούν τα είδη των λαθών που έγιναν με βάση την κατηγοριοποίηση των λαθών του πεδίου.
2. Με βάση το ιεραρχικό μοντέλο ανάδρασης και την απάντηση του εκπαιδευόμενου καθορίζεται ο τύπος της ανάδρασης.
3. Αναγνωρίζονται οι τύποι των λαθών, προσδιορίζεται ο κατάλληλος τύπος ανάδρασης και δημιουργείται το κατάλληλο μήνυμα ανάδρασης ως εξής:

- Για κάθε τύπο λάθους, γίνεται ανάλυση της άσκησης για να καθοριστούν τα αίτια του λάθους και δημιουργούνται οι κατάλληλες επεξηγήσεις.
 - Για κάθε στοιχείο που είναι σωστό, δημιουργείται το κατάλληλο θετικό μήνυμα ανάδρασης (positive feedback message)
4. Παράγονται τα κατάλληλα μηνύματα κειμένου ακολουθώντας προσέγγιση που βασίζεται σε πρότυπα μηνυμάτων (template-based).
 5. Παρέχονται μηνύματα ανάδρασης στον εκπαιδευόμενο, σύμφωνα με το ιεραρχικό μοντέλο ανάδρασης και με τον κατάλληλο τρόπο, δηλαδή άμεσα (immediately) ή κατά απαίτηση του εκπαιδευόμενου (on demand).
 6. Επιτρέπει στον εκπαιδευόμενο να παρέχει μια νέα απάντηση και πηγαίνει στο βήμα 1.

2.2 Εφαρμογή του Πλαισίου Ανάδρασης στο Ευφύες Εκπαιδευτικό Σύστημα NLtoFOL

Το πλαίσιο της παροχής ανάδρασης εφαρμόστηκε στο ευφύες εκπαιδευτικό σύστημα NLtoFOL (Hatzilygeroudis & Perikos; 2009), το οποίο χρησιμοποιείται για την διδασκαλία της μετατροπής φυσικής γλώσσας (ΦΓ) σε κατηγορηματική λογική πρώτης τάξης (ΚΛΠΤ). Το σύστημα έχει ως στόχο την παροχή ενός δομημένου τρόπου μετατροπής εκφράσεων φυσικής γλώσσας σε κατηγορηματική λογική και την αλληλεπιδραστική καθοδήγηση του χρήστη κατά την διάρκεια των σταδίων της μετατροπής. Η παροχή ανάδρασης προσαρμόστηκε κατάλληλα στο πεδίο της εφαρμογής και δημιουργήθηκε το κατάλληλο σχήμα κατηγοριοποίησης των λαθών του πεδίου. Επίσης, δημιουργήθηκαν τα κατάλληλα πρότυπα μηνύματα ανάδρασης προσαρμοσμένα στα χαρακτηριστικά της διαδικασίας. Στην συνέχεια παρουσιάζονται τα βήματα της διαδικασίας μετατροπής ΦΓ σε ΚΛΠΤ καθώς και η προσαρμογή και λειτουργία του πλαισίου ανάδρασης στο σύστημα.

2.2.1 Διαδικασία Μετατροπής Φυσικής Γλώσσας σε ΚΛΠΤ

Η ΚΛΠΤ ως γλώσσα αναπαράστασης γνώσης και αυτομάτου συλλογισμού έχει πολλές πτυχές. Μια από αυτές είναι η μετατροπή ΦΓ σε εκφράσεις ΚΛΠΤ που φέρουν την σημασιολογική πληροφορία των προτάσεων. Πρόκειται για μια ad-hoc διαδικασία που δεν μπορεί να αυτοματοποιηθεί μέσω υπολογιστή. Στα πλαίσια του εκπαιδευτικού συστήματος NLtoFOL σχεδιάστηκε και υλοποιήθηκε μια μοντελοποίηση της διαδικασίας σε βήματα (Hatzilygeroudis, 2007; Perikos & Hatzilygeroudis, 2009). Πιο συγκεκριμένα, η διαδικασία μετατροπής μοντελοποιήθηκε σε 10 βήματα τα οποία είναι τα ακόλουθα:

1. Προσδιορισμός των κατηγορημάτων και των συναρτήσεων της πρότασης.
2. Προσδιορισμός του αριθμού, του τύπου και των συμβόλων των συναρτησιακών εκφράσεων και των κατηγορημάτων.
3. Προσδιορισμός των ποσοδεικτών των μεταβλητών.
4. Σχηματισμός των δομικών ατομικών εκφράσεων.
5. Σχηματισμός ομάδων ατόμων ίδιου επιπέδου.
6. Προσδιορισμός των συνδετικών των ομάδων ατόμων και σχηματισμός τύπων.
7. Προσδιορισμός ομάδων τύπων ίδιου επιπέδου.
8. Αν έχει προσδιοριστεί μόνο μία ομάδα τύπων, σχηματισμός τύπου επόμενου επιπέδου και μετάβαση στο βήμα 10.
9. Προσδιορισμός συνδετικών τύπων ομάδων, σχηματισμός τύπων επόμενου επιπέδου και μετάβαση στο βήμα 7.
10. Προσδιορισμός θέσης ποσοδεικτών και σχηματισμός τελικής έκφρασης.

Το πλήθος των βημάτων καθώς και ποια βήματα της διαδικασίας χρειάζεται να εφαρμοστούν για να μετατραπεί μια πρόταση ΦΓ σε ΚΛΠΤ εξαρτάται από τα χαρακτηριστικά και τη σημασιολογία της πρότασης. Επομένως, διαφορετικές προτάσεις μπορεί να απαιτούν την εφαρμογή διαφορετικών βημάτων της διαδικασίας και μπορεί με βάση τα χαρακτηριστικά της πρότασης να χρειαστεί να εφαρμοστούν μερικά βήματα της διαδικασίας περισσότερες από μία φορές. Η διαδικασία παρουσιάζεται μέσω ενός παραδείγματος φορμαλισμού της ακόλουθης πρότασης:

“Every city has a dog-catcher that has been bitten by every dog living in the city”

Βήμα 1: Εντοπίζονται τα ρήματα, τα ουσιαστικά και τα επίθετα μέσα στην πρόταση και προσδιορίζονται αντίστοιχα τα κατηγορήματα ή τα συναρτησιακά σύμβολα.

Στην συγκεκριμένη πρόταση υπάρχουν πέντε κατηγορήματα τα οποία είναι τα εξής:

city → κατηγορήμα: city

has a dog-catcher → κατηγορήμα: dog-catcher

has been bitten → κατηγορήμα: has-bitten

dog → κατηγορήμα: dog

living in →κατηγορία: lives-in

Βήμα 2: Προσδιορισμός του αριθμού, του τύπου και των συμβόλων των ορισμάτων των συναρτησιακών εκφράσεων και των κατηγορημάτων.

Η εφαρμογή του βήματος παρουσιάζεται στον ακόλουθο Πίνακα:

Πίνακας 2.4 Ανάλυση του βήματος 2 της διαδικασίας NLtoFOL

Κατηγορία/Συνάρτηση	Πλήθος ορισμάτων	Είδος ορισμάτων	Σύμβολο
city	1	variable	x
dog-catcher	2	variable, variable	y, x
dog	1	variable	z
lives-in	2	variable, variable	z, x
has-bitten	2	variable, variable	y, z

Βήμα 3: Προσδιορισμός των ποσοδεικτών των μεταβλητών.

$x \rightarrow \forall$, η πρόταση περιέχει το “every”

$y \rightarrow \exists$, η πρόταση περιέχει το “has a”

$z \rightarrow \forall$, η πρόταση περιέχει το “every”

Βήμα 4: Δημιουργούνται τόσα άτομα όσα και τα κατηγορήματα και γίνεται σχηματισμός ατομικών εκφράσεων:

Atom 1: city(x), Atom 2: dog-catcher(y, x),

Atom 3: dog(z), Atom 4: lives-in(z, x), Atom 5: has-bitten(z, y)

Βήμα 5: Σχηματισμός ομάδων ατόμων ίδιου επιπέδου.

Αυτό αναφέρεται κυρίως σε ομαδοποίηση ατόμων που θα πρέπει να συνδεθούν μεταξύ τους με κάποια συνδετικά, επειδή ανήκουν στο ίδιο σύνθετο συστατικό μιας πρότασης:

AtomGroup 1: {city(x)}, AtomGroup 2: {dog-catcher(y, x)},

AtomGroup 3: {dog(z), lives-in(z, x)} AtomGroup 4: {has-bitten(z, y)}

Βήμα 6: Προσδιορισμός των συνδετικών μεταξύ των ατόμων κάθε ομάδας και σχηματισμός των αντίστοιχων λογικών εκφράσεων.

AtomGroup 1 $\rightarrow$ Form 1: city(x), AtomGroup 2 $\rightarrow$ Form 2: dog-catcher(y, x),

AtomGroup 3 $\rightarrow$ Form 3: dog(z) $\wedge$ lives-in(z, x), AtomGroup 4 $\rightarrow$ Form 4: has-bitten(z, y)

Βήμα 7: Προσδιορισμός ομάδων τύπων ίδιου επιπέδου.

Αυτό συνήθως αφορά στον προσδιορισμό των αριστερών και δεξιών τμημάτων μιας συνεπαγωγής:

FormGroup1-1: {city(x)}, FormGroup1-2: {dog-catcher(y, x)},

FormGroup1-3: {dog(z) $\wedge$ lives-in(z, x), has-bitten(z, y)}

Βήμα 9: Προσδιορισμός των συνδετικών μεταξύ των τύπων της κάθε ομάδας, σχηματισμός εκφράσεων επόμενου επιπέδου και μετάβαση στο βήμα 7

FormGroup1-1 $\rightarrow$ Form1-1: city(x), FormGroup1-2 $\rightarrow$ Form1-2: dog-catcher(y, x),

FormGroup1-3 $\rightarrow$ Form1-3: (dog(z) $\wedge$ lives-in(z, x)) $\Rightarrow$ has-bitten(z, y)

Βήμα 7: Προσδιορισμός ομάδων τύπων ίδιου επιπέδου.

FormGroup 2-1: {city(x)}

FormGroup 2-2: {dog-catcher(y, x), (dog(z) $\wedge$ lives-in(z, x)) $\Rightarrow$ has-bitten(z, y)}

Βήμα 9: Προσδιορισμός των συνδετικών μεταξύ των τύπων της κάθε ομάδας και σχηματισμός του επόμενου επιπέδου εκφράσεις και μετάβαση στο βήμα 7

FormGroup 2-1 $\rightarrow$ Form2-1: city(x)

FormGroup 2-2 $\rightarrow$ Form2-2: dog-catcher(y, x) $\wedge$ (dog(z) $\wedge$ lives-in(z, x)) $\Rightarrow$ has-bitten(z, y)}

Βήμα 7: Προσδιορισμός ομάδων τύπων ίδιου επιπέδου.

FormGroup 3-1: {city(x), dog-catcher(y, x) $\wedge$ (dog(z) $\wedge$ lives-in(z, x)) $\Rightarrow$ has-bitten(z, y)}

Βήμα 8: Αν έχει προσδιοριστεί μόνο μία ομάδα τύπων, σχηματισμός τύπου επόμενου επιπέδου και μετάβαση στο βήμα 10.

FormGroup 3-1 $\rightarrow$ Form 2-3: $\text{city}(x) \Rightarrow (\text{dog-catcher}(y, x) \wedge (\text{dog}(z) \wedge \text{lives-in}(z, x)) \Rightarrow \text{has-bitten}(z, y))$

Βήμα 10: Προσδιορισμός της θέσης των ποσοδεικτών και δημιουργία της τελικής έκφρασης

Η παραγόμενη τελική έκφραση ΚΛΠΤ είναι:

$(\forall x) \text{city}(x) \Rightarrow ((\exists y) \text{dog-catcher}(y, x) \wedge (\forall z) (\text{dog}(z) \wedge \text{lives-in}(z, x)) \Rightarrow \text{has-bitten}(z, y))$

2.2.2 Κατηγοριοποίηση Λαθών στο Σύστημα NLtoFOL

Για την εφαρμογή του πλαισίου ανάδρασης όπως αναφέρθηκε παραπάνω είναι αναγκαίο να γίνει αναγνώριση των λαθών και δημιουργία του σχήματος κατηγοριοποίησης των λαθών με βάση το πεδίο της εφαρμογής. Για το σκοπό αυτό, έγινε στα πλαίσια του συστήματος καταγραφή όλων των λαθών και προσδιορίστηκαν διάφοροι τύποι λαθών που μπορούν να γίνουν στα πλαίσια της διαδικασίας μετατροπής ΦΓ σε ΚΛΠΤ. Τα λάθη που μπορούν να γίνουν κατά την μετατροπή ΦΓ σε ΚΛΠΤ μέσω του συστήματος NLtoFOL κατηγοριοποιούνται σε έξι κύριες κατηγορίες όπως παρουσιάζονται στον Πίνακα 2.5.

Πίνακας 2.5 Κατηγοριοποίηση λαθών της διαδικασίας φορμαλισμού

Κατηγορία Λάθους	Τύπος Λάθους	Επεξηγήσεις Λαθών
Predicate Error	Number Error	Ουσιαστικό σε πληθυντικό αριθμό αντί σε ενικό. Π.χ. “men” αντί “man”.
	Person Error	Κατηγορία σε διαφορετικό πρόσωπο αντί για τρίτο πρόσωπο. Π.χ. “eat” αντί “eats”.
	Semantics Error	Μια λέξη π.χ. χρησιμοποιείται λανθασμένα ως κατηγορία αντί για συνάρτηση ή σταθερά
Term Error	Type Error	Λανθασμένος τύπος έχει προσδιοριστεί για ένα όρισμα ενός κατηγορήματος ή συνάρτησης. Π.χ. προσδιορίστηκε “σταθερά” αντί “μεταβλητή”
	Semantics Error	Μια λέξη λανθασμένα χρησιμοποιείται ως συναρτησιακό σύμβολο (π.χ. “brother-of”) ή σταθερά αντί κατηγορήματος.
Argument Error	Cardinality Error	Ο αριθμός των ορισμάτων ενός ατόμου δεν είναι σωστός.
	Order Error	Η σειρά των ορισμάτων ενός ατόμου δεν είναι σωστή.

Connective Error	Connective Error	Ένα διασυνδετικό (π.χ. “^”) χρησιμοποιείται αντί κάποιου άλλου (π.χ. “⇒”) για να συνδέσει τα άτομα μιας έκφρασης ή μια άρνηση λείπει.
	Scope Error	Η εμβέλεια ενός διασυνδετικού δεν είναι σωστή.
Quantifier Error	Type Error	Ο τύπος ενός ποσοδείκτη δεν είναι σωστός.
	Scope Error	Η εμβέλεια ενός ποσοδείκτη δεν καλύπτει όλες τις εμφανίσεις της μεταβλητής του.
	Order Error	Η θέση των ποσοδεικτών δεν είναι σωστή.
Group Error	Atoms Error	Τα άτομα που ορίζονται για να σχηματιστεί μια ομάδα είναι λανθασμένα.
	Order Error	Η σειρά των ατόμων για να σχηματίσουν μια ομάδα είναι λανθασμένη..
	Cardinality Error	Ο αριθμός των ατόμων που σχηματίζουν μια ομάδα είναι λανθασμένος..
	Number Error	Ο αριθμός των ομάδων που δημιουργήθηκαν είναι λανθασμένος.

Στην συνέχεια, παρουσιάζονται μερικά ενδεικτικά παραδείγματα απαντήσεων εκπαιδευόμενων και ανταπόκρισης του συστήματος κατά την διάρκεια της μετατροπής της ακόλουθης πρότασης ΦΓ:

“Every city has a dog-catcher that has been bitten by every dog living in the city”

σε ΚΛΠΤ μέσω του ευφυούς συστήματος NLtoFOL.

Στο πρώτο στάδιο της διαδικασίας, οι εκπαιδευόμενοι καλούνται να προσδιορίσουν τα κατηγορήματα και τις συναρτήσεις. Ας υποθέσουμε ότι ένας εκπαιδευόμενος προσδιόρισε τα εξής τρία κατηγορήματα: “city”, “dog” και “lives-in”. Ο προσδιορισμός των κατηγορημάτων είναι σωστός, αλλά ο εκπαιδευόμενος δεν προσδιόρισε όλα τα κατηγορήματα της πρότασης. Επομένως, το σύστημα χαρακτηρίζει την απάντηση του εκπαιδευόμενου ως *Incomplete-Accurate*. Ας υποθέσουμε μια άλλη απάντηση εκπαιδευόμενου, όπου προσδιορίστηκαν τα εξής πέντε κατηγορήματα : “city”, “dog-catcher”, “dogs”, “lives-in”, “has-bitten”. Η απάντηση του εκπαιδευόμενου χαρακτηρίζεται ως *complete* διότι προσδιορίστηκαν όλα τα στοιχεία της σωστής απάντησης, αλλά χαρακτηρίζεται ως *inaccurate* δεδομένου ότι το κατηγορήμα “dogs” έχει προσδιοριστεί με

λάθος τρόπο. Επομένως, η απάντηση του εκπαιδευόμενου χαρακτηρίζεται ως *Complete-Inaccurate*.

Στο δεύτερο στάδιο της διαδικασίας, οι εκπαιδευόμενοι θα πρέπει να προσδιορίσουν τον αριθμό, τον τύπο και το σύμβολο για κάθε όρισμα της συνάρτησης και των κατηγορημάτων. Ας υποθέσουμε ότι ένας εκπαιδευόμενος έχει δώσει την απάντηση που παρουσιάζεται στον Πίνακα 2.6. Το σύστημα χαρακτηρίζει την απάντηση του εκπαιδευόμενου ως *Complete-Inaccurate*. Η απάντηση χαρακτηρίζεται ως *complete*, διότι ο εκπαιδευόμενος έχει προσδιορίσει όλα τα στοιχεία της σωστής απάντησης, και *Inaccurate* διότι δυο στοιχεία της απάντησης δεν έχουν προσδιοριστεί σωστά. Στη συνέχεια, η απάντηση του εκπαιδευόμενου αναλύεται από τον μηχανισμό αναγνώρισης λαθών (ED) και αναγνωρίζονται τα λάθη που έχουν γίνει. Ο μηχανισμός αναγνωρίζει ότι ο εκπαιδευόμενος έχει ορίσει το κατηγορημα “dog” λανθασμένα. Πιο συγκεκριμένα, ο εκπαιδευόμενος έχει κάνει ένα *argument error* και ειδικότερα ένα *cardinality error* έχοντας ορίσει δυο ορίσματα αντί για ένα όρισμα. Επίσης, ο εκπαιδευόμενος έχει κάνει ένα *order error* για το κατηγορημα “lives-in” έχοντας προσδιορίσει λανθασμένη θέση για τις μεταβλητές των ορισμάτων.

Πίνακας 2.6 Απάντηση Εκπαιδευόμενου και σωστή απάντηση για το δεύτερο βήμα της διαδικασίας.

Απάντηση Εκπαιδευόμενου (SA)				Response
<i>Predicate/ Function</i>	<i>Arity</i>	<i>Argument Types</i>	<i>Symbols</i>	
city	1	variable	x	Correct ✓
dog-catcher	2	variable, variable	y, x	Correct ✓
dog	2	variable, variable	x, z	Incorrect ✗
				Incorrect ✗
lives-in	2	variable, variable	x, z	Correct ✓
has-bitten	2	variable, variable	y, z	

Σωστή Απάντηση (CA)			
<i>Predicate/ Function</i>	<i>Arity</i>	<i>Argument Types</i>	<i>Symbols</i>
city	1	variable	x
dog-catcher	2	variable, variable	y, x
dog	1	variable	z
lives-in	2	variable, variable	z, x
has-bitten	2	variable, variable	y, z

Στην συνέχεια παρουσιάζεται ένα παράδειγμα απάντησης ενός εκπαιδευόμενου για το τρίτο βήμα, όπου γίνεται ο προσδιορισμός των ποσοδεκτών των μεταβλητών:

“ $x \rightarrow \exists$ ”, “ $y \rightarrow \exists$ ”, “ $z \rightarrow \forall$ ”, “ $w \rightarrow \exists$ ”.

Ο εκπαιδευόμενος έχει προσδιορίσει τέσσερις ποσοδείκτες και το σύστημα χαρακτηρίζει την απάντηση του εκπαιδευόμενου ως *Superfluous* και πιο συγκεκριμένα ως *Superfluous-Inaccurate*. Η απάντηση χαρακτηρίζεται ως *Superfluous* διότι ο εκπαιδευόμενος έχει προσδιορίσει έναν επιπλέον ποσοδείκτη (τέσσερις ποσοδείκτες αντί για τρεις που απαιτούνται) και *Inaccurate* επειδή ο ποσοδείκτης της μεταβλητής x δεν έχει προσδιοριστεί σωστά. Ο εκπαιδευόμενος έχει προσδιορίσει “ $x \rightarrow \exists$ ” αντί “ $x \rightarrow \forall$ ”, που αποτελεί τη σωστή απάντηση. Ο μηχανισμός αναγνώρισης λαθών αναγνωρίζει ότι ο εκπαιδευόμενος έχει κάνει ένα *quantifier error* και πιο συγκεκριμένα ένα *type error*.

Στο έκτο βήμα της διαδικασίας, ο εκπαιδευόμενος θα πρέπει να καθορίσει τα διασυνδεδετικά μεταξύ των ατόμων της κάθε ομάδας και θα πρέπει να δημιουργήσει τις αντίστοιχες λογικές εκφράσεις. Μια απάντηση του εκπαιδευόμενου ήταν τα εξής:

AtomGroup 1 $\rightarrow$ Form 1: $city(x)$, AtomGroup 2 $\rightarrow$ Form 2: $dog-catcher(y, x)$

AtomGroup 3 $\rightarrow$ Form 3: $dog(z) \vee lives-in(z, x)$

Το σύστημα χαρακτηρίζει την απάντηση του εκπαιδευόμενου ως *Incomplete-Inaccurate*. Ο εκπαιδευόμενος δεν έχει προσδιορίσει σωστά όλες τις ομάδες ατόμων της πρότασης και η απάντησή χαρακτηρίζεται ως *incomplete*. Επίσης, η απάντηση χαρακτηρίζεται ως *inaccurate*, επειδή ένα άτομο της ομάδας έχει προσδιοριστεί λανθασμένα. Στην συνέχεια, ο μηχανισμός αναγνώρισης λαθών αναγνωρίζει ένα *connective error*, επειδή το διασυνδετικό “ $\vee$ ” που καθορίστηκε από τον εκπαιδευόμενο για το AtomGroup3 δεν είναι σωστό. Επιπλέον, η απάντηση χαρακτηρίζεται ως *incomplete*, επειδή ο εκπαιδευόμενος δεν έχει προσδιορίσει το απαραίτητο AtomGroup4: “*has-bitten(z, y)*”.

Τέλος, ας θεωρήσουμε το ακόλουθο παράδειγμα απάντησης για το βήμα10:

“($\forall x$) *City(x)* $\Rightarrow$ (($\exists y$) *dog-catcher(y, x)* $\wedge$ ($\forall z$) (*dog(z)* $\wedge$ *lives-in(z, x)*) $\Rightarrow$ *has-bitten(z, y)*)”

Το σύστημα κατηγοριοποιεί την απάντηση ως *Complete-Accurate*. Η απάντηση χαρακτηρίζεται ως *complete*, επειδή ο εκπαιδευόμενος έχει καθορίσει όλα τα στοιχεία της έκφρασης ΚΛΠΤ και *accurate* επειδή, όλα τα στοιχεία που προσδιορίστηκαν είναι σωστά. Επομένως, η απάντηση είναι σωστή.

2.2.3 Παραγωγή Κειμενικών Μηνυμάτων Ανάδρασης

Ο μηχανισμός παραγωγής ανάδρασης (Feedback Generation FG) χρησιμοποιείται για να παράγει και να παρέχει στους εκπαιδευόμενους διάφορους τύπους ανάδρασης ανάλογα με τις ενέργειές τους. Η παραγωγή των κειμένων των μηνυμάτων ανάδρασης γίνεται χρησιμοποιώντας μια προσέγγιση βασισμένη στα πρότυπα μηνυμάτων (template-based approach). Πιο συγκεκριμένα, οι τύποι των λαθών και οι υποδείξεις του συστήματος σχετίζονται με ένα αντίστοιχο πρότυπο (template), το οποίο αποτελείται από τρία κύρια μέρη: (i) την κατηγορία του λάθους, (ii) τον τύπο του λάθους και (iii) το πρότυπο (pattern) με βάση το οποίο θα δημιουργηθούν τα κείμενα των μηνυμάτων. Στον Πίνακα 2.7, παρουσιάζονται μερικά παραδείγματα προτύπων ανάδρασης (feedback texture templates).

Πίνακας 2.7 Παραδείγματα προτύπων μηνυμάτων

Κατηγορία	Τύπος	Pattern
Λάθους	Λάθους	
Predicate Error	Number Error	[element of the student answer] is not correct because [explanation]. You should [remedial action].

Predicate Error	Person Error	[element of the student answer] is not correct. It is not in third person.
Quantifier Error	Type Error	Quantifier [element of the student answer] is not correct because [explanation] is associated with the word [explanation].

Όταν ο εκπαιδευόμενος υποβάλλει μια απάντηση που δεν είναι σωστή, μετά την ανάλυση της απάντησης και τον προσδιορισμό των λαθών που έγιναν, ο μηχανισμός ανάδρασης ανασύρει το κατάλληλο πρότυπο προκειμένου να δημιουργηθούν τα μηνύματα κειμένου που θα δοθούν στον εκπαιδευόμενο. Η συμπλήρωση του προτύπου γίνεται με βάση τα στοιχεία της απάντησης του εκπαιδευόμενου. Για παράδειγμα, ας θεωρήσουμε την μετατροπή της πρότασης “Maria loves all dogs that love their master”. Στο πρώτο βήμα της διαδικασίας, ένα συνηθισμένο λάθος αφορά τον προσδιορισμό των κατάλληλων κατηγορημάτων. Ας υποθέσουμε ότι ο εκπαιδευόμενος έχει προσδιορίσει ως κατηγορήματα τις λέξεις “love”, “dogs” και τη λέξη “master-of” ως συνάρτηση. Ο μηχανισμός αναγνώρισης λαθών αναγνωρίζει τα λάθη και ενημερώνει τη μονάδα παραγωγής ανάδρασης σχετικά με τα είδη των λαθών. Η πρόταση αναλύεται και η λέξη “love” αναγνωρίζεται ως ρήμα και ότι δεν είναι στο τρίτο πρόσωπο στην απάντηση του εκπαιδευόμενου, κάτι που καθιστά την απάντηση λανθασμένη. Επίσης, αναγνωρίζεται ότι η λέξη “dogs” είναι ένα ουσιαστικό στον πληθυντικό αριθμό, κάτι που επίσης δεν είναι σωστό. Ο μηχανισμός παραγωγής ανάδρασης ενεργοποιεί τα σχετικά πρότυπα:

[element of the student answer] is not correct because [explanation]

[element of the student answer] is not correct. It is not in third person

Στη συνέχεια αυτά συμπληρώνονται με τα αντίστοιχα στοιχεία και δημιουργούνται τα ακόλουθα μηνύματα ανάδρασης που θα δοθούν στους εκπαιδευόμενους:

The word dogs is not correct because it is in plural number.

The word love is not correct. It isn't in third person.

Ως ένα δεύτερο παράδειγμα, ας θεωρήσουμε την περίπτωση όπου κατά την διάρκεια του τρίτου σταδίου ο εκπαιδευόμενος έχει καθορίσει λάθος τους ποσοδείκτες για τις μεταβλητές που αναφέρονται σε κατηγορήματα της άσκησης. Για παράδειγμα, ας υποθέσουμε ότι εκπαιδευόμενος στην απάντησή του έχει καθορίσει για τη μεταβλητή “x” που αντιπροσωπεύει την οντότητα “σκύλος” τον υπαρξιακό ποσοδείκτη “ $\exists$ ”, κάτι που δεν

αποτελεί μια σωστή απάντηση. Ο μηχανισμός ανάδρασης με βάση την ανάλυση της απάντησης θα ενεργοποιήσει το ακόλουθο πρότυπο:

Quantifier [element of the student answer] is not correct because is [explanation] is associated with the word [explanation]

το οποίο συμπληρώνει με τα κατάλληλα στοιχεία της άσκησης και το μήνυμα που θα εμφανιστεί στον εκπαιδευόμενο είναι το εξής :

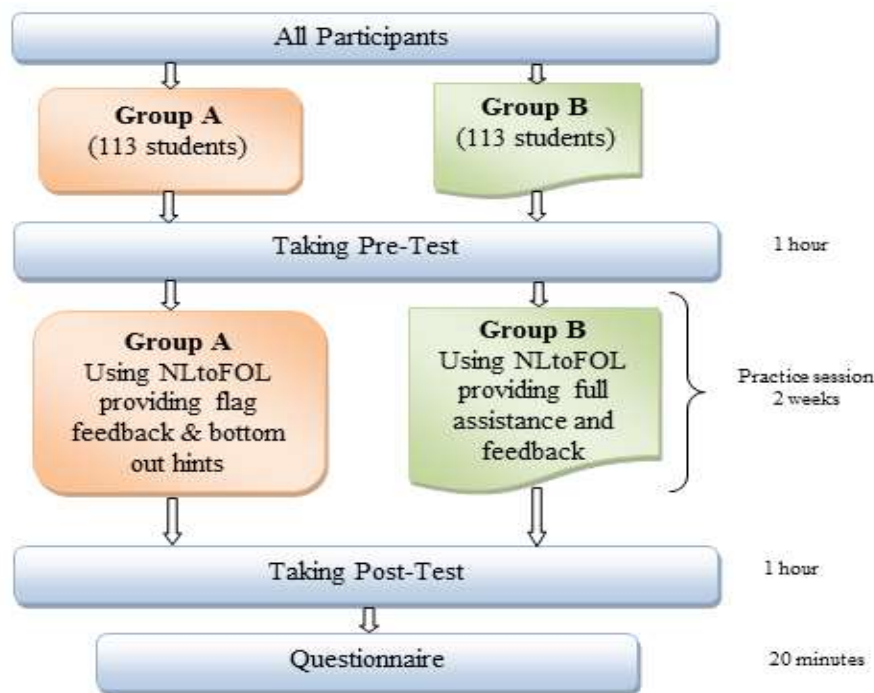
The quantifier specified for the variable “x” which represents “dogs” is not correct because the entity dogs is associated with the word “all”.

2.2.4 Πειραματική Αξιολόγηση

Για να αξιολογήσουμε τον μηχανισμό ανάδρασης έγινε εφαρμογή του μηχανισμού στο σύστημα NLtoFOL. Πιο συγκεκριμένα, σχεδιάστηκαν και διεξάχθηκαν δυο διαφορετικά πειράματα. Οι συμμετέχοντες στα δυο πειράματα ήταν φοιτητές του τμήματος Μηχανικών Η/Υ και Πληροφορικής του Πανεπιστημίου Πατρών και βρίσκονταν στο 4ο έτος σπουδών. Όλοι οι συμμετέχοντες παρακολουθούσαν το μάθημα της τεχνητής νοημοσύνης και είχαν τις απαραίτητες γνώσεις στο συγκεκριμένο πεδίο. Ο κύριος στόχος της πειραματικής μελέτης ήταν να αξιολογηθεί η αποτελεσματικότητα του μηχανισμού παραγωγής ανάδρασης που εφαρμόστηκε στο σύστημα NLtoFOL, να γίνει σύγκριση των διαφόρων σεναρίων ανάδρασης και να προσδιοριστεί η επίδρασή τους στη μάθηση των εκπαιδευόμενων. Για την πειραματική διαδικασία χρησιμοποιήθηκε η προσέγγιση pretest/posttest.

2.2.4.1 Πρώτη πειραματική μελέτη

Οι συμμετέχοντες στην πρώτη πειραματική μελέτη ήταν 226 προπτυχιακοί φοιτητές (άνδρες και γυναίκες) και η διαδικασία της πειραματικής μελέτης που ακολουθήθηκε παρουσιάζεται στην Εικόνα 2.2.



Εικόνα 2.2 Διαδικασία Πειράματος

Αρχικά, οι φοιτητές χωρίστηκαν τυχαία σε δύο ομάδες, GroupA και GroupB, όπου η κάθε ομάδα περιείχε 113 φοιτητές. Στην συνέχεια, όλοι οι φοιτητές συμμετείχαν σε ένα αρχικό τεστ (pretest), με στόχο να προσδιοριστεί το αρχικό επίπεδο γνώσεων πάνω στη μετατροπή της ΦΓ σε ΚΛΠΤ. Το pretest περιελάμβανε 10 ασκήσεις φυσικής γλώσσας, οι οποίες ήταν διαφόρων επιπέδων δυσκολίας (Perikos et al., 2011b). Πιο συγκεκριμένα, το pretest περιελάμβανε μια άσκηση επιπέδου “very easy”, δύο ασκήσεις “easy”, τέσσερις ασκήσεις “medium”, δύο ασκήσεις “difficult” και μία άσκηση επιπέδου “advanced”. Η βαθμολογία κυμαινόταν από 0 έως 10 και η διάρκεια του τεστ ήταν 1 ώρα. Στην συνέχεια, όλες οι απαντήσεις (εκφράσεις ΚΛΠΤ) των εκπαιδευόμενων του pretest βαθμολογήθηκαν αυτόματα από τον μηχανισμό αυτόματης αξιολόγησης (Perikos et al., 2012).

Μετά την διεξαγωγή του pretest, όλοι οι φοιτητές έκαναν εγγραφή και χρησιμοποίησαν το σύστημα NLtoFOL για την εκμάθηση της διαδικασίας μετατροπής ΦΓ σε ΚΛΠΤ. Κατά την εγγραφή, οι φοιτητές συμπλήρωσαν τα προσωπικά και δημογραφικά τους στοιχεία, όπως το όνομα, το φύλο, την ηλικία, το έτος σπουδών κ.α. Στην συνέχεια, οι φοιτητές αλληλεπίδρασαν με το σύστημα για δύο εβδομάδες, μετατρέποντας προτάσεις φυσικής γλώσσας σε ΚΛΠΤ, χρησιμοποιώντας τα παραπάνω βήματα της διαδικασίας. Πιο συγκεκριμένα, οι φοιτητές της ομάδας GroupA χρησιμοποίησαν μια έκδοση του

συστήματος που παρείχε ανάδραση μόνο των τύπων “*flag feedback*” και “*bottom out hints*”, ενώ οι φοιτητές της ομάδας GroupB αλληλεπιδράσαν με μια έκδοση του συστήματος που παρείχε πλήρη καθοδήγηση και ανάδραση. Επίσης, καθ’ όλη τη διάρκεια της αλληλεπίδρασης των εκπαιδευόμενων με το σύστημα και της διαδικασίας μάθησης, όλες οι ενέργειες καταγράφονταν από το σύστημα. Μετά από το πέρας της φάσης μάθησης των δύο εβδομάδων, όλοι οι εκπαιδευόμενοι συμμετείχαν σε ένα τελικό τεστ (posttest). Το posttest περιελάμβανε 10 ασκήσεις και τα επίπεδα δυσκολίας τους ήταν ίδια με το pretest. Πιο συγκεκριμένα, οι ασκήσεις ήταν ίδιου επιπέδου δυσκολίας και κάλυπταν τις ίδιες πτυχές της διαδικασίας με αυτές του pretest. Για κάθε άσκηση του posttest, όπως και στο pretest, οι εκπαιδευόμενοι είχαν την πρόταση σε φυσική γλώσσα και έπρεπε να την μετατρέψουν στην αντίστοιχη έκφραση ΚΛΠΤ χωρίς τη βοήθεια του συστήματος. Στην συνέχεια όλες οι απαντήσεις (εκφράσεις ΚΛΠΤ) των εκπαιδευόμενων στο posttest όπως και στο pretest βαθμολογήθηκαν αυτόματα από τον μηχανισμό αξιολόγησης (Perikos et al., 2012).

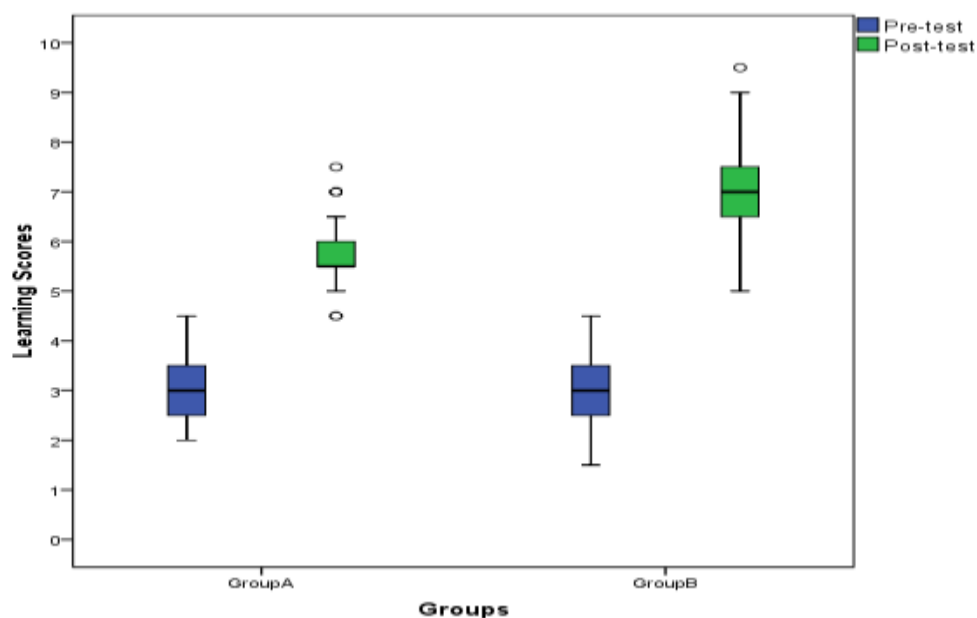
Για να αξιολογήσουμε τις επιδόσεις των εκπαιδευόμενων, πραγματοποιήθηκε πρώτα ανάλυση διακύμανσης ANOVA, για να ελέγξουμε αν υπήρχαν διαφορές στις δύο ομάδες (GroupA και GroupB). Πιο συγκεκριμένα, για την ομάδα GroupA η μέση τιμή στο pretest ήταν 3.12 (SD = 0.66) και για το Group B η μέση τιμή ήταν 2.30 (SD = 0.65). Από τα αποτελέσματα φαίνεται ότι δεν υπάρχει στατιστικά σημαντική διαφορά μεταξύ των ομάδων ($F = 2.988 > 1$, $p = 0.085 > 0.05$). Μάλιστα, η δοκιμή είχε στατιστική δύναμη 40.53% και το effect size (Cohen, 1988) ήταν μικρό ($d = 0.23$). Από τα αποτελέσματα προκύπτει ότι οι δύο ομάδες δεν διέφεραν ως προς το αρχικό επίπεδο γνώσεων τους και μάλιστα είχαν το ίδιο επίπεδο γνώσεων πριν από την έναρξη της διαδικασίας της μετατροπής προτάσεων φυσικής γλώσσας σε ΚΛΠΤ.

Επίσης, ένας επιπρόσθετος στόχος της παρούσας μελέτης ήταν να εξετάσει την αποτελεσματικότητα της ανάδρασης σχετικά με τη βελτίωση των γνώσεων των εκπαιδευόμενων. Για το σκοπό αυτό πραγματοποιήσαμε μια ανάλυση συνδιακύμανσης ANCOVA, για να εξετάσουμε τη διαφορά μεταξύ των επιδόσεων των δύο ομάδων, χρησιμοποιώντας τα αποτελέσματα του pretest ως συμμεταβλητή και τα αποτελέσματα του posttest ως εξαρτημένες μεταβλητές. Από τα αποτελέσματα προκύπτει ότι στο posttest υπάρχουν στατιστικά σημαντικές διαφορές μεταξύ των δύο ομάδων ($F = 14.98$, $p < .001$, partial eta = 0.06)

Πίνακας 2.8 Πειραματικά Αποτελέσματα ANCOVA

Groups	N	Mean	SD	Adjusted mean
Group A (Flag Feedback & Bottom out hints)	113	5.72	0.57	5.68
Group B (Full Feedback)	113	7.50	0.90	7.53

Στον Πίνακα 2.8 παρουσιάζονται τα αποτελέσματα της ανάλυσης ANCOVA. Οι προσαρμοσμένες μέσες τιμές (adjusted mean) των βαθμολογιών ήταν 5.71 για την ομάδα GroupA (flag feedback & bottom out hints) και 7.5 για την ομάδα GroupB (full feedback). Από τα αποτελέσματα προκύπτει μια σημαντική διαφορά μεταξύ των δύο ομάδων ($F = 14.98$ και $p = 0.00 < 0.05$), η οποία δείχνει ότι το σύστημα συνέβαλε αποτελεσματικά στη μάθηση για το GroupB. Πιο συγκεκριμένα, η χρήση του πλήρους σχήματος ανάδρασης (full feedback) είχε μεγαλύτερη και καλύτερη επίδραση στην μάθηση των εκπαιδευόμενων (GroupB), απ' ότι ένα μερικό σχήμα ανάδρασης (GroupA).



Εικόνα 2.3 Τα κέρδη μάθησης για τις δύο ομάδες

Στην εικόνα 2.3 παρουσιάζεται το θηκόγραμμα της βαθμολογίας του Group A και του Group B κατά το pretest και επίσης μετά την αλληλεπίδραση με το σύστημα και την διεξαγωγή του posttest. Από τα αποτελέσματα προκύπτει αύξηση στις επιδόσεις των εκπαιδευόμενων και στις δυο ομάδες και ότι το σύστημα συνέβαλε στη μάθηση, δείχνοντας ότι οι εκπαιδευόμενοι έχουν μάθει και στις δύο περιπτώσεις. Οι εκπαιδευόμενοι και των δύο ομάδων είχαν καλύτερη απόδοση στο posttest, επιτυγχάνοντας μεγαλύτερη βαθμολογία από ότι στο pretest. Ωστόσο, οι εκπαιδευόμενοι της ομάδας GroupB είχαν πολύ καλύτερες επιδόσεις από εκείνες της ομάδας GroupA. Πιο συγκεκριμένα, οι εκπαιδευόμενοι της ομάδας GroupB είχαν καλύτερες επιδόσεις στο posttest, έκαναν λιγότερα λάθη και κατανόησαν καλύτερα την διαδικασία της μετατροπής.

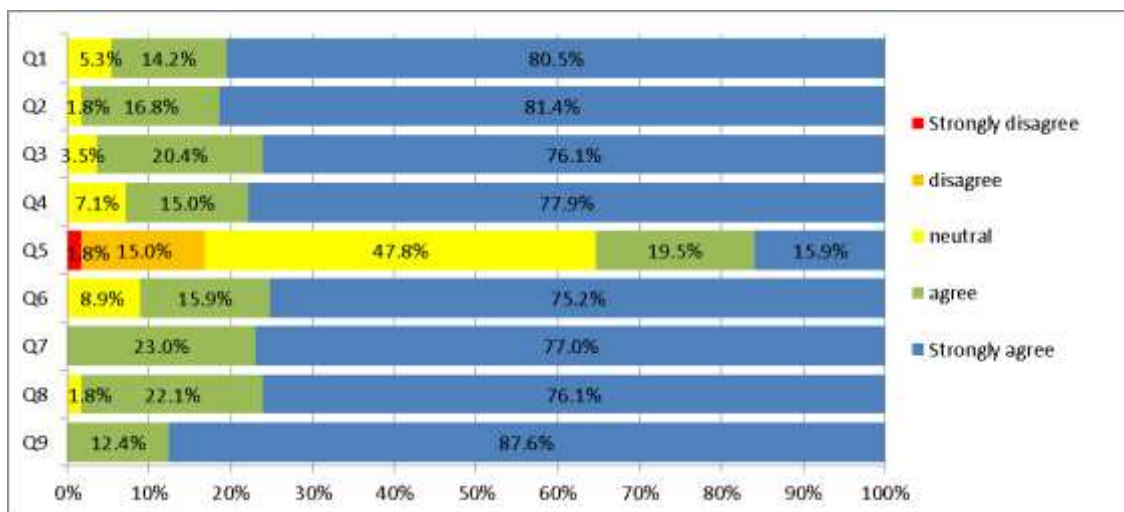
Στην τελευταία φάση της πειραματικής μελέτης, οι φοιτητές και των δύο ομάδων συμπληρώσαν ένα ερωτηματολόγιο. Το ερωτηματολόγιο περιελάμβανε ερωτήσεις για την αξιολόγηση του μηχανισμού ανάδρασης του συστήματος, καθώς και ερωτήσεις σχετικά με την εκμάθηση της διαδικασίας μέσω του συστήματος NLtoFOL. Το ερωτηματολόγιο του GroupA περιελάμβανε 12 ερωτήσεις, ενώ το ερωτηματολόγιο του GroupB περιελάμβανε 18 ερωτήσεις, δηλαδή έξι επιπλέον ερωτήσεις σχετικές με την γνώμη των φοιτητών για τα χαρακτηριστικά των κειμένων των μηνυμάτων ανάδρασης του συστήματος. Από τις 12 κοινές ερωτήσεις, οι 9 ερωτήσεις ήταν τύπου Likert-scale (1:διαφωνώ με 5:συμφωνώ απόλυτα) και οι 3 ερωτήσεις ήταν ανοιχτού τύπου. Οι ερωτήσεις παρουσιάζονται στον Πίνακα 2.9 και η ανάλυση των απαντήσεων των δύο ομάδων παρουσιάζεται στις εικόνες 2.4 και 2.5.

Πίνακας 2.9 Ερωτήσεις ερωτηματολογίου και για τις δυο ομάδες

Ερώτηση	
Q1	Βρήκατε την ενασχόληση με το NLtoFOL σύστημα ευχάριστη και ενδιαφέρουσα;
Q2	Το σύστημα σας βοήθησε στην μετατροπή των προτάσεων και στη διαδικασία της μάθησης;
Q3	Η ανάδραση και η καθοδήγηση του συστήματος ήταν ικανοποιητική;
Q4	Η ανάδραση που παράχθηκε από το σύστημα ήταν ακριβής και κατάλληλη;
Q5	Το σύστημα σας βοήθησε να κατανοήσετε τα λάθη σας;
Q6	Οι υποδείξεις που παρείχε το σύστημα σας βοήθησε να διορθώσετε τα λάθη σας;
Q7	Πως θα αξιολογούσατε την συνολική μαθησιακή σας εμπειρία με το σύστημα;
Q8	Νιώθετε με την χρησιμοποίηση του συστήματος πιο σίγουροι στην διαδικασία της μετατροπής;

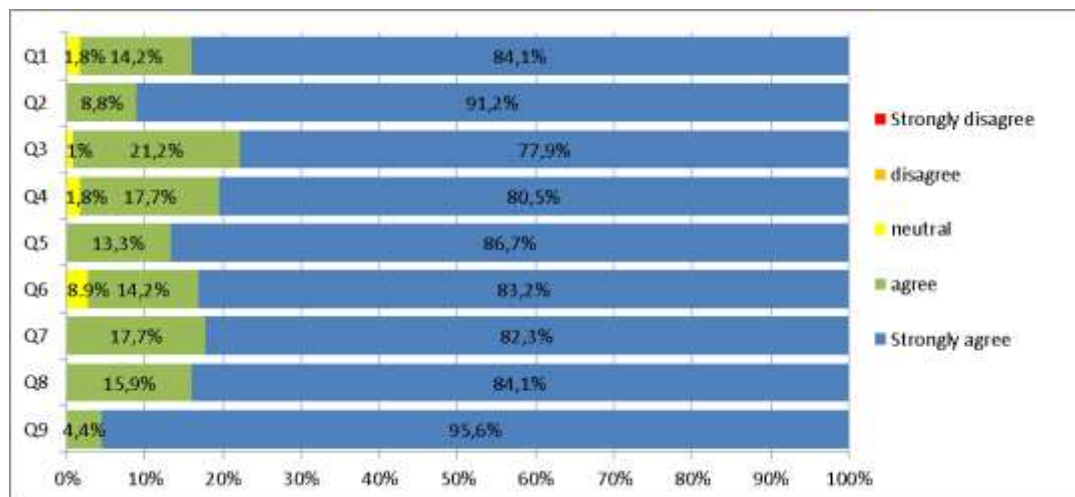
Q9 Θα συνιστούσατε το σύστημα στους συμφοιτητές σας και θα θέλατε να ενσωματωθεί στο μάθημα;

Αρχικά, για τον έλεγχο της αξιοπιστίας, υπολογίστηκε ο δείκτης *Cronbach's α* (alpha) ή *δείκτης εσωτερικής συνέπειας* (*internal consistency coefficient*) (Cronbach, 1951). Οι τιμές του δείκτη κυμαίνονται από 0 μέχρι και 1 και αποτελεί έναν από τους πιο γνωστούς και ευρέως χρησιμοποιούμενους δείκτες εσωτερικής συνέπειας (Hatcher, 1994). Η αξιοπιστία των ερωτηματολογίων υπολογίστηκε να είναι $\alpha = 0.73$ για το ερωτηματολόγιο της ομάδας groupA και $\alpha = 0.82$ για το ερωτηματολόγιο της ομάδας GroupB, κάτι που δείχνει ότι τα ερωτηματολόγια είχαν υψηλή αξιοπιστία. Στην συνέχεια, έγινε η ανάλυση των απαντήσεων των ερωτηματολογίων, η οποία ανέδειξε πολύ θετικά αποτελέσματα τα οποία παρουσιάζονται στις ακόλουθες εικόνες.



Εικόνα 2.4 Τα αποτελέσματα του ερωτηματολογίου της ομάδας GroupA

Τα αποτελέσματα των ερωτηματολογίων έδειξαν ότι το 80.5% των εκπαιδευόμενων του GroupA και το 84% του GroupB βρήκαν ευχάριστη και ενδιαφέρουσα τη χρήση του συστήματος NLtoFOL. Επίσης, το 81.5% των εκπαιδευόμενων του GroupA και 91% του GroupB απάντησαν ότι το σύστημα τους βοήθησε στον φορμαλισμό των προτάσεων και στη εκμάθηση της διαδικασίας. Σχετικά, με την ανάδραση που παράχθηκε από το σύστημα, το 93% των εκπαιδευόμενων του GroupA και 80.5% του GroupB απάντησε ότι η ανάδραση ήταν ακριβής και κατάλληλη.



Εικόνα 2.5 Τα αποτελέσματα του ερωτηματολογίου της ομάδας GroupB

Η διαφορά μεταξύ των απαντήσεων των εκπαιδευόμενων των δυο ομάδων οφείλεται στη χρήση του απλού τύπου της ανάδρασης (flag & bottom out feedback) που ήταν διαθέσιμος για τους εκπαιδευόμενους του GroupA. Όσον αφορά την αναγνώριση των λαθών, το 35.5% του GroupA και το 86.5% της GroupB ανέφερε ότι η ανάδραση βοήθησε στην κατανόηση των λαθών. Επιπλέον, το 76% του GroupA και το 83% του GroupB δήλωσαν ότι οι υποδείξεις που παρείχε το σύστημα τους βοήθησαν να διορθώσουν τα λάθη τους στις απαντήσεις τους. Επιπλέον, οι εκπαιδευόμενοι του GroupB βρήκαν τα μηνύματα κειμένου να είναι υψηλής ποιότητας (80.5%), γραμματικά ορθά και κατανοητά (76%), καθώς επίσης και κατάλληλης έκτασης (84%). Επίσης, βρήκαν τις παρεμβάσεις του συστήματος για υποδείξεις και βοήθειες να είναι ακριβείς και χρήσιμες (79.5%). Τέλος, το 87.5% του GroupA και το 95.5% του Group B πρότειναν ότι το σύστημα πρέπει να ενσωματωθεί στο πρόγραμμα σπουδών και να παρέχεται στους φοιτητές των επόμενων ετών.

2.2.4.2 Ανάλυση των μαθησιακών ενεργειών των εκπαιδευόμενων

Κατά τη διάρκεια της αλληλεπίδρασης των εκπαιδευόμενων με το σύστημα, όλες οι ενέργειές τους καταγράφονταν και αποθηκεύονταν στη βάση του συστήματος. Τα εκπαιδευτικά αυτά δεδομένα αποτελούν μια πολύτιμη πηγή πληροφοριών για την κατανόηση της διαδικασίας μάθησης των εκπαιδευόμενων. Στα πλαίσια των πειραματικών μελετών, πραγματοποιήθηκε μια πολύπλευρη και εις βάθος ανάλυση των εκπαιδευτικών δεδομένων για να αποτιμηθεί η αποτελεσματικότητα και η επίδραση που είχε η ανάδρασης στην μάθηση των εκπαιδευόμενων.

Αρχικά, επιλέχθηκαν 25 ασκήσεις του συστήματος, στις οποίες όλοι οι εκπαιδευόμενοι και των δύο ομάδων είχαν αλληλεπιδράσει και έγινε ανάλυση των ενεργειών των εκπαιδευόμενων. Πιο συγκεκριμένα, υπολογίστηκαν (i) οι συνολικές προσπάθειες και οι απαντήσεις που υπέβαλαν οι εκπαιδευόμενοι και των δύο ομάδων, (ii) τα μηνύματα ανάδρασης που παράχθηκαν από το σύστημα, (iii) οι bottom out hints υποδείξεις που δόθηκαν από το σύστημα καθώς και (iv) η απόδοση των εκπαιδευόμενων μετά την παροχή διαφόρων ειδών υποδείξεων. Στον πίνακα 2.10 παρουσιάζονται οι μετρικές ανάλυσης που προέκυψαν.

Πίνακας 2.10 Ανάλυση συμπεριφοράς εκπαιδευόμενων

	Group A	Group B
Μέσος όρος προσπαθειών	7.2	3.9
Παροχή υποδείξεων (εκτός από flag feedback)	51%	76%
Bottom out	51%	12%

Από την ανάλυση προκύπτει ότι οι εκπαιδευόμενοι του GroupA υπέβαλαν κατά μέσο όρο 7.2 απαντήσεις ανά βήμα κάθε άσκησης. Αυτό δείχνει ότι υπέβαλαν κατά μέσο όρο 6.2 απαντήσεις μετά από μια αρχική απάντηση που δεν ήταν σωστή. Δηλαδή υπέβαλαν κατά μέσο όρο 5.2 λανθασμένες απαντήσεις, ως ότου καθορίσουν και υποβάλουν την τελική σωστή απάντηση. Από την άλλη πλευρά, οι εκπαιδευόμενοι του GroupB υπέβαλαν κατά μέσο όρο 3.9 απαντήσεις ανά βήμα της διαδικασίας. Δηλαδή υπέβαλαν 2.9 απαντήσεις μετά από μια πρώτη λανθασμένη απάντηση και πιο συγκεκριμένα υπέβαλαν κατά μέσο όρο 1.9 απαντήσεις ως ότου καθορίσουν την σωστή απάντηση.

Μια μετρική ιδιαίτερης αξίας αφορά τον υπολογισμό της παρεχόμενης υποστήριξης του συστήματος στις δύο ομάδες εκπαιδευόμενων. Στα πλαίσια της ανάλυσης, θεωρούμε ως υποστήριξη του συστήματος τα μηνύματα ανάδρασης που παρέχονται στον εκπαιδευόμενο κατά τα βήματα των ασκήσεων. Επομένως, στο σενάριο μάθησης της ομάδας GroupA, η υποστήριξη του συστήματος αφορά την παροχή ανάδρασης bottom-out hints, η οποία παρέχει στον εκπαιδευόμενο ένα στοιχείο της σωστής απάντησης ή ολόκληρη τη σωστή απάντηση. Στο σενάριο μάθησης της ομάδας GroupB, η βοήθεια του συστήματος αφορά την παροχή ανάδρασης που βασίζεται στο πλήρες πλαίσιο ανάδρασης (full feedback). Από

τα αποτελέσματα προκύπτει ότι οι εκπαιδευόμενοι της ομάδας GroupA ολοκλήρωσαν το 51% των βημάτων της διαδικασίας με την παροχή ανάδρασης bottom-out hints από το σύστημα. Από την άλλη πλευρά, οι εκπαιδευόμενοι της ομάδας GroupB ολοκλήρωσαν το 76% των βημάτων των ασκήσεων με τις υποδείξεις του δεύτερου και του τρίτου επιπέδου ανάδρασης. Όσον αφορά την παροχή ανάδρασης bottom-out hints, οι εκπαιδευόμενοι της ομάδας GroupB ολοκλήρωσαν το 12% των βημάτων των ασκήσεων βασιζόμενοι σε ανάδραση bottom-out hints σε αντίθεση με την ομάδα GroupA, η οποία ολοκλήρωσε το 51% των ασκήσεων βασιζόμενη σε ανάδραση bottom-out hints. Μάλιστα, από τα αποτελέσματα προκύπτει ότι το δεύτερο και το τρίτο επίπεδο ανάδρασης βοήθησε τους εκπαιδευόμενους της ομάδας GroupB να ολοκληρώσουν το 64% των βημάτων των ασκήσεων χωρίς υποδείξεις τύπου bottom out hints από το 76% των βημάτων των ασκήσεων στις οποίες έλαβαν βοήθεια. Στο δεύτερο και τρίτο επίπεδο της βοήθειας, το σύστημα παρέχει στους εκπαιδευόμενους υποδείξεις τύπου positive, error-specific και procedure feedback. Αυτοί οι τύποι των υποδείξεων βοήθησαν τους εκπαιδευόμενους να διορθώσουν τα λάθη τους και να υποβάλουν σωστές απαντήσεις. Επιπλέον, μπορεί να φανεί ότι οι εκπαιδευόμενοι της ομάδας GroupA είχαν ανάγκη από βοήθεια και καθοδήγηση καθώς υπέβαλαν πολλές λανθασμένες απαντήσεις, οι οποίες ήταν πολύ περισσότερες από τις απαντήσεις των εκπαιδευόμενων της ομάδας GroupB. Επίσης, ολοκλήρωσαν το 51% των βημάτων των ασκήσεων με την παράδοση ενός τμήματος (ή ολόκληρης) της σωστής απάντησης (bottom out hints). Επιπροσθέτως, οι εκπαιδευόμενοι της ομάδας GroupA, μετά από την υποβολή λανθασμένης απάντησης και την παροχή από το σύστημα ανάδρασης flag feedback, στο 51% των βημάτων των ασκήσεων δεν ήταν σε θέση να προσδιορίσουν τη σωστή απάντηση και ζήτησαν βοήθεια bottom out hints. Αυτό δείχνει ότι η ανάδραση flag feedback δεν ήταν τόσο αποτελεσματική στο να τους βοηθήσει να διορθώσουν τις λανθασμένες απαντήσεις τους. Σε αντίθεση, οι εκπαιδευόμενοι της ομάδας GroupB ολοκλήρωσαν το 64% των βημάτων με τις υποδείξεις του δεύτερου και του τρίτου επιπέδου ανάδρασης και μόνο το 12% των βημάτων των ασκήσεων με βοήθεια bottom out hints.

Μια επιπρόσθετη ανάλυση που διεξήχθη, είχε ως στόχο να αποτυπώσει την απόδοση των εκπαιδευόμενων αμέσως μετά τη παροχή κάποιας βοήθειας από το σύστημα. Πιο συγκεκριμένα, στις περιπτώσεις που οι εκπαιδευόμενοι είχαν υποβάλει μια απάντηση που δεν ήταν σωστή, προσδιορίστηκε σε ποιο βαθμό συνέβαλλαν οι υποδείξεις των διαφορετικών επιπέδων ανάδρασης στο να διορθώσουν τα λάθη τους και να προσδιορίσουν τη σωστή απάντηση. Για το σκοπό αυτό έγινε ανάλυση των απαντήσεων που υποβλήθηκαν

αμέσως μετά την παροχή κάποιου είδους ανάδρασης από το σύστημα. Στον Πίνακα 2.11 παρουσιάζονται τα ποσοστά των σωστών απαντήσεων μετά την παροχή υποδείξεων ανά επίπεδο ανάδρασης.

Πίνακας 2.11 Ποσοστό σωστών απαντήσεων ανά επίπεδο παρεχόμενης βοήθειας

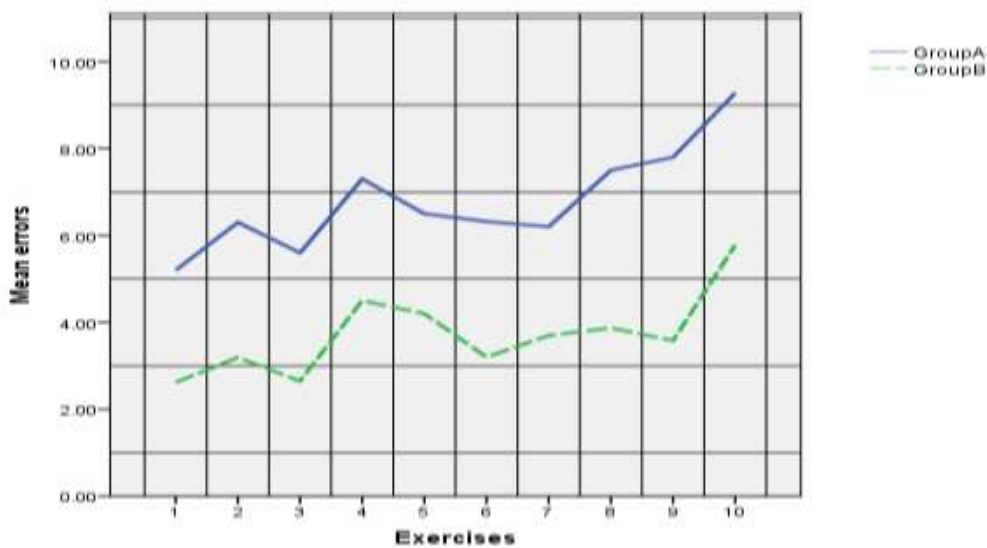
Επίπεδο ανάδρασης	GroupA	GroupB
Flag Feedback	17%	20%
Δεύτερο επίπεδο	-	64%
Τρίτο επίπεδο	-	34%
Τέταρτο επίπεδο	81%	92%

Τα αποτελέσματα δείχνουν ότι η παροχή flag feedback ανάδρασης βοήθησε τους εκπαιδευόμενους της ομάδας GroupA να προσδιορίσουν τη σωστή απάντηση στο 17% των περιπτώσεων και τους εκπαιδευόμενους της ομάδας GroupB στο 20% των περιπτώσεων στις οποίες προσφέρθηκε. Η ανάδραση του δεύτερου επιπέδου όπου παρέχονται υποδείξεις positive feedback και *error specific feedback*, βοήθησε τους εκπαιδευόμενους του GroupB να προσδιορίσουν την σωστή απάντηση στο 64% των περιπτώσεων που προσφέρθηκε. Επίσης, η ανάλυση έδειξε ότι το τρίτο επίπεδο βοήθειας βοήθησε τους εκπαιδευόμενους της ομάδας GroupB να διορθώσουν την απάντησή τους στο 34% των περιπτώσεων που προσφέρθηκε. Τέλος, οι υποδείξεις *bottom out hints* βοήθησαν τους εκπαιδευόμενους της ομάδας GroupA να προσδιορίσουν την απάντησή τους στο 81% των περιπτώσεων και τους εκπαιδευόμενους της ομάδας GroupB στο 92% των περιπτώσεων στις οποίες προσφέρθηκε. Οι περιπτώσεις αυτές αφορούσαν κυρίως απαντήσεις τύπου *inaccurate*, όπου οι υποδείξεις *bottom out hints* διόρθωναν ένα στοιχείο που είχε καθοριστεί λανθασμένα ή απαντήσεις τύπου *incomplete* όπου πρόσφεραν ένα στοιχείο που έλειπε από την απάντηση του εκπαιδευόμενου. Η διαφορά της αποτελεσματικότητας των υποδείξεων *bottom-out hints* στις δύο ομάδες οφείλεται κυρίως στο γεγονός ότι στους εκπαιδευόμενους του GroupB είχαν δοθεί προηγουμένως υποδείξεις από το δεύτερο και το τρίτο επίπεδο ανάδρασης και έτσι σε γενικές γραμμές είχαν λιγότερα λάθη στις απαντήσεις τους, τα οποία ήταν πιο εύκολο να διορθώσουν με τις υποδείξεις *bottom-out hints* του τετάρτου επιπέδου.

Επίσης, η ανάλυση των δεδομένων αποκάλυψε μερικές καταστάσεις παροχής ανάδρασης που μπορεί να μπερδέψουν τους εκπαιδευόμενους. Για παράδειγμα, μια ενδιαφέρουσα

κατάσταση είναι όταν οι εκπαιδευόμενοι έχουν υποβάλει μια απάντηση τύπου *superfluous-accurate*. Από τα αποτελέσματα της ανάλυσης προέκυψε η υποβολή πολλών εσφαλμένων απαντήσεων μετά την παροχή ανάδρασης *flag feedback* σε απαντήσεις τύπου *superfluous-accurate*. Αναλύοντας τις απαντήσεις των εκπαιδευόμενων που υποβλήθηκαν μετά από μια *superfluous-incorrect* απάντηση και της παροχής ανάδρασης *flag feedback*, βρέθηκε ότι στις περισσότερες περιπτώσεις (περίπου 81%) οι εκπαιδευόμενοι προσπάθησαν να διορθώσουν το επιπλέον στοιχείο της απάντησης, το οποίο έχει καθοριστεί με κόκκινο χρώμα ως λανθασμένο. Πιο συγκεκριμένα, προκύπτει ότι η ανάδραση τύπου *flag feedback* δημιούργησε την εσφαλμένη αντίληψη στους εκπαιδευόμενους ότι πρέπει να τροποποιήσουν και να διορθώσουν το στοιχείο που έχει χρωματιστεί με κόκκινο χρώμα αντί να το αφαιρέσουν. Κατά συνέπεια, φαίνεται ότι η παροχή ανάδρασης *flag feedback* σε περιπτώσεις *superfluous-incorrect* απαντήσεων δεν είναι μια καλή επιλογή.

Επιπλέον, πραγματοποιήθηκε μια εις βάθος ανάλυση των δεδομένων των *posttest* των εκπαιδευόμενων των δύο ομάδων. Πιο συγκεκριμένα, έγινε ανάλυση των επιδόσεων των εκπαιδευόμενων των δύο ομάδων στις ασκήσεις του *posttest* και καταγράφηκαν τα λάθη που έγιναν. Ο αριθμός των λαθών που έγιναν μπορεί να αναδείξει την κατανόηση των εκπαιδευόμενων. Το *posttest* περιείχε 10 ασκήσεις και τα επίπεδα δυσκολίας των ασκήσεων ήταν ως εξής: η άσκηση 1 ήταν πολύ εύκολη, οι ασκήσεις 2 και 3 ήταν εύκολες, οι ασκήσεις 4-7 ήταν μέσου επιπέδου, οι ασκήσεις 8 και 9 ήταν δύσκολες και τέλος η άσκηση 10 ήταν προχωρημένου επιπέδου. Για κάθε άσκηση (πρόταση σε φυσική γλώσσα), οι απαντήσεις που δόθηκαν από τους εκπαιδευόμενους (έκφραση ΚΛΠΤ) και των δύο ομάδων αναλύθηκαν και υπολογίστηκε ο αριθμός των λαθών που έγιναν σε κάθε μια άσκηση. Στην εικόνα 2.6 παρουσιάζεται ο μέσος αριθμός των λαθών από τους εκπαιδευόμενους και των δύο ομάδων.

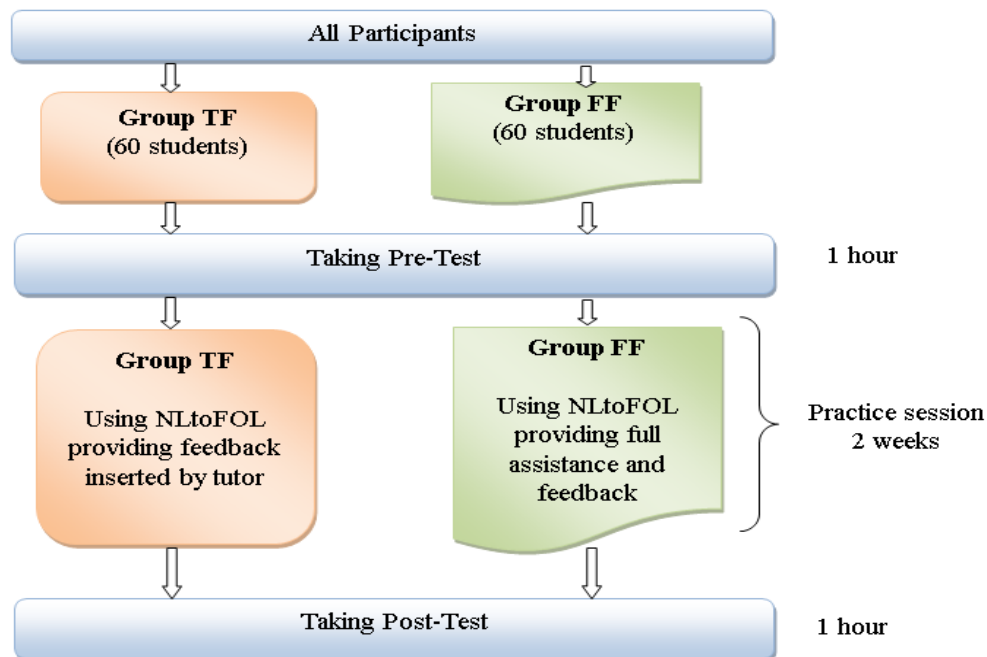


Εικόνα 2.6 Μέσος όρος των λαθών ανά εκπαιδευόμενο για τις δυο ομάδες

Τα αποτελέσματα δείχνουν ότι οι εκπαιδευόμενοι της ομάδας GroupA έκαναν στο posttest περισσότερα λάθη από τους εκπαιδευόμενους της ομάδας GroupB. Ο μέσος αριθμός των λαθών που έγιναν ανά άσκηση από τους εκπαιδευόμενους της ομάδας GroupA κυμαίνονταν από 5 έως 9 λάθη, ενώ ο μέσος όρος λαθών που έγιναν για το GroupB κυμαίνονταν από 2 έως 7 λάθη ανά άσκηση. Η ανάλυση των λαθών δείχνει ότι το πλήρες πλαίσιο ανάδρασης (full feedback framework) βοήθησε τους εκπαιδευόμενους της ομάδας GroupB να έχουν μια καλύτερη κατανόηση της εκπαιδευτικής διαδικασίας.

2.2.4.3 Δεύτερη Πειραματική Μελέτη

Στο δεύτερο πείραμα συμμετείχαν 120 φοιτητές, οι οποίοι χωρίστηκαν τυχαία σε δύο ομάδες. Η πρώτη ομάδα (GroupFF) είχε 60 φοιτητές, οι οποίοι χρησιμοποίησαν το σύστημα NLtoFOL, το οποίο παρείχε πλήρη ανάδραση με βάση το πλαίσιο παροχής ανάδρασης. Η δεύτερη ομάδα (GroupTF) περιείχε 60 φοιτητές που χρησιμοποίησαν το σύστημα NLtoFOL, το οποίο παρείχε πλήρη ανάδραση, η οποία είχε δημιουργηθεί και εισαχθεί στο σύστημα από τον διδάσκοντα. Πιο συγκεκριμένα, ο διδάσκοντας κωδικοποίησε τα μηνύματα ανάδρασης ως πρότυπα κειμένου (text templates), μια παρόμοια προσέγγιση με αυτή που χρησιμοποιείται από τους γνωστικούς εκπαιδευτές (cognitive tutors) (Aleven et al., 2006). Η δομή του δεύτερου πειράματος απεικονίζεται στην Εικόνα 2.7.



Εικόνα 2.7 Διαδικασία Πειράματος

Αρχικά, οι δύο ομάδες συμμετείχαν σε ένα αρχικό τεστ (pretest) το οποίο περιλάμβανε 10 ασκήσεις διαφόρων επιπέδων δυσκολίας. Η βαθμολογία της κάθε απάντησης κυμαίνονταν από 0 έως 10 και το τεστ είχε διάρκεια 1 ώρα. Στην συνέχεια, δόθηκε πρόσβαση στις δύο ομάδες (GroupFF και GroupTF) να χρησιμοποιήσουν το σύστημα NLtoFOL και να μελετήσουν τη διαδικασία του φορμαλισμού για μία εβδομάδα. Οι εκπαιδευόμενοι της ομάδας GroupFF αλληλεπίδρασαν με την έκδοση του συστήματος όπου η παραγωγή της ανάδρασης έγινε σύμφωνα με το πλήρες πλαίσιο ανατροφοδότησης που παράχθηκε από τον μηχανισμό παραγωγής ανάδρασης, ενώ η ομάδα GroupTF αλληλεπίδρασε με την έκδοση του συστήματος όπου τα μηνύματα ανάδρασης είχαν εισαχθεί στο σύστημα από τον διδάσκοντα. Οι δύο ομάδες αλληλεπίδρασαν με το σύστημα για μία εβδομάδα και στην συνέχεια οι εκπαιδευόμενοι συμμετείχαν στο τελικό τεστ (posttest). Το posttest περιελάμβανε 10 ασκήσεις διαφόρων επιπέδων δυσκολίας και είχε διάρκεια μία ώρα.

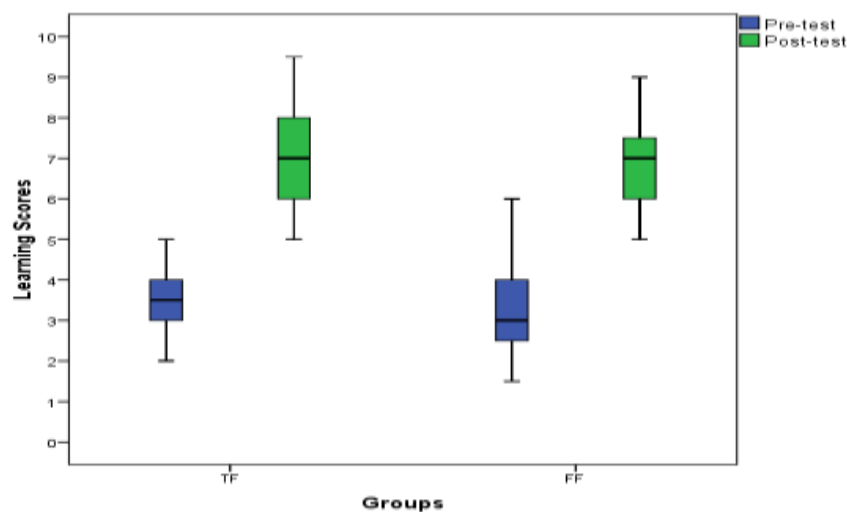
Για να ελέγξουμε και να αξιολογήσουμε τις επιδόσεις των εκπαιδευόμενων, πραγματοποιήθηκε μια ανάλυση συνδιακύμανσης ANCOVA με στόχο να εξεταστεί η διαφορά μεταξύ των δύο ομάδων στα τεστ χρησιμοποιώντας το pretest ως συμμεταβλητή και το posttest ως εξαρτημένη μεταβλητή. Τα αποτελέσματα παρουσιάζονται στον πίνακα 2.12. Οι προσαρμοσμένες μέσες (adjusted mean) τιμές είναι 7.02 για την ομάδα GroupTF και 6.90 για την ομάδα GroupFF.

Πίνακας 2.12 Αποτελέσματα ANCOVA

Groups	N	Mean	SD	Adjusted mean
Group TF (Tutor feedback)	60	7.02	1.01	6.95
Group FF (Full Feedback)	60	6.90	0.98	6.96

Από τα αποτελέσματα προκύπτει ότι δεν υπάρχουν σημαντικές διαφορές μεταξύ των δύο ομάδων (Group TF και GroupFF) ($F = 1.39$ και $p = 0.24 > 0.05$, $\eta^2 = 0.012$), το μέγεθος της επίδρασης υπολογίστηκε ότι ήταν 0.11 και η δύναμη δοκιμής 0.22. Από τα αποτελέσματα προκύπτει ότι και οι δυο ομάδες κατανόησαν και έμαθαν την διαδικασία το ίδιο.

Στην εικόνα 2.8 παρουσιάζεται το θηκόγραμμα των επιδόσεων των δύο ομάδων πριν την δοκιμή (pretest) και μετά την πειραματική δοκιμή (posttest).



Εικόνα 2.8 Κέρδος μάθησης για τις δύο ομάδες.

Τα αποτελέσματα δείχνουν αύξηση στην βαθμολογία των δύο ομάδων, η οποία αυξήθηκε σημαντικά κατά τη χρήση του συστήματος NLtoFOL και στα δύο σενάρια μάθησης. Το πιο ενθαρρυντικό στοιχείο, που δείχνουν τα αποτελέσματα, είναι ότι η ανάδραση και οι υποδείξεις που παράχθηκαν από τον μηχανισμό παραγωγής ανάδρασης του συστήματος και η ανάδραση που εισήχθη από τον διδάσκοντα βοήθησαν τους εκπαιδευόμενους στον ίδιο βαθμό.

3

Αυτόματη Μετατροπή Προτάσεων Φυσικής Γλώσσας σε Λογική

3.1 Εισαγωγή

Τα τελευταία χρόνια, το διαδίκτυο έχει εξελιχθεί σε μια τεράστια αποθήκη ανθρώπινης γνώσης, η οποία επεκτείνεται με εκθετικό ρυθμό. Ένα μεγάλο μέρος της γνώσης στο διαδίκτυο είναι εκφρασμένο σε μορφή κειμένου και η αποτελεσματική ανάλυση και κατανόησή του θέτει σημαντικές προκλήσεις. Καθημερινά, ένας τεράστιος αριθμός νέων άρθρων ειδήσεων, μηνυμάτων και δημοσιεύσεων αποθηκεύονται σε ειδησεογραφικές ιστοσελίδες, blogs και κοινωνικά δίκτυα. Η τεράστια ποσότητα του ηλεκτρονικού περιεχομένου κειμένου στο διαδίκτυο καθιστά αναγκαία την χρήση αυτοματοποιημένων μεθόδων για την ανάλυση, την κατανόηση και την εξαγωγή της γνώσης από αυτό. Ένα βασικό θέμα στην επεξεργασία φυσικής γλώσσας αποτελεί η κατανόηση των προτάσεων και η μετατροπή τους σε δομημένη μορφή. Οι αυτοματοποιημένες μέθοδοι που αναλύουν και μετατρέπουν τη γνώση αυτή σε μια δομημένη μορφή μπορούν να ανοίξουν νέους ορίζοντες στην κατανόηση της φυσικής γλώσσας και στην συλλογιστική πάνω σε προτάσεις φυσικής γλώσσας (Gaur et al., 2014; Janssen, 1996). Η ανάπτυξη προσεγγίσεων για τον φορμαλισμό κειμένου και την εξαγωγή πληροφορίας από κειμενικά δεδομένα είναι εξαιρετικής σημασίας και έχει ένα ευρύ φάσμα εφαρμογών. Η αυτόματη μετατροπή προτάσεων φυσικής γλώσσας σε δομημένη μορφή λογικής αναπαράστασης θα μπορούσε να βοηθήσει σε μεγάλο βαθμό στο σχεδιασμό αυτόματων συστημάτων απάντησης ερωτημάτων όπου με την αναπαράσταση του περιεχομένου των κειμενικών δεδομένων θα μπορούσαν να προσδιοριστούν αυτόματα

ακριβείς απαντήσεις σε ερωτήματα (Weston et al., 2015). Επίσης, εργασίες επεξεργασίας κειμένου, όπως η αυτόματη εξαγωγή περίληψης και η αυτόματη μετάφραση κειμένου θα μπορούσαν να βελτιωθούν σε μεγάλο βαθμό μέσω της αναπαράστασης των κειμενικών δεδομένων φυσικής γλώσσας σε λογική.

Ωστόσο, η ανάπτυξη συστημάτων και προσεγγίσεων για την αυτόματη κατανόηση φυσικής γλώσσας, με στόχο να δομήσουν το περιεχόμενό της, είναι μια πολύ δύσκολη και πολύπλοκη διαδικασία (MacCartney & Manning, 2007). Αρκετές μελέτες έχουν δείξει ότι η ανάλυση και η κατανόηση της φυσικής γλώσσας θεωρείται ένα πολύ σύνθετο πρόβλημα και από τα δυσκολότερα στον τομέα της Τεχνητής Νοημοσύνης, το οποίο θεωρείται ότι είναι AI-complete (Shapiro, 1992; Wong & Mooney, 2007; Yampolskiy, 2012). Πράγματι, η μετατροπή της φυσικής γλώσσας σε Κατηγορηματική Λογική Πρώτης Τάξης (ΚΛΠΤ) είναι μια ad-hoc διαδικασία για την οποία δεν υπάρχει κανένας αλγόριθμος που να την μοντελοποιεί αυτόματα σε ένα υπολογιστικό σύστημα, ενώ απαιτεί εις βάθος κατανόηση του περιεχομένου και του γνωστικού υπόβαθρου των εννοιών που εμφανίζονται στο κείμενο.

Ο φορμαλισμός προτάσεων φυσικής γλώσσας και η επιτυχημένη μετατροπή τους σε Κατηγορηματική Λογική Πρώτης Τάξης (ΚΛΠΤ) μπορεί να συμβάλλει σε πολλούς τομείς, όπως για παράδειγμα σε υψηλού επιπέδου γνωστικές λειτουργίες των ευφυών πρακτόρων και στην καλύτερη επικοινωνία μεταξύ ανθρώπου και υπολογιστικών συστημάτων (Kamp & Reyle, 2013). Οι λογικές εκφράσεις κατηγορηματικής λογικής πρώτης τάξης που εξάγονται από προτάσεις φυσικής γλώσσας μπορούν να αντιπροσωπεύουν την σημασιολογία και πληροφορίες του κειμένου και μπορούν να επιτρέπουν τον αυτόματο συλλογισμό και την εξαγωγή συμπερασμάτων από κειμενικά δεδομένα. Δεδομένου ότι η πλήρης κατανόηση της φυσικής γλώσσας παραμένει ένας μακρινός στόχος, μια γρήγορη και αξιόπιστη προσέγγιση για την μετατροπή προτάσεων φυσικής γλώσσας σε λογική θα μπορούσε να επιτρέψει ένα ευρύ φάσμα άμεσων εφαρμογών.

Στα πλαίσια της ενότητας αυτής, παρουσιάζουμε μία νέα, καινοτόμα προσέγγιση για τον φορμαλισμό προτάσεων φυσικής γλώσσας και την αυτόματη μετατροπή τους σε κατηγορηματική λογική πρώτης τάξης. Η προσέγγιση που εισάγεται, βασίζεται στην αρχή ότι προτάσεις φυσικής γλώσσας που έχουν ίδια ή παρόμοια συντακτική δομή και δέντρα εξαρτήσεων (dependency trees) θα έχουν ίδια ή παρόμοια δομή έκφρασης σε κατηγορηματική λογική πρώτης τάξης. Η προσέγγιση βασίζεται σε κλασικές αρχές της λογικής και της φιλοσοφίας της γλώσσας και κυρίως στην αρχή της σύνθεσης (compositionality) του νοήματος, όπου η σημασιολογία μιας πρότασης προκύπτει ως

συνάρτηση της σημασιολογίας των τμημάτων της (De Mantaras et al., 2005; Janseen, 1996). Αρχικά, στα πλαίσια της προσέγγισης, γίνεται μια εις βάθος ανάλυση της νέας πρότασης φυσικής γλώσσας προς μετατροπή σε κατηγορηματική λογική και δημιουργείται το δέντρο εξαρτήσεων της. Στην συνέχεια ακολουθείται μια μεθοδολογία συλλογιστικής βασισμένη στις περιπτώσεις (case-based reasoning), όπου γίνεται αξιοποίηση υπάρχουσας γνώσης, η οποία αφορά προτάσεις που έχουν μετατραπεί σε λογική, προκειμένου να ανακτηθούν παρόμοιες προτάσεις και να κατευθύνουν τον φορμαλισμό της νέας πρότασης. Η ομοιότητα μεταξύ των προτάσεων φυσικής γλώσσας υπολογίζεται με βάση την *απόσταση μετασχηματισμού δέντρου* (tree edit distance), στα δέντρα εξαρτήσεων τους. Επίσης, στην ενότητα αυτή παρουσιάζεται μια μεθοδολογία για τον προσδιορισμό της δυσκολίας μετατροπής προτάσεων φυσικής γλώσσας σε ΚΛΠΤ. Ο προσδιορισμός της δυσκολίας των προτάσεων αποτελεί μια βασική παράμετρο για την αυτόματη μετατροπή προτάσεων φυσικής γλώσσας σε ΚΛΠΤ και βασίζεται στα χαρακτηριστικά της πρότασης φυσικής γλώσσας καθώς και σε χαρακτηριστικά της αντίστοιχης ΚΛΠΤ της πρότασης.

3.2 Σχετικές Εργασίες

Η μετατροπή προτάσεων φυσικής γλώσσας σε δομημένες αναπαραστάσεις και σε λογική μπορεί να συμβάλει στην καλύτερη απόδοση εφαρμογών και διαδικασιών σε πολλά πεδία της επεξεργασίας φυσικής γλώσσας. Δεδομένου της σημασίας και της συμβολής της διαδικασίας μετατροπής, διάφορες προσεγγίσεις έχουν προταθεί στην διεθνή βιβλιογραφία.

Στην εργασία (Zettlemoyer & Collins, 2012), οι συγγραφείς παρουσιάζουν μια εργασία πάνω στην μετατροπή προτάσεων σε λάμδα-λογισμό (lambda calculus) η οποία βασίζεται στις κατηγορηματικές γραμματικές. Πιο συγκεκριμένα, στα πλαίσια της προσέγγισης, δίνεται ως είσοδο μια σειρά από προτάσεις φυσικής γλώσσας μαζί με τις αντίστοιχες εκφράσεις τους σε λάμδα-λογισμό και γίνεται προσδιορισμός συνδυαστικής κατηγορικής γραμματικής (combinatory categorical grammar) η οποία χρησιμοποιείται για να γίνουν αντιστοιχίσεις μεταξύ της φυσικής γλώσσας και των λογικών εκφράσεων. Στην εργασία (Wong & Mooney 2007), οι συγγραφείς παρουσιάζουν μια προσέγγιση στην αυτόματη εκμάθηση σύγχρονων γραμματικών (synchronous grammars) οι οποίες θα μπορούν να χρησιμοποιηθούν για να παράγουν λογικές εκφράσεις. Στα πλαίσια της προσέγγισης, προτάσεις φυσικής γλώσσας μαζί με τις αντίστοιχες λογικές εκφράσεις τους χρησιμοποιούνται για να εκπαιδευτεί ένας σημασιολογικό αναλυτής (Semantic Parser) με βάση τεχνικής στατιστικής μηχανικής μετάφρασης. Στην εργασία (Costantini et al., 2011), οι συγγραφείς παρουσιάζουν ένα πλαίσιο για την ανάλυση προτάσεων φυσικής γλώσσας

και την μετατροπή τους σε λογισμό. Η ανάλυση των προτάσεων φυσικής γλώσσας γίνεται με εργαλεία μορφοσυντακτικής ανάλυσης και αναλυτές που προσδιορίζουν την γραμματική δομή της πρότασης και τις εξαρτήσεις. Στην συνέχεια, οι συγγραφείς χρησιμοποιούν την βάση DBPedia για να ανακτήσουν πληροφορίες για τις λέξεις και χρησιμοποιούν αξιώματα για να μετατρέψουν τις προτάσεις σε λ-λογισμό. Στην εργασία (Gaur et al., 2014) οι συγγραφείς παρουσιάζουν ένα σύστημα για την μετατροπή απλών προτάσεων φυσικής γλώσσας σε λ-λογισμό. Η προσέγγιση των συγγραφέων βασίζεται στην εκμάθηση συνδυαστικής κατηγορικής γραμματικής (Combinatory Categorical Grammar), από ζεύγη προτάσεων φυσικής γλώσσας μαζί με τις λογικές εκφράσεις τους. Στην εργασία, (Baral et al., 2011) οι συγγραφείς παρουσιάζουν ένα σύστημα για την μετατροπή προτάσεων φυσικής γλώσσας σε λ-λογισμό. Το σύστημα βασίζεται σε αντίστροφες διαδικασίες μετάφρασης λ-λογισμού και βασιζόμενο σε αυτές μπορεί να δέχεται ως είσοδο σημασιολογικές αναπαραστάσεις λέξεων και προτάσεων και να προσδιορίζει με βάσει αυτές τις σημασιολογικές αναπαραστάσεις άλλων λέξεων και προτάσεων. Το σύστημα χρησιμοποιεί αναλυτή συνδυαστικής κατηγορικής γραμματικής για να αναλύσει τις προτάσεις και να δημιουργήσει τις σημασιολογικές αναπαραστάσεις τους σε συνδυασμό με στατιστικές τεχνικές που χρησιμοποιούνται για να αναπαραστήσουν την σημασία λέξεων που εμφανίζουν πολλές έννοιες. Οι συγγραφείς στην εργασία (Kate & Mooney, 2006) παρουσιάζουν μια προσέγγιση για την μετατροπή προτάσεων φυσικής γλώσσας σε λογική αναπαράσταση η οποία βασίζεται σε SVM ταξινομητές. Στα πλαίσια της προσέγγισης, οι ταξινομητές χρησιμοποιούν ως πυρίνα (kernel) την ομοιότητα ακολουθιών (string similarity) οι λογικές αναπαραστάσεις νέων προτάσεων προσδιορίζονται από τους ταξινομητές με βάση την πιο πιθανή σημασιολογική ανάλυση.

Ένα βασικό χαρακτηριστικό των εργασιών και των σχετικών ερευνητικών προσπαθειών, είναι ότι δεν χρησιμοποιούν την γνώση που απορρέει από την αρχή ότι προτάσεις με παρόμοια δομή και δέντρα εξαρτήσεων θα έχουν επίσης παρόμοιες λογικές αναπαραστάσεις. Σε αυτή την κατεύθυνση, η προσέγγιση που αναπτύχθηκε και παρουσιάζεται στο κεφάλαιο αυτό αποτελεί μια καινοτόμα προσέγγιση που προσπαθεί να βασιστεί και να χρησιμοποιήσει με αποτελεσματικό τρόπο την γνώση που απορρέει από αυτή την αρχή. Πιο συγκεκριμένα, στα πλαίσια της προσέγγισης αρχικά πραγματοποιείται μια εις βάθος ανάλυση της φυσικής γλώσσας για να εξαχθούν κατάλληλα χαρακτηριστικά από μια πρόταση και να καθοριστεί η δομή της και στην συνέχεια ακολουθείται μια μεθοδολογία συλλογιστικής βασισμένης στις περιπτώσεις (Case based Reasoning) για να

αξιοποιήσει την υπάρχουσα γνώση που σχετίζεται με ήδη φορμαρισμένες προτάσεις και να οδηγήσει την διαδικασία μετασχηματισμού νέων προτάσεων.

3.3 Βασικές Έννοιες

3.3.1 Σύνταξη και Σημασιολογία Κατηγορηματικής Λογικής

Η λογική παρέχει έναν τρόπο για την αποσαφήνιση και την τυποποίηση της διαδικασίας της ανθρώπινης σκέψης και προσφέρει μια σημαντική και εύχρηστη μεθοδολογία για την αναπαράσταση και επίλυση προβλημάτων. Η ανάγκη χρήσης μιας αυστηρά ορισμένης γλώσσας, με τη μαθηματική έννοια, προήλθε από την ακαταλληλότητα της χρήσης της φυσικής γλώσσας σε υπολογιστικά συστήματα. Σε αντίθεση με τη φυσική γλώσσα, η λογική προσφέρει μια σαφή, ακριβή και απλή στη σύνταξη γλώσσα, καθώς και την δυνατότητα παραγωγής νέας γνώσης από την ήδη υπάρχουσα. Η λογική ορίζεται σαν τη μελέτη της σωστής εξαγωγής συμπερασμάτων. Μια κατ' ελάχιστον απαίτηση, για τη σωστή εξαγωγή συμπερασμάτων, είναι η διατήρηση της αλήθειας, δηλαδή η απαίτηση από αληθείς υποθέσεις - προτάσεις να εξάγονται αληθή συμπεράσματα - προτάσεις.

Η κατηγορηματική Λογική Πρώτης Τάξης (ΚΛΠΤ) είναι η πιο ευρέως χρησιμοποιούμενη λογική ως γλώσσα αναπαράσταση γνώσης. Οι λογικές εκφράσεις ή προτάσεις αντιπροσωπεύουν ιδιότητες ή σχέσεις μεταξύ των οντοτήτων του κόσμου, που αντιπροσωπεύονται από το σύνολο D . Η σύνταξη της ΚΛΠΤ ορίζεται ως εξής:

- *Σταθερές (constants)*: κάθε σταθερά (constant) αναπαριστά μια συγκεκριμένη οντότητα του συνόλου D (π.χ. Maria, Pluto).
- *Μεταβλητές (variables)*: κάθε μεταβλητή (variable) αναπαριστά μια ή περισσότερες οντότητες του συνόλου D . Κάθε μεταβλητή (π.χ. x, y, z) μπορεί να πάρει μία τιμή κάθε φορά από ένα υποσύνολο του D .
- *Συναρτησιακά σύμβολα (function symbols)*: κάθε συνάρτηση αναπαριστά μια μια-προς-μια σχέση μεταξύ δυο οντοτήτων του συνόλου D (π.χ. mother-of(John), color-of(sky)).
- *Κατηγορήματα (predicates)*: κάθε κατηγορηματικό αναπαριστά μια ιδιότητα της οντότητας ή μια σχέση μεταξύ δυο ή περισσότερων οντοτήτων του D (π.χ. tall, loves).
- *Λογικά διασυνδετικά (connectives)*: η άρνηση “ $\neg$ ”, η διάζευξη “ $\vee$ ” (λογικό «ή»), η σύζευξη “ $\wedge$ ” (λογικό «και»), η συνεπαγωγή “ $\Rightarrow$ ” και η διπλή συνεπαγωγή “ $\Leftrightarrow$ ” (λογική ισοδυναμία).

- *Ποσοδείκτες* (quantifiers): Καθολικός ποσοδείκτης “ $\forall$ ” (universal quantifier) και Υπαρξιακός “ $\exists$ ” (existential quantifier).
 - $(\forall x)P(x)$ σημαίνει ότι η έκφραση P ισχύει για όλες τις τιμές του x στο πεδίο που σχετίζεται με την εν λόγω μεταβλητή. Π.χ.
 $(\forall x) \text{dolphin}(x) \Rightarrow \text{mammal}(x)$
 - $(\exists x)P(x)$ σημαίνει ότι η έκφραση P ισχύει για τουλάχιστον μία τιμή του x στο πεδίο που συσχετίζεται με την εν λόγω μεταβλητή. Π.χ.
 $(\exists x) \text{number}(x) \wedge \text{even}(x)$

Μια *ατομική έκφραση* ή ένα *άτομο* στην ΚΛΠΤ είναι μια έκφραση της μορφής $P(t_1, \dots, t_n)$, όπου P ένα κατηγορημα n ορισμάτων και $t_1, \dots, t_n$ είναι όροι. Ένας όρος (term) μπορεί να είναι μια σταθερά ή μια μεταβλητή ή μια συνάρτηση. Αν f είναι ένα συναρτησιακό σύμβολο n τάξης και $t_1, \dots, t_n$ είναι όροι, τότε $f(t_1, \dots, t_n)$ είναι όρος.

Μια *ορθά δομημένη έκφραση* (well formed formula-wff) στην ΚΛΠΤ ορίζεται ως εξής:

- Ένα άτομο είναι μια wff.
- Αν P είναι μια wff, τότε $\neg P$ είναι μια wff.
- Αν P και Q είναι wffs, τότε $(P \vee Q)$, $(P \wedge Q)$, $(P \Rightarrow Q)$ και $(P \Leftrightarrow Q)$ είναι wffs.
- Αν P είναι μια wff και x μια ελεύθερη μεταβλητή στην P , τότε $(\forall x)P$, $(\exists x)P$ είναι wff
- Το σύνολό των wffs δημιουργείται μόνο από πεπερασμένο αριθμό εφαρμογών των (i), (ii), (iii).

Κάθε ποσοδείκτης έχει μια εμβέλεια (scope), η οποία είναι η έκφραση στην οποία εφαρμόζεται ο ποσοδείκτης. Κάθε ποσοδείκτης συσχετίζεται με μια μεταβλητή και όλες οι εμφανίσεις μιας μεταβλητής συσχετίζονται με ένα ποσοδείκτη.

Οι ακόλουθες προτάσεις ΚΛΠΤ είναι wffs, όπου δίπλα σε κάθε μια φαίνεται η αντίστοιχη πρόταση ΦΓ που αναπαριστούν:

“ $(\exists x) \text{apple}(x) \wedge \text{red}(x)$ ” $\rightarrow$ “There is a red apple”

“ $(\forall x) \text{gardener}(x) \Rightarrow \text{likes}(x, \text{sun})$ ” $\rightarrow$ “Every gardener likes the sun”

“ $(\forall x) (\text{cat}(x) \vee \text{dog}(x)) \Rightarrow \text{pet}(x)$ ” $\rightarrow$ “Cats and dogs are pets”

3.3.2 Οι Δυσκολίες Αυτοματοποίησης του Φορμαλισμού Φυσικής Γλώσσας

Η μετατροπή προτάσεων φυσικής γλώσσας (ΦΓ) σε εκφράσεις ΚΛΠΤ είναι μια διαδικασία που ονομάζεται *λογικός φορμαλισμός* της φυσικής γλώσσας και είναι μια ad-hoc διαδικασία, για την οποία δεν υπάρχει κάποιος αλγόριθμος που να μπορεί να την αυτοματοποιήσει. Αυτό οφείλεται κυρίως στο γεγονός, ότι η φυσική γλώσσα δεν έχει σαφή σημασιολογία, σε αντίθεση με την ΚΛΠΤ. Σε πολλές περιπτώσεις, οι προτάσεις στον γραπτό και κυρίως στον προφορικό λόγο δεν ακολουθούν γραμματικούς κανόνες με αποτέλεσμα να μην είναι καλά δομημένες και να έχουν ασάφειες, που προσθέτουν επιπλέον δυσκολία στον φορμαλισμό τους. Σε ορισμένες περιπτώσεις, οι προτάσεις φυσικής γλώσσας μπορεί να αναφέρονται σε αφηρημένες οντότητες, όπως αυτές που εκφράζονται με λέξεις όπως “someone”, “something”, “somebody”, και επομένως ο προσδιορισμός του πλήθους και του είδους των ορισμάτων των κατηγορημάτων στην ΚΛΠΤ μπορεί να μην είναι σαφές πώς μπορεί να πραγματοποιηθεί. Επιπλέον, οι ποσοδείκτες των κατηγορημάτων μπορεί να αναφέρονται σε οντότητες και αντικείμενα της πρότασης τα οποία μπορεί να είναι και αυτά ασαφή στον προσδιορισμό. Για παράδειγμα, στις ακόλουθες προτάσεις “All animals eat some plants” και “John like someone who is vegetarian”, οι λέξεις “some” και “someone” παρότι έχουν την ίδια σημασία, οδηγούν σε διαφορετικούς ποσοδείκτες, δεδομένου ότι στην δεύτερη πρόταση η λέξη “someone” υποδηλώνει το γενικό “all”. Επομένως, οι ΚΛΠΤ εκφράσεις των παραπάνω προτάσεων είναι οι εξής:

$$(\forall x)(\exists y)animal(x) \wedge plant(y) \Rightarrow eats(x, y)$$

$$(\forall x)vegetarian(x) \Rightarrow likes(John, x)$$

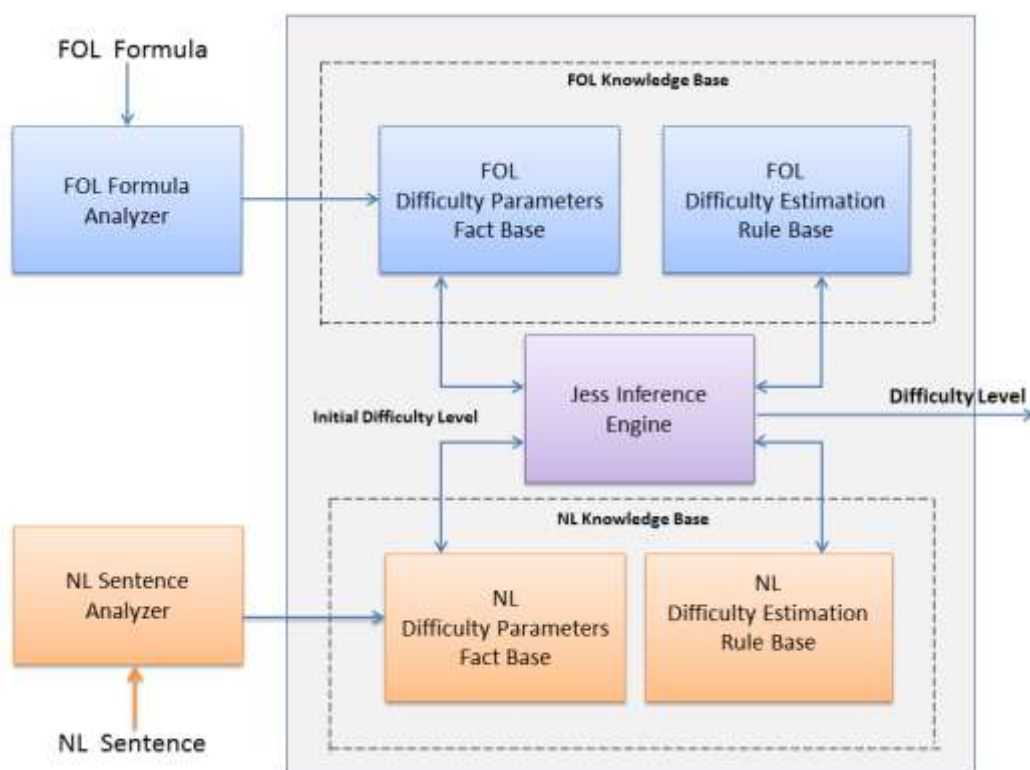
Επίσης, λέξεις που εκφράζουν άρνηση μπορούν να προκαλέσουν επιπλέον δυσκολία στον φορμαλισμό της πρότασης. Πιο συγκεκριμένα, παράγοντες που συμβάλουν στην δυσκολία φορμαλισμού αποτελούν (α) η ύπαρξη ή όχι άρνησης σε μια έκφραση ΚΛΠΤ και (β) η εμβέλεια της άρνησης, δηλαδή το αν η άρνηση αναφέρεται σε ένα άτομο, μια ομάδα ατόμων ή ακόμα και σε όλη την έκφραση. Για παράδειγμα, η ακόλουθη πρόταση “There are no mushrooms that are poisonous and purple” μπορεί λανθασμένα να μεταφραστεί ως “ $(\forall x) mushroom(x) \Rightarrow \neg poisonous(x) \wedge \neg purple(x)$ ” αντί της σωστής έκφρασης “ $(\forall x) mushroom(x) \Rightarrow \neg (poisonous(x) \wedge purple(x))$ ”. Η εμβέλεια της άρνησης δημιουργεί δυσκολίες κατά την διάρκεια του φορμαλισμού της πρότασης. Μια άλλη περίπτωση ασάφειας αφορά τον προσδιορισμό των διασυνδετικών. Για παράδειγμα, η ύπαρξη ενός “and” σε μια πρόταση φυσικής γλώσσας μπορεί να οδηγήσει λανθασμένα στη χρήση του

διασυνδετικού “ $\wedge$ ” στον φορμαλισμό της. Ας θεωρήσουμε την πρόταση “Apples and oranges are fruits” και τη μετάφρασή της στην έκφραση ΚΛΠΤ “ $(\forall x) (\text{apple}(x) \wedge \text{orange}(x)) \Rightarrow \text{fruit}(x)$ ”. Η μετάφραση αυτή δεν είναι σωστή καθώς δεν μπορεί μια οντότητα να είναι ταυτόχρονα “apple” και “orange” αντί της σωστής “ $(\forall x) (\text{apple}(x) \vee \text{orange}(x)) \Rightarrow \text{fruit}(x)$ ”. Τα παραπάνω παραδείγματα δείχνουν ότι η ακριβής μετάφραση μιας πρότασης φυσικής γλώσσας σε ΚΛΠΤ μπορεί να είναι πολύπλοκη και δύσκολη διαδικασία, η οποία απαιτεί εις βάθος κατανόηση της σημασίας της φυσικής γλώσσας, γεγονός που δημιουργεί δυσκολίες στις αυτοματοποιημένες προσεγγίσεις.

3.4 Προσδιορισμός Πολυπλοκότητας Φορμαλισμού Προτάσεων Φυσικής Γλώσσας

Σε αυτή την ενότητα, παρουσιάζουμε την προσέγγιση και το σύστημα που αναπτύχθηκε για τον προσδιορισμό του επιπέδου δυσκολίας της διαδικασίας φορμαλισμού προτάσεων φυσικής γλώσσας. Πιο συγκεκριμένα, γίνεται περιγραφή των αρχών που βασίστηκε και των μεθόδων που χρησιμοποιήθηκαν για τον ακριβή προσδιορισμό της δυσκολίας φορμαλισμού των προτάσεων. Ο προσδιορισμός του επιπέδου δυσκολίας πραγματοποιεί εις βάθος ανάλυση της πρότασης φυσικής γλώσσας καθώς και ανάλυση της αντίστοιχης έκφρασης ΚΛΠΤ της πρότασης και εξαγωγή των κατάλληλων χαρακτηριστικών που σχετίζονται με τη δυσκολία και την πολυπλοκότητα φορμαλισμού της πρότασης.

Αρχικά, γίνεται μια εκτίμηση της δυσκολίας με βάση την πολυπλοκότητα της έκφρασης ΚΛΠΤ. Το σύστημα παίρνει ως είσοδο παραμέτρους, όπως τον αριθμό, το είδος και τη θέση των ποσοδεικτών, τον αριθμό των συνεπαγωγών και τον αριθμό των διαφορετικών συνδέσμων. Στη συνέχεια, ο τελικός προσδιορισμός του επιπέδου δυσκολίας γίνεται με βάση τόσο σημασιολογικές πτυχές της πρότασης φυσικής γλώσσας όσο και με βάση την έκφραση ΚΛΠΤ. Το σύστημα προσδιορισμού δυσκολίας ταξινομεί μια πρόταση με βάση τις παραμέτρους σε ένα από πέντε επίπεδα δυσκολίας, που είναι τα εξής: “very easy”, “easy”, “medium”, “difficult”, “advanced”.



Εικόνα 3.1. Αρχιτεκτονική Συστήματος Προσδιορισμού Δυσκολίας Φορμαλισμού

Αρχικά, το σύστημα δέχεται ως είσοδο μια έκφραση ΚΛΠΤ, την οποία αναλύει μέσω της μονάδας *Αναλυτής Εκφράσεων ΚΛΠΤ (FOL Formula Analyzer)* (Perikos et al., 2011a). Η μονάδα αυτή αναλύει την έκφραση ΚΛΠΤ και προσδιορίζει τις τιμές των παραμέτρων που σχετίζονται με την πολυπλοκότητα της άσκησης και στην συνέχεια εισάγει τα αντίστοιχα γεγονότα στην *Βάση Γεγονότων Παραμέτρων Δυσκολίας ΚΛΠΤ (FOL Difficulty Parameters Fact Base)*. Από την άλλη πλευρά, η μονάδα *Αναλυτής Προτάσεων ΦΓ (NL Sentence Analyzer)* δέχεται ως είσοδο τη πρόταση φυσικής γλώσσας και προσδιορίζει τις τιμές των παραμέτρων που σχετίζονται με τη σημασιολογία της πρότασης και εισάγει τα αντίστοιχα γεγονότα στη *Βάση Γεγονότων Παραμέτρων Δυσκολίας ΦΓ (NL Difficulty Parameters Fact Base)*.

Το σύστημα αρχικά προσδιορίζει ένα πρωταρχικό επίπεδο δυσκολίας με βάση την έκφραση ΚΛΠΤ, χρησιμοποιώντας τη *Βάση Γεγονότων Παραμέτρων Δυσκολίας ΚΛΠΤ (FOL Difficulty Parameters Fact Base)* και τους κανόνες από τη *Βάση Κανόνων Εκτίμησης Δυσκολίας ΚΛΠΤ (FOL Difficulty Estimation Rule Base)*. Στην συνέχεια, η εξαγωγή του τελικού επιπέδου δυσκολίας βασίζεται στον αρχικό προσδιορισμό του επιπέδου δυσκολίας, στα γεγονότα από την βάση παραμέτρων γεγονότων δυσκολίας ΦΓ (*NL difficulty*

Parameters Fact Base) και σε κανόνες από τη βάση κανόνων εκτίμησης δυσκολίας (*NL Difficulty Estimation Rule Base*).

Επομένως, ο προσδιορισμός της δυσκολίας αποτελείται από δυο έμπειρα συστήματα. Το πρώτο έμπειρο σύστημα εξάγει συμπεράσματα για το επίπεδο δυσκολίας χρησιμοποιώντας τη βάση γνώσης που σχετίζεται με την πολυπλοκότητα της έκφρασης ΚΛΠΤ και το δεύτερο έμπειρο σύστημα εξάγει συμπεράσματα για το επίπεδο δυσκολίας βασιζόμενο στον αρχικό προσδιορισμό της δυσκολίας του πρώτου έμπειρου συστήματος και σε παραμέτρους που προέκυψαν από την ανάλυση της δομής και της σημασιολογίας της πρότασης φυσικής γλώσσας.

Η ανάπτυξη του έμπειρου συστήματος του βασισμένου στην ΚΛΠΤ έγινε σε συνεργασία με εμπειρογνώμονες-ειδικούς. Στα πλαίσια της συνεργασίας αυτής, καθορίστηκαν ποιοι παράμετροι μιας έκφρασης ΚΛΠΤ και σε ποιο βαθμό είναι σημαντικές για τον προσδιορισμό της δυσκολίας της μετατροπής των προτάσεων φυσικής γλώσσας σε ΚΛΠΤ. Μετά από διεξοδική μελέτη, καταλήξαμε στις ακόλουθες παραμέτρους: (α) Το πλήθος και ο τύπος των ποσοδεικτών, (β) Το πλήθος των συνεπαγωγών και (γ) Το πλήθος των διαφορετικών διασυνδετικών ($\Rightarrow, \vee, \wedge, \neg$). Ο προσδιορισμός του επιπέδου δυσκολίας της διαδικασίας μετατροπής γίνεται με βάση τους κανόνες που παρουσιάζονται στον Πίνακα 3.1. Πιο συγκεκριμένα, μια έκφραση ΚΛΠΤ, σε αντίθεση με μια πρόταση ΦΓ, έχει ελεγχόμενη και αυστηρή σύνταξη, κάτι που καθιστά ευκολότερη την ανάλυση της και σύμφωνα με τους ειδικούς αποτελεί μια καλή επιλογή για την εξόρυξη σημαντικών χαρακτηριστικών γνωρισμάτων.

Πίνακας 3.1 Κανόνες για τον προσδιορισμό της δυσκολίας με βάση την ΚΛΠΤ έκφραση

#	Πλήθος των συνεπαγωγών ($\Rightarrow$)	Πλήθος των Ποσοδεικτών ($\exists \forall$)	Καθολικός Ποσοδείκτης ($\forall$)	Υπαρξιακός Ποσοδείκτης ($\exists$)	Πλήθος των διαφορετικών διασυνδετικών ($\wedge \vee$)	Επίπεδο Δυσκολίας
1	<2	0	*	*	*	Very Easy
2	≥ 2	0	*	*	*	Easy
3	0	1	*	*	<3	Easy
4	1		yes	no	≤ 1	Easy
5	1 ^d	>1 ^d	no	yes	$\geq 3^d$	Medium
6	0	>1 ^d	*	*	$\geq 3^d$	Medium
7	1	*	yes	no	2	Medium
8	1	≤ 2	yes	no	≥ 3	Medium

9	≥ 2	1	*	*	<2	Medium
10	1	*	yes	yes	*	Difficult
11	1	>2	yes	no	≥ 3	Difficult
12	≥ 2	1	*	*	≥ 2	Difficult
13	≥ 2	>1	*	*	*	Advanced

^d These rule premises are combined with disjunction

Ο προσδιορισμός της δυσκολίας του φορμαλισμού των προτάσεων βασίζεται στην ανάλυση των χαρακτηριστικών της έκφρασης ΚΛΠΤ καθώς και στην ανάλυση της φυσικής γλώσσας. Πιο συγκεκριμένα, γίνεται ανάλυση της πρότασης ΦΓ και εξάγονται μια σειρά από χαρακτηριστικά, τα οποία οργανώνονται σε δύο κατηγορίες. Η πρώτη κατηγορία περιλαμβάνει χαρακτηριστικά τα οποία εξάγονται αποκλειστικά από τη ανάλυση της φυσικής γλώσσας, ενώ η δεύτερη κατηγορία περιλαμβάνει χαρακτηριστικά, οι τιμές των οποίων προσδιορίζονται αναλύοντας την πρόταση φυσικής γλώσσας και συγκρίνοντας την με την αντίστοιχη έκφραση ΚΛΠΤ. Τα χαρακτηριστικά της δεύτερης κατηγορίας προσδιορίζουν την πολυπλοκότητα της πρότασης και απαιτούν επιπρόσθετη επεξεργασία για την μετάφραση της πρότασης σε ΚΛΠΤ. Τα γνωρίσματα που εξάγονται προκύπτουν από τη μορφοσυντακτική ανάλυση της πρότασης φυσικής γλώσσας και είναι τα εξής:

Anaphoric Connectives: Η ύπαρξη αναφορικών εκφράσεων σε μια πρόταση προσθέτει δυσκολία στον φορμαλισμό της, καθώς απαιτείται να γίνει απόαναφοροποίηση και προσδιορισμός της οντότητας στην οποία αναφέρεται το κάθε αναφορικό συνδετικό. Τα αναφορικά συνδετικά αφορούν την ύπαρξη λέξεων σε μια πρόταση όπως : “he”, “she”, “it”, “they”, “himself”, “herself”, “who”, “that”, “which”, “who”, “whose”. Συνεπώς, η τιμή της παραμέτρου είναι ο αριθμός των αναφορικών συνδετικών που εμφανίζονται στην πρόταση. Για παράδειγμα, ας θεωρήσουμε την πρόταση:

“There is a song **that** every person **who** hears **it** loves **it** or hates **it**.”

και την αντίστοιχη έκφραση ΚΛΠΤ:

$$(\exists x) \text{ song}(x) \wedge ((\forall y) ((\text{person}(y) \wedge \text{hears}(y,x)) \Rightarrow (\text{loves}(y,x) \vee \text{hates}(y,x))))$$

Η πρόταση περιέχει πέντε αναφορικά συνδετικά και συνεπώς η τιμή της παραμέτρου anaphoric connectives είναι 5.

Special words or phrases: Πρόκειται για συγκεκριμένες λέξεις, οι οποίες μπορούν να αυξήσουν την πολυπλοκότητα της πρότασης, λόγω της δυσκολίας που εισάγουν στον καθορισμό της σημασιολογίας της πρότασης. Μερικά χαρακτηριστικά παραδείγματα

τέτοιων λέξεων: “only”, “exactly one”, “unless”, “provided that”, “as long as”, “while”. Για παράδειγμα, ας θεωρήσουμε την ακόλουθη πρόταση :

“Some friends of Paul **only** like him **while** he is rich”

και την αντίστοιχη έκφραση ΚΛΠΤ :

“(∃x) friend(x,Paul) ∧ (rich(Paul) ⇒ likes(x,Paul))”.

Για την συγκεκριμένη πρόταση προκύπτει ότι η τιμή της παραμέτρου special words είναι 2.

Negating keywords: Αναφέρεται στις περιπτώσεις που η πρόταση περιέχει λέξεις με αρνητική σημασία όπως, “not”, “no”, “neither”, “none”, “except” και “never”, που δηλώνουν κάποιο είδος άρνησης. Όταν χρησιμοποιούνται τέτοιες λέξεις διαδοχικά μέσα σε μια πρόταση, πολλές φορές μπορεί να δημιουργήσει σύγχυση, δεδομένου ότι θα πρέπει να καθοριστεί αν η μια άρνηση αναιρεί την άλλη. Ως εκ τούτου, ο αριθμός των αρνητικών λέξεων-κλειδιών μπορεί να χρησιμοποιηθεί ως ένα χαρακτηριστικό για τον προσδιορισμό της δυσκολίας της μετατροπής. Για παράδειγμα, ας θεωρήσουμε την πρόταση:

“**Not** all cats will betray their master if he **doesn't** feed them”

με την αντίστοιχη έκφραση ΚΛΠΤ:

“(∃x) cat(x) ∧ ¬feeds(master-of(x),x) ∧ ¬betrays(x,master-of(x))”

Για την συγκεκριμένη πρόταση προκύπτει ότι η τιμή της παραμέτρου Negating keywords είναι 2.

Word order matching: Ένα κύριο χαρακτηριστικό πολυπλοκότητας στην διαδικασία μετατροπής σχετίζεται με τις απαιτούμενες αναδιατάξεις των λέξεων της αρχικής πρότασης προκειμένου να δημιουργηθεί η έκφραση ΚΛΠΤ. Για παράδειγμα, αν η πρόταση εκφράζει μια συνθήκη υπό όρους, συχνά συμβαίνει η σειρά της συνθήκης και του συμπεράσματος της πρότασης να αντιστρέφονται στην αντίστοιχη έκφραση ΚΛΠΤ. Το γεγονός αυτό δημιουργεί πολλές φορές σύγχυση κατά την διαδικασία της μετάφρασης. Τα βήματα για τον προσδιορισμό της παραμέτρου *Word order matching* είναι τα εξής:

1. Εξάγονται τα κατηγορήματα της έκφρασης ΚΛΠΤ, τα λήμματα (stems) αυτών και δημιουργείται μια λίστα (L_{FOL}), η οποία περιέχει τα κατηγορήματα με τη σειρά που εμφανίζονται στην πρόταση.
2. Γίνεται ανάλυση της πρότασης φυσικής γλώσσας και, λαμβάνοντας υπόψη μόνο λέξεις που έχουν λήμμα που περιλαμβάνεται στην πρώτη λίστα, δημιουργείται μια δεύτερη λίστα (L_{NL}), η οποία περιέχει τα κατηγορήματα με τη

σειρά που εμφανίζονται στην πρόταση φυσικής γλώσσας. Επίσης, αφαιρούνται οι διπλές εμφανίσεις ενός κατηγορήματος και λαμβάνεται υπόψη μόνο η πρώτη εμφάνιση του.

3. Γίνεται απαλοιφή των κατηγορημάτων από την λίστα L_{FOL} , τα οποία δεν βρέθηκαν στην πρόταση φυσικής γλώσσας, εξασφαλίζοντας ότι υπάρχει ο ίδιος αριθμός διακριτών στοιχείων και στις δύο λίστες.
4. Για κάθε λέξη στην L_{FOL} υπολογίζεται η διαφορά της από τη σειρά εμφάνισής της στην λίστα L_{NL} . Αν μια λέξη είναι στη θέση n_1 στη λίστα L_{FOL} και στη θέση n_2 στην L_{NL} , τότε η *διαφορά θέσης* (order difference) ορίζεται ως $|n_1 - n_2|$.
5. Ως *Word order matching* ορίζεται το άθροισμα των απόλυτων τιμών των διαφορών θέσης όλων των λέξεων:

$$Word_order_matching = \sum_{i=1}^n |pos_1(x_i) - pos_2(x_i)|$$

όπου n είναι ο αριθμός των λέξεων στις λίστες και $pos_1(x)$ και $pos_2(x)$ είναι συναρτήσεις που επιστρέφουν τις θέσεις n_1 και n_2 του στοιχείου x στις λίστες L_{FOL} και L_{NL} αντίστοιχα. Για παράδειγμα, ας θεωρήσουμε την πρόταση:

“Cats are **happy** as long as someone **feeds** them”

και την αντίστοιχη ΚΛΠΤ έκφραση:

$$“(\forall x) (\text{cat}(x) \wedge ((\exists y) \text{feeds}(y, x))) \Rightarrow \text{happy}(x)”$$

Αρχικά, εξάγονται τα κατηγορήματα της έκφρασης ΚΛΠΤ και εισάγονται με τη σειρά εμφάνισής τους στη λίστα $L_{FOL} = \{\text{cat}, \text{feed}, \text{happy}\}$ (βήμα 1). Στη συνέχεια γίνεται ανάλυση της πρότασης φυσικής γλώσσας και δημιουργείται η λίστα $L_{NL} = \{\text{cat}, \text{happy}, \text{feed}\}$ (βήμα 2). Δεν γίνεται καμία διαγραφή, αφού δεν υπάρχουν διαφορές (βήμα 3). Κατόπιν, υπολογίζονται οι ‘διαφορές θέσης’ των λέξεων και υπολογίζεται το $Word_order_matching = |1-1| + |2-3| + |3-2| = 2$ (βήματα 4 και 5).

Connective Mismatch: Σε μια πρόταση φυσικής γλώσσας ορισμένες λέξεις-κλειδιά υποδηλώνουν ότι θα πρέπει να χρησιμοποιηθούν συγκεκριμένα διασυνδετικά. Για παράδειγμα, προτάσεις που περιέχουν λέξεις όπως “and” ή “both” υποδηλώνουν ότι θα πρέπει να χρησιμοποιηθεί το λογικό διασυνδετικό “Λ”. Υπάρχουν όμως περιπτώσεις που η εμφάνιση τέτοιων λέξεων-κλειδιών όχι μόνο δεν συμβάλλει στον καθορισμό του διασυνδετικού που θα πρέπει να χρησιμοποιηθεί στην έκφραση ΚΛΠΤ, αλλά αποπροσανατολίζουν όταν απαιτείται να χρησιμοποιηθεί το αντίθετο διασυνδετικό. Πιο συγκεκριμένα, για κάθε πιθανό διασυνδετικό, δημιουργείται μια λίστα με τις αντίστοιχες λέξεις-κλειδιά που υπονοούν τη χρήση του. Εάν η πρόταση φυσικής γλώσσας περιέχει μία

από αυτές τις λέξεις, αλλά το αντίστοιχο συνδετικό λείπει στη μορφή ΚΛΠΤ, τότε η τιμή του συγκεκριμένου χαρακτηριστικού είναι *αληθής (true)*. Για παράδειγμα, ας θεωρήσουμε την πρόταση: “Apples and oranges are fruits”. Παρότι η ύπαρξη της λέξης “and” υπονοεί τη χρήση του διασυνδετικού “Λ”, ωστόσο λόγω της σημασιολογίας της πρότασης θα πρέπει να χρησιμοποιηθεί το διασυνδετικού “V”, και η σωστή έκφραση ΚΛΠΤ είναι: “ $(\forall x)(\text{apple}(x) \vee \text{orange}(x)) \Rightarrow \text{fruit}(x)$ ”.

Quantifier Mismatch: Όπως αναφέρθηκε και στις παραπάνω περιπτώσεις, ορισμένες λέξεις-κλειδιά σε μια πρόταση, μπορεί να συνδέονται με τη χρήση συγκεκριμένων ποσοδεικτών. Για παράδειγμα λέξεις όπως “something”, “some”, “a” υπονοούν την χρήση του υπαρξιακού ποσοδείκτη “Ε”. Ωστόσο, υπάρχουν περιπτώσεις όπου, παρά την ύπαρξη τέτοιων λέξεων, μπορεί να απαιτείται η χρήση του καθολικού ποσοδείκτη. Για παράδειγμα, θεωρούμε την πρόταση φυσικής γλώσσας: “A banana is both a fruit and a weapon”, και την αντίστοιχη έκφραση ΚΛΠΤ: “ $(\forall x)(\text{banana}(x) \Rightarrow (\text{fruit}(x) \wedge \text{weapon}(x)))$ ”. Παρότι η ύπαρξη της λέξης “a” υπονοεί τη χρήση του υπαρξιακού ποσοδείκτη, ωστόσο θα πρέπει να χρησιμοποιηθεί ο καθολικός ποσοδείκτης και η τιμή της παραμέτρου Quantifier Mismatch είναι true. Εναλλακτικά, για την πρόταση “All bananas are fruits and weapons” που έχει την ίδια έκφραση ΚΛΠΤ, η λέξη-κλειδί “all” υποδηλώνει τη χρήση του καθολικού ποσοδείκτη “V”, ο οποίος περιέχεται στη σωστή απάντηση.

Για τον τελικό προσδιορισμό του επιπέδου δυσκολίας μιας πρότασης-άσκησης, ο μηχανισμός βασίζεται και στις παραμέτρους της έκφρασης ΚΛΠΤ και στα χαρακτηριστικά της πρότασης φυσικής γλώσσας. Πιο συγκεκριμένα, για τον προσδιορισμό του τελικού επιπέδου δυσκολίας, λαμβάνεται υπόψη και η δυσκολία που προκύπτει από την σημασιολογία της πρότασης Φυσικής Γλώσσας. Για τον σκοπό αυτό αναπτύξαμε ένα έμπειρο σύστημα βασισμένο σε κανόνες.

Στον Πίνακα 3.2 παρουσιάζονται οι κανόνες για τον προσδιορισμό της δυσκολίας των προτάσεων με βάση την έκφραση ΚΛΠΤ και τα σημασιολογικά χαρακτηριστικά της πρότασης ΦΓ. Η πρώτη στήλη του πίνακα αναπαριστά την αρχική εκτίμηση του επιπέδου δυσκολίας της άσκησης από το Βασισμένο στην ΚΛΠΤ έμπειρο σύστημα, ενώ οι υπόλοιπες στήλες αφορούν τα χαρακτηριστικά που σχετίζονται με τη σημασιολογία της πρότασης ΦΓ. Στην συνέχεια, με βάση τα γεγονότα εισόδου, ενεργοποιείται ο κατάλληλος κανόνας και γίνεται επαναπροσδιορισμός του επιπέδου δυσκολίας, ενώ σε διαφορετική περίπτωση το τελικό επίπεδο δυσκολίας είναι το αρχικό. Στον Πίνακα 3.2, το σύμβολο “*” δηλώνει ότι η τιμή της συγκεκριμένης παραμέτρου δεν παίζει ρόλο, το SUM είναι το άθροισμα των τιμών

των τεσσάρων αριθμητικών χαρακτηριστικών και ο εκθέτης “d” δηλώνει ότι οι παράμετροι αυτοί συνδυάζονται διαζευκτικά μεταξύ τους, ενώ με τις υπόλοιπες συζευκτικά.

Πίνακας 3.2 Κανόνες για τον προσδιορισμό του επιπέδου δυσκολίας με βάση την έκφραση ΚΛΠΤ και τη σημασιολογία της Φυσικής Γλώσσας

FOL-Based Difficulty level	Word order matching	Anaphoric Connectives	Negating Keywords	Special words/ phrases	Quantifier Mismatch	Connective Mismatch	Difficulty Level (Output)
Difficult	*	>2	*	*	*	*	Advanced
Difficult	SUM >3				*	*	Advanced
Medium	*	>2	*	*	*	*	Difficult
Medium	*	*	>1	*	*	*	Difficult
Medium	SUM >3				*	*	Advanced
Easy	0	>1	*	*	yes	*	Medium
Very Easy	>2 ^d	>2 ^d	*	*	*	yes ^d	Easy

3.5 Μεθοδολογία Αυτόματης Μετατροπής Φυσικής Γλώσσας σε ΚΛΠΤ

Σε αυτή την ενότητα, παρουσιάζουμε μια προσέγγιση για την αυτόματη μετατροπή προτάσεων ΦΓ σε ΚΛΠΤ και ανάλυση της λειτουργικότητάς της. Η προσέγγιση έχει ως στόχο να μεταφράζει προτάσεις φυσικής γλώσσας σε ΚΛΠΤ και βασίζεται στην αρχή ότι προτάσεις φυσικής γλώσσας που έχουν ίδια ή παρόμοια συντακτική δομή και δέντρα εξαρτήσεων θα έχουν ίδια ή παρόμοια δομή έκφρασης σε ΚΛΠΤ. Στο πλαίσιο αυτό χρησιμοποιείται η *συλλογιστική βασισμένη σε περιπτώσεις (case based reasoning)*.

3.5.1 Συλλογιστική Βασισμένη σε Περιπτώσεις

Γενικά, η *συλλογιστική βασισμένη σε περιπτώσεις (case based reasoning)* βασίζεται στην αξιοποίηση παρόμοιων περιπτώσεων του παρελθόντος για την επίλυση νέων. Η μεθοδολογία της συλλογιστικής της βασισμένης σε περιπτώσεις είναι στενά συνδεδεμένη με τον τρόπο που οι άνθρωποι επιλύουν προβλήματα, και πιο συγκεκριμένα με την ανάσυρση παλαιότερων περιπτώσεων προβλημάτων και την επαναχρησιμοποίηση των λύσεων τους. Αυτή η προσέγγιση είναι εμπνευσμένη από μοντέλα της γνωστικής ψυχολογίας, σχετικά με το πώς οι άνθρωποι βγάζουν συμπεράσματα, προσδιορίζουν λύσεις νέων προβλημάτων αξιοποιώντας εμπειρίες και παλιότερες περιπτώσεις και λύσεις τους και έχει αποδειχθεί ότι

λειτουργεί πολύ καλά σε ένα ευρύ φάσμα εφαρμογών. Η συλλογιστική βασισμένη στις περιπτώσεις βασίζεται στην αρχή ότι παρόμοια προβλήματα έχουν παρόμοιες λύσεις. Πιο συγκεκριμένα, η εμπειρία από προηγούμενες επιλύσεις σε προβλήματα αποθηκεύεται ως περιπτώσεις και αποτελούν τη βάση γνώσης.

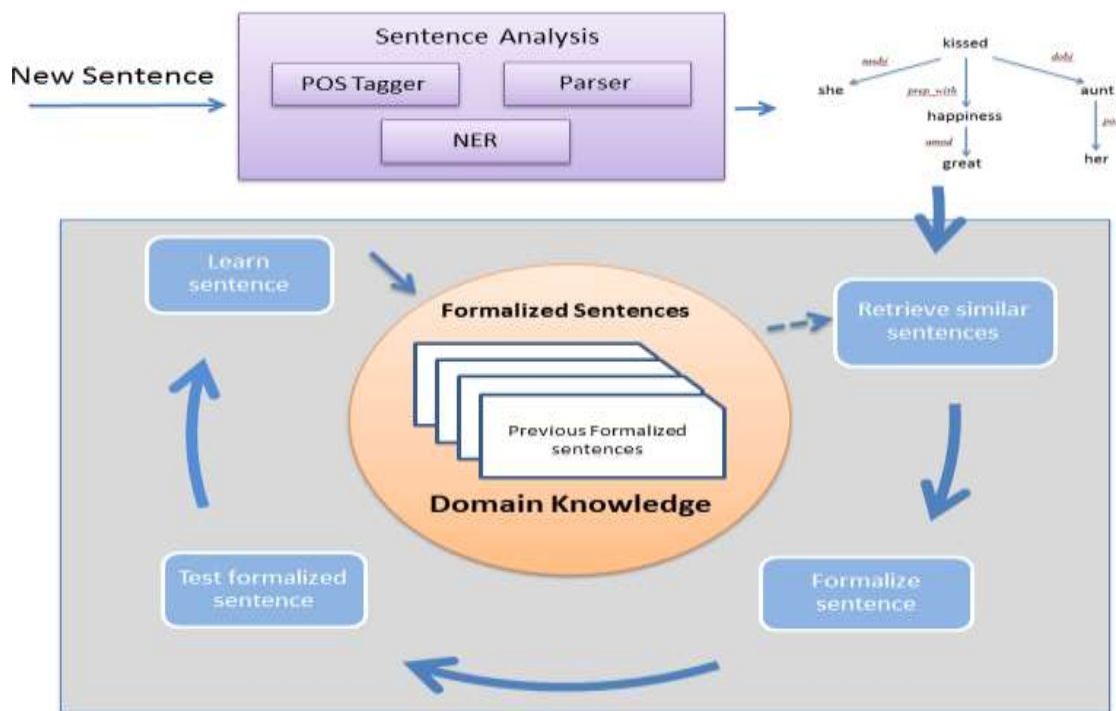
Σε ένα σύστημα βασισμένο στη συλλογιστική βασισμένη σε περιπτώσεις, μια βασική συνιστώσα είναι η *βάση περιπτώσεων* (*case base*), που αποτελεί τη βάση γνώσης του συστήματος. Τα νέα προβλήματα συγκρίνονται με προηγούμενες επιλυθείσες περιπτώσεις, που είναι αποθηκευμένες στη βάση, προκειμένου να προσδιοριστεί μια λύση για το νέο πρόβλημα και να γίνει λήψη αποφάσεων. Για κάθε νέο πρόβλημα, παρόμοιες περιπτώσεις ανακτώνται από τη βάση γνώσης και όταν βρεθεί η καλύτερη και πιο κατάλληλη περίπτωση, η λύση της χρησιμοποιείται είτε όπως είναι είτε προσαρμόζεται για να ταιριάζει στο νέο πρόβλημα. Η εύρεση της πιο κατάλληλης περίπτωσης στηρίζεται σε μετρικές ομοιότητας, που αποτελούν μια από τις σημαντικές ερευνητικές κατευθύνσεις του πεδίου.

Η μεθοδολογία της βασισμένης σε περιπτώσεις συλλογιστικής περιλαμβάνει τέσσερα στάδια. Το πρώτο στάδιο είναι το στάδιο της “ανάκτησης” (*retrieve*), όπου προσδιορίζονται από την βάση γνώσης περιπτώσεις, παρόμοιες με την τρέχουσα. Ο προσδιορισμός αυτός στηρίζεται, όπως προαναφέρθηκε, σε κάποια ή κάποιες μετρικές ομοιότητας. Στην συνέχεια, στο στάδιο της “επαναχρησιμοποίησης” (*reuse*), γίνεται επαναχρησιμοποίηση των λύσεων των περιπτώσεων αυτών, δηλαδή προσδιορίζεται η ζητούμενη λύση με βάση τις προηγούμενες. Στο επόμενο στάδιο της “αναθεώρησης” (*revise*) γίνεται αναθεώρηση της προταθείσας λύσης, αν απαιτείται, μετά από κάποια αξιολόγηση. Τέλος, στο τέταρτο στάδιο της μεθοδολογίας, το οποίο είναι το στάδιο της “διατήρησης” (*retain*), η παραχθείσα λύση, μετά από επιτυχή εφαρμογή ή αξιολόγηση, ενσωματώνεται στη βάση γνώσης ως νέα περίπτωση.

Ένα βασικό χαρακτηριστικό της μεθοδολογίας της συλλογιστικής με βάση τις περιπτώσεις είναι ότι αποτελεί μια προσέγγιση για αυξητική και συνεχή εκμάθηση, δεδομένου ότι μια νέα εμπειρία διατηρείται κάθε φορά που έχει λυθεί ένα νέο πρόβλημα, καθιστώντας την αμέσως διαθέσιμη για συνεισφορά σε επίλυση νέων μελλοντικών προβλημάτων. Όταν ένα νέο πρόβλημα λύνεται επιτυχώς, η εμπειρία διατηρείται προκειμένου να ληφθεί υπόψη σε παρόμοια προβλήματα στο μέλλον.

3.5.2 Αυτόματη Μετατροπή ΦΓ σε ΚΛΠΤ Βασισμένη σε Περιπτώσεις

Στην εικόνα 3.2 παρουσιάζεται ο γενικός κύκλος λειτουργίας της προσέγγισής μας για την αυτόματη μετατροπή ΦΓ σε ΚΛΠΤ.



Εικόνα 3.2 Κύκλος λειτουργίας μετατροπής ΦΓ σε ΚΛΠΤ

Στα πλαίσια αυτής της προσέγγισης, για τον φορμαλισμό μιας νέα πρότασης φυσικής γλώσσας, αρχικά πραγματοποιείται εις βάθος ανάλυση της πρότασης, όπου αναλύεται η γραμματική της δομή και σχηματίζεται το δέντρο εξαρτήσεων (dependency tree). Στην συνέχεια, ακολουθείται η μεθοδολογία της *συλλογιστικής βασισμένης σε περιπτώσεις* (case-based reasoning), προκειμένου να καθοριστούν παρόμοιες προτάσεις και να χρησιμοποιηθούν οι αντίστοιχες εκφράσεις τους σε ΚΛΠΤ για να καθοδηγήσουν τον φορμαλισμό της νέας πρότασης. Πιο συγκεκριμένα, ο φορμαλισμός μιας νέας πρότασης βασίζεται σε παρόμοιες προτάσεις φυσικής γλώσσας που έχουν ήδη μετατραπεί, και οι οποίες αποτελούν μαζί με τις αντίστοιχες εκφράσεις τους σε ΚΛΠΤ, τη βάση γνώσης του συστήματος.

Αρχικά, πραγματοποιείται συντακτική ανάλυση της νέας πρότασης με στόχο να καθοριστεί η μορφοσυντακτική δομή της και να εξαχθεί γνώση σχετικά με την σημασιολογία της. Η σημασιολογική ανάλυση περιλαμβάνει τον προσδιορισμό διαφόρων

τύπων λέξεων ή φράσεων, αναγνωρίζοντας αν μια λέξη είναι κύριο όνομα καθώς και τον συντακτικό ρόλο που έχει μέσα στην πρόταση, όπως για παράδειγμα αν είναι το υποκείμενο ή το αντικείμενο ενός ρήματος. Δοθείσας μια πρόταση, αρχικά ο Tree Tagger χρησιμοποιείται για να καθορίσει για κάθε λέξη τη βασική της μορφή (λήμμα) και να προσδιορίσει τον γραμματικό ρόλο της στην πρόταση. Στη συνέχεια, ο Stanford parser χρησιμοποιείται για να εκτελέσει μια εις βάθος ανάλυση της δομής της πρότασης, ώστε να προσδιοριστούν οι σχέσεις μεταξύ των λέξεων της πρότασης και να καθοριστούν οι αντίστοιχες εξαρτήσεις (dependencies). Το δέντρο εξαρτήσεων (dependency tree) που σχηματίζεται αντιπροσωπεύει τις γραμματικές σχέσεις μεταξύ των λέξεων της πρότασης (tree based approach). Οι σχέσεις αναπαρίστανται ως τριπλέτες, που κάθε μια αποτελείται από το όνομα της σχέσης, τον κυβερνήτη (governor), που πραγματοποιεί την σχέση, και τον εξαρτώμενο (dependent), που δέχεται την ενέργεια της σχέσης. Οι εξαρτήσεις αυτές δείχνουν τον τρόπο που οι λέξεις της πρότασης συνδέονται και αλληλεπιδρούν μεταξύ τους. Έπειτα, γίνεται προσδιορισμός των κυρίων ονομάτων, που ενδεχομένως υπάρχουν στην πρόταση. Ο προσδιορισμός των κυρίων ονομάτων πραγματοποιείται από το εργαλείο Named Entity Recognizer (NER), το οποίο αναγνωρίζει λέξεις σε μια πρόταση, οι οποίες είναι ονόματα πραγμάτων (π.χ. person, company κ.λπ.), και προσδιορίζει το είδος της οντότητας τους. Για παράδειγμα, καθορίζει για κάθε κύριο όνομα αν αυτό αναπαριστά ένα πρόσωπο (π.χ. John), αν αναπαριστά το όνομα μιας χώρας (π.χ. Hellas), ή μιας πόλης (π.χ. Athens). Τα κύρια ονόματα οντοτήτων της πρότασης θα αποτελέσουν τις σταθερές της αντίστοιχης έκφρασης ΚΛΠΤ και επομένως η ακριβής αναγνώρισή τους σε μια πρόταση είναι εξαιρετικής σημασίας για τον καθορισμό των ορισμάτων των κατηγορημάτων.

Μετά την ανάλυση της νέας πρότασης φυσικής γλώσσας και τον καθορισμό του δέντρου εξαρτήσεών της, ακολουθεί η προσέγγιση *συλλογιστικής βασισμένης σε κειμενικές περιπτώσεις* (textual case based reasoning), η οποία αποτελείται από τα τέσσερα στάδια που προαναφέρθηκαν. Το πρώτο στάδιο είναι το στάδιο της “ανάκτησης” (retrieve), όπου προσδιορίζονται παρόμοιες προτάσεις από την βάση γνώσης μέσω κάποιας μετρικής ομοιότητας. Στην συνέχεια, το στάδιο της “επαναχρησιμοποίησης” (reuse) έχει ως στόχο να γίνει επαναχρησιμοποίηση των λύσεων τους, δηλαδή των εκφράσεων ΚΛΠΤ που μεταφέρουν το νόημά τους και να καθοριστεί η προτεινόμενη έκφραση ΚΛΠΤ της νέας πρότασης. Έπειτα, το στάδιο της “αναθεώρησης” (revise) έχει ως στόχο να επιβεβαιώσει την συντακτική της ορθότητα και τέλος, στο τέταρτο στάδιο της μεθοδολογίας, το οποίο είναι το στάδιο της “διατήρησης” (retain), η παραγόμενη έκφραση ΚΛΠΤ παρουσιάζεται

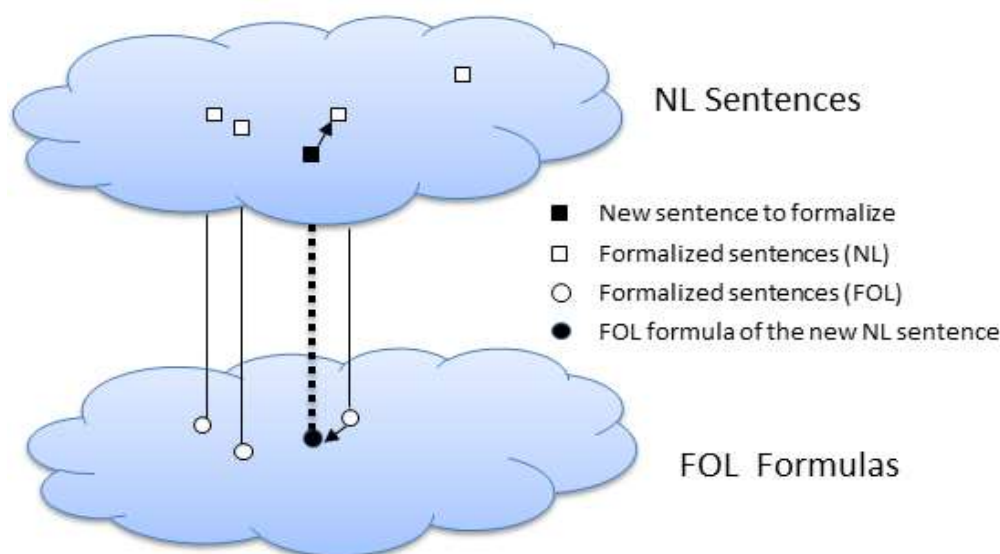
στον ειδικό-εμπειρογνώμονα για να ελέγξει την καταλληλότητα της και στη συνέχεια να ενσωματωθεί στη βάση γνώσης ως νέα περίπτωση.

Όταν οι περιπτώσεις αφορούν ελεύθερη μορφή κειμένου, η μέτρηση της ομοιότητας ανάμεσα σε ένα νέο πρόβλημα (πρόταση ΦΓ) και σε προηγούμενα προβλήματα (προτάσεις ΦΓ) που έχουν επιλυθεί αποτελεί μια δύσκολη υπόθεση. Η ομοιότητα των προτάσεων ΦΓ στην προσέγγισή μας προσδιορίζεται με βάση την Tree Edit Distance (Bille, 2005), μετρική η οποία υπολογίζεται στα δένδρα εξαρτήσεων των προτάσεων. Λόγω του κύριου ρόλου της ανάκτησης σε μεθοδολογίες συλλογιστικής βασισμένες στις περιπτώσεις, ένα σημαντικό μέρος της έρευνας έχει επικεντρωθεί στην ανάκτηση και τον υπολογισμό της ομοιότητας (similarity assessment) (De Mantaras et al., 2005).

Στο πλαίσιο της μελέτης μας, η tree edit distance χρησιμοποιείται για τη μέτρηση της ομοιότητας μεταξύ των δέντρων εξαρτήσεων μιας δεδομένης πρότασης και αυτών της βάσης γνώσης του συστήματος. Η tree edit distance έχει αποδειχθεί ότι είναι αποτελεσματική σε μια ποικιλία εφαρμογών, συμπεριλαμβανομένων των κειμενικών συνεπαγωγών, της αναγνώρισης παραφράσεων, της κατάταξης απαντήσεων και στην ανάκτηση πληροφοριών. Επιπλέον, είναι κατάλληλη κατά τη διαδικασία καθορισμού της ομοιότητας μεταξύ των εξαρτήσεων των προτάσεων. Η tree edit distance βασίζεται στην απόσταση μετασχηματισμού (edit distance) (Levenshtein, 1966) και αποτελεί ένα ενδεικτικό μέτρο της ομοιότητας μεταξύ δύο ή περισσότερων δεντρικών δομών. Σε γενικές γραμμές, η tree edit distance ανάμεσα σε δύο δένδρα T_1 και T_2 αντιπροσωπεύει το συνολικό κόστος της μετατροπής από το T_1 στο T_2 με τις βασικές λειτουργίες μετασχηματισμού σύμφωνα με τη συνάρτηση κόστους τους. Οι τρεις στοιχειώδεις πράξεις μετασχηματισμού είναι, η εισαγωγή (insert), η διαγραφή (delete) και η μετονομασία (relabeling). Σε ένα δέντρο εξαρτήσεων, κάθε κόμβος συνδέεται με τρία πεδία πληροφοριών, τα οποία είναι το λήμμα της λέξης, το μέρος του λόγου της και το είδος της σχέσης εξάρτησης με άλλες λέξεις. Ο αλγόριθμος υπολογισμού της tree edit distance έχει ως στόχο να βρει την καλύτερη και λιγότερο δαπανηρή ακολουθία στοιχειωδών πράξεων μετασχηματισμού που μετατρέπουν την δεντρική δομή T_1 στην T_2 . Η ομοιότητα μεταξύ των δύο προτάσεων φυσικής γλώσσας υπολογίζεται με τον ακόλουθο τρόπο. Υπολογίζεται το κόστος μετασχηματισμού από τον τύπο:

$$Cost(T_1, T_2) = \frac{(ted(T_1, T_2))}{(ted(_, T_2))}$$

όπου $ted(T1, T2)$ αναπαριστά την απόσταση tree edit distance μεταξύ των δέντρων $T1$ και $T2$ και όπου $ted(_, T2)$ αναπαριστά το κόστος εισαγωγής ολόκληρου του δέντρου $T2$. Στη μελέτη μας, δεδομένου των χαρακτηριστικών του δέντρου εξαρτήσεων της πρότασης, μια εισαγωγή κόμβου ή διαγραφή έχει οριστεί να είναι 3, δεδομένου ότι τρία κύρια πεδία των πληροφοριών έχουν εισαχθεί ή διαγραφεί. Όσο μικρότερο είναι το κόστος μετασχηματισμού τόσο μεγαλύτερη είναι η ομοιότητα μεταξύ των δέντρων των προτάσεων. Όπως παρουσιάζεται στην Εικόνα 3.3, ο στόχος είναι να επιλεγθεί μια παρόμοια πρόταση, με τους όρους των δέντρων εξαρτήσεων και την δομή της, και την αντίστοιχη έκφραση ΚΛΠΤ στη διαδικασία του φορμαλισμού της νέας πρότασης.



Εικόνα 3.3 Σχηματική αναπαράσταση του προσδιορισμού όμοιων προτάσεων

Στην συνέχεια, στο στάδιο της “επαναχρησιμοποίησης” (reuse), γίνονται κατάλληλες τροποποιήσεις με τη χρήση κανόνων, προκειμένου να καθοριστεί η έκφραση ΚΛΠΤ της νέας πρότασης. Οι κανόνες που σχεδιάστηκαν, απορρέουν από τους γενικούς περιορισμούς και τις αρχές του πεδίου της κατηγορηματικής λογικής πρώτης τάξης, καθώς και τις αρχές της διαδικασίας μετατροπής προτάσεων φυσικής γλώσσας σε ΚΛΠΤ και παρουσιάζονται στον Πίνακα 3.3.

Πίνακας 3.3 Κανόνες δημιουργίας εκφράσεων ΚΛΠΤ

	Κανόνας	Επεξήγηση
Κανόνες	If word is verb or word is noun or word	Κατηγορήματα είναι τα ρήματα, ουσιαστικά και επίθετα της πρότασης.

κατηγορημάτων	is adjective then word is predicate	
	If word is proper_noun then word is constant.	Τα κύρια ονόματα αποτελούν σταθερές στην έκφραση
	If word is predicate and word is noun then word is in_singular	Τα ουσιαστικά-κατηγορήματα είναι στον ενικό αριθμό.
	If word is predicate and word is verb then word is in_third_person	Τα ρήματα-κατηγορήματα είναι στο τρίτο ενικό πρόσωπο και σε ενεστώτα χρόνο.
	If word is predicate and word is verb then word is in_present	
	If word is predicate and word is verb then word is connect_predicates	Τα ρήματα συνδέουν μεταξύ τους τα κατηγορήματα της πρότασης.
Κανόνες Σχέσεων Διασυνδέσεων	If word is verb and word is predicate then arity is num_of_dependencies	Το πλήθος των ορισμάτων των ρημάτων καθορίζεται από τον αριθμό των εξαρτήσεων/σχέσεων του ρήματος με άλλες οντότητες της πρότασης.
	If word is predicate and word has dependency and dependency is neg then has_negation is yes	Αν ένα κατηγορήμα εμφανίζεται σε εξάρτηση τύπου “neg” τότε δέχεται άρνηση.
Κανόνες ποσοδεικτών	If word is predicate and word is noun and word has dependency and (dependency is det or dependency is predet) and dependency word is in_universal_word_list then quantifier is universal	Αν ένα κατηγορήμα αναπαριστά σχέση ουσιαστικού και εμφανίζεται σε εξάρτηση τύπου “det” ή “predet” και η εξάρτηση είναι με λέξη από την λίστα των «καθολικών» λέξεων, τότε η οντότητα ποσοτικοποιείται από καθολικό ποσοδείκτη

If word is predicate and word is noun and word has dependency and (dependency is det or dependency is predet) and dependency word is in_existential_word_list then quantifier is existential	Αν ένα κατηγορημα αναπαριστά σχέση ουσιαστικού και εμφανίζεται σε εξάρτηση τύπου “det” ή “predet” και η εξάρτηση είναι με λέξη από την λίστα των «υπαρξιακών» λέξεων, τότε η οντότητα ποσοτικοποιείται από υπαρξιακό ποσοδείκτη
--	---

Για παράδειγμα, οι κανόνες αυτοί καθορίζουν ότι τα ουσιαστικά αντιστοιχούν σε κατηγορήματα και ότι πρέπει να είναι σε ενικό αριθμό, ενώ αντίστοιχα τα ρήματα, που και αυτά αντιστοιχούν σε κατηγορήματα, πρέπει να είναι στο τρίτο πληθυντικό πρόσωπο. Επίσης, οι κανόνες αυτοί καθορίζουν και συγκεκριμένες εξαρτήσεις μεταξύ των λέξεων στο δέντρο εξαρτήσεων, που μπορούν να γίνουν αντιληπτές στην έκφραση ΚΛΠΤ. Για παράδειγμα, ένας κανόνας μπορεί να εξεταστεί και να εφαρμοστεί σε μια σχέση “det”, που υποδηλώνει μια σχέση ποσοτικοποίησης, προκειμένου να προσδιοριστεί ο κατάλληλος λογικός ποσοδείκτης (π.χ. υπαρξιακός, καθολικός) της μεταβλητής που αντιπροσωπεύει το αντίστοιχο ουσιαστικό (MacCartney & Manning, 2007). Οι λέξεις που ποσοτικοποιούν τις μεταβλητές παρουσιάζονται στον Πίνακα 3.4

Πίνακας 3.4 Λέξεις ποσοτικοποίησης μεταβλητών ατομικών εκφράσεων

Universal list	Existential list
all, every, any, anyone, everyone, everything,	somebody, someone, some, something

Στην συνέχεια, το στάδιο της “αναθεώρησης” (revise), αναλύεται η παραγόμενη έκφραση ΚΛΠΤ με στόχο να επιβεβαιωθεί η συντακτική της ορθότητα. Η ανάλυση της δομής της έκφρασης ΚΛΠΤ πραγματοποιείται από το εργαλείο FOL analyzer (Perikos et al., 2011a).

Τέλος, στο στάδιο της “διατήρησης” (retain), η παραγόμενη έκφραση ΚΛΠΤ παρουσιάζεται στον ειδικό-εμπειρογνώμονα για να ελέγξει την καταλληλότητα της και στη συνέχεια να ενσωματωθεί στη βάση γνώσης ως νέα περίπτωση (case).

Για παράδειγμα, ας θεωρήσουμε την ακόλουθη πρόταση φυσικής γλώσσας προς μετατροπή σε ΚΛΠΤ:

“Gardeners love all flowers”.

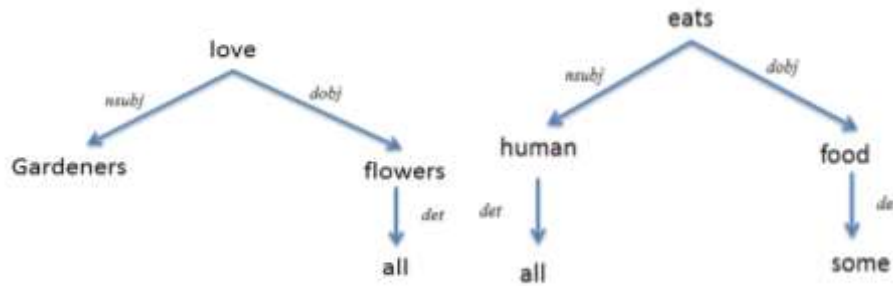
Επίσης, ας θεωρήσουμε ότι η βάση γνώσης περιέχει μεταξύ άλλων και την ακόλουθη πρόταση φυσικής γλώσσας:

“All humans eat some food”

μαζί με την αντίστοιχη ΚΛΠΤ έκφρασή της :

$"(\forall x)(\exists y)human(x) \wedge food(y) \Rightarrow eats(x,y)"$

Τα δέντρα εξαρτήσεων για τις δύο προτάσεις παρουσιάζονται στην Εικόνα 3.4.



Εικόνα 3.4 Τα δέντρα Εξαρτήσεων των προτάσεων

Η tree edit distance μεταξύ δυο προτάσεων υπολογίζεται ως: $ted(T_1, T_2) = 3$, δεδομένου ότι η εισαγωγή που απαιτείται έχει κόστος 3. Η tree edit distance μεταξύ της T_2 και της κενής είναι 15 και επομένως το κόστος μετασχηματισμού είναι:

$$Cost(T_1, T_2) = \frac{(ted(T_1, T_2))}{(ted(_, T_2))} = \frac{3}{15}$$

Η μικρή τιμή κόστους υποδηλώνει μεγάλη ομοιότητα των δύο περιπτώσεων. Υποθέτουμε ότι, μεταξύ όλων των προτάσεων της βάσης γνώσης, η νέα πρόταση έχει μεγαλύτερη ομοιότητα με την παραπάνω πρόταση και επομένως η αντίστοιχη έκφραση ΚΛΠΤ ανακτάται κατά το στάδιο της ανάκτησης για να καθοδηγήσει την διαδικασία του φορμαλισμού της νέας πρότασης.

Έπειτα, στο στάδιο της επαναχρησιμοποίησης, γίνεται προσαρμογή της έκφρασης ΚΛΠΤ στα χαρακτηριστικά την νέας πρότασης. Γενικοί κανόνες, οι οποίοι αφορούν τις βασικές αρχές της διαδικασίας μετατροπής και του πεδίου της λογικής, ορίζουν ότι τα ουσιαστικά αποτελούν τα κατηγορήματα και πρέπει να είναι στον ενικό αριθμό και αντίστοιχα ότι τα ρήματα πρέπει να είναι στο τρίτο πρόσωπο. Επομένως, η ενδιάμεση ΚΛΠΤ έκφραση που προσδιορίζεται είναι η εξής:

$$(\forall x)(\exists y)gardner(x) \wedge flower(y) \Rightarrow loves(x, y)$$

Στη συνέχεια, γίνεται ανάλυση του δέντρου εξαρτήσεων της νέας πρότασης, προκειμένου να γίνει προσαρμογή της έκφρασης στις ακριβείς σχέσεις εξαρτήσεων της νέας πρότασης. Γενικοί κανόνες ορίζουν τις αρχές αντιστοίχισης των στοιχείων των εξαρτήσεων στην έκφραση ΚΛΠΤ και με βάση την ανάλυση κάθε εξάρτησης καθορίζονται πιθανές αλλαγές. Πιο συγκεκριμένα, η εξάρτηση *det (flowers, all)* αναλύεται και λόγω της σχέσης τύπου “det” της λέξης “flowers” με την λέξη “all”, ο κανόνας συσχετίζει τον καθολικό ποσοδεικτή “ $\forall$ ”, με την μεταβλητή που αναπαριστά την λέξη “flowers”. Επομένως, μετά την τροποποίηση, η τελική έκφραση που δίνεται ως έξοδος είναι η εξής:

$$(\forall x)(\forall y)gardner(x) \wedge flower(y) \Rightarrow loves(x, y)$$

η οποία είναι η σωστή έκφραση ΚΛΠΤ. Στην συνέχεια, η έκφραση ΚΛΠΤ αξιολογείται συντακτικά και κατόπιν παρουσιάζεται στον εμπειρογνώμονα για να προσδιορίσει την καταλληλότητά της από άποψη σημασιολογίας. Αν η αξιολόγηση είναι θετική αποθηκεύεται στη βάση γνώσης ως νέα περίπτωση.

3.6 Πειραματικά Αποτελέσματα

Για την αξιολόγηση της απόδοσης της μεθοδολογίας του φορμαλισμού προτάσεων φυσικής γλώσσας καθώς και του προσδιορισμού του επιπέδου δυσκολίας του φορμαλισμού τους πραγματοποιήθηκε μια πειραματική μελέτη η οποία είχε δύο κύρια μέρη. Αρχικά, για τις ανάγκες της μελέτης, έγινε συλλογή διάφορων ειδών προτάσεων φυσικής γλώσσας και δημιουργήθηκε ένα σύνολο δεδομένων που περιείχε τις προτάσεις ΦΓ καθώς και τις αντίστοιχες εκφράσεις ΚΛΠΤ για κάθε μια πρόταση. Στα πλαίσια του πρώτου μέρους, έγινε μελέτη και αξιολόγηση της απόδοσης της προσέγγισης για τον προσδιορισμό του επιπέδου δυσκολίας των προτάσεων φυσικής γλώσσας. Στην συνέχεια, στο δεύτερο μέρος της πειραματικής μελέτης, αξιολογήθηκε η απόδοσή της μεθοδολογίας μετατροπής προτάσεων φυσικής γλώσσας σε κατηγορηματική λογική πρώτης τάξης.

3.6.1 Δημιουργία Συνόλου Δεδομένων

Το σύνολο δεδομένων που δημιουργήθηκε περιείχε συνολικά 160 προτάσεις. Οι προτάσεις συλλέχτηκαν από βιβλία λογικής καθώς και από τίτλους ειδήσεων. Οι τίτλοι των ειδήσεων χρησιμοποιήθηκαν κυρίως επειδή είναι μικροί και εκφραστικοί και μεταφέρουν μεγάλη ποσότητα πληροφορίας και επιλέχθηκαν από διάφορες ειδησεογραφικές ιστοσελίδες όπως το Euronews, το BBC και το CNN. Επίσης, οι προτάσεις επιλέχθηκαν ώστε να μην περιέχουν πολλές ασάφειες και αναφορικές εκφράσεις και να μπορούν να εκφραστούν με ακρίβεια σε κατηγορηματική λογική πρώτης τάξης.

3.6.2 Πειραματική Αξιολόγηση Προσδιορισμού Δυσκολίας

Για να αξιολογήσουμε την μεθοδολογία του μηχανισμού προσδιορισμού δυσκολίας της διαδικασίας μετατροπής προτάσεων που αναπτύχθηκε, πραγματοποιήσαμε μια πειραματική μελέτη χρησιμοποιώντας το σύνολο δεδομένων. Στα πλαίσια της μελέτης, έγινε εκτίμηση του επίπεδου δυσκολίας όλων των προτάσεων από τον εμπειρογνώμονα και στη συνέχεια δόθηκε το σύνολο στο σύστημα για να προσδιοριστεί αυτόματα το επίπεδο δυσκολίας της κάθε πρότασης. Για την αξιολόγηση του μηχανισμού υπολογίστηκαν οι μετρικές μέση ορθότητα (*average accuracy*), ακρίβεια (*precision*), ανάκληση (*recall*) και *F-measure* (Sokolova & Lapalme, 2009). Στον Πίνακα 3.5 παρουσιάζονται τα πειραματικά αποτελέσματα που προέκυψαν.

Πίνακας 3.5 Πειραματικά αποτελέσματα μηχανισμού εκτίμησης δυσκολίας προ

Average accuracy	0.94
Precision	0.88
Recall	0.89
F-measure	0.88

Τα αποτελέσματα δείχνουν μια πολύ καλή απόδοση του μηχανισμού προσδιορισμού της δυσκολίας των ασκήσεων. Πιο συγκεκριμένα, ο μηχανισμός στο 94% των περιπτώσεων προσδιόρισε σωστά το επίπεδο δυσκολίας των ασκήσεων κατατάσσοντάς αυτές στην κατάλληλη κατηγορία. Στην συνέχεια, υπολογίσαμε το στατιστικό δείκτη Cohen's Kappa (Cohen, 1960), για την αποτίμηση της απόδοσης του συστήματος και για τον προσδιορισμό της αξιοπιστίας του βαθμού συμφωνίας μεταξύ του διδάσκοντα-εμπειρογνώμονα και του αυτόματου μηχανισμού προσδιορισμού της δυσκολίας. Τα αποτελέσματα της μελέτης έδειξαν ότι υπήρξε μια σχεδόν τέλεια συμφωνία μεταξύ εκπαιδευτών και συστήματος,

δεδομένου ότι η μετρική ήταν $k = 0.85$ (confidence interval 95%, $p < 0.0005$) (Viera & Garret, 2005). Τα αποτελέσματα επιβεβαιώνουν την πολύ καλή απόδοση του συστήματος. Επιπλέον, τα χαρακτηριστικά που εξάγονται από την ανάλυση της δομής και της γραμματικής σύνταξη της φυσικής γλώσσας συμβάλουν σημαντικά στην εκτίμηση της πολυπλοκότητας φορμαλισμού των προτάσεων. Από τα αποτελέσματα φαίνεται ότι η προσέγγιση που ακολουθήθηκε και το σύστημα που αναπτύχθηκε προσδιορίζουν με ακρίβεια και με συνέπεια το επίπεδο δυσκολίας των προτάσεων ΦΓ.

3.6.3 Πειραματική Αξιολόγηση του Μηχανισμού Αυτόματης Μετατροπής ΦΓ σε ΚΛΠΤ

Στα πλαίσια του δεύτερου μέρους της πειραματικής μελέτης, έγινε μελέτη της μεθοδολογίας φορμαλισμού προτάσεων φυσικής γλώσσας και αξιολογήθηκε η απόδοσή της στην μετατροπή προτάσεων σε ΚΛΠΤ. Οι προτάσεις του συνόλου δεδομένων μαζί με τις αντίστοιχες εκφράσεις ΚΛΠΤ αποτέλεσαν την βάση γνώσης της μεθοδολογίας, όπου κάθε μια πρόταση ΦΓ μαζί με την αντίστοιχη έκφραση ΚΛΠΤ αποτελούν μια περίπτωση (case).

Για την αξιολόγηση της μεθοδολογίας, έγινε επαναληπτικά επιλογή κάθε μιας πρότασης από το σύνολο δεδομένων, απομάκρυνσή της ως περίπτωση από την βάση γνώσης και προσπάθεια φορμαλισμού της με βάση τις υπόλοιπες περιπτώσεις της βάσης γνώσης. Στην συνέχεια, με βάση τον υπολογισμό της ομοιότητας της πρότασης με τις προτάσεις της βάσης γνώσης γίνεται η επιλογή της πρότασης με την μεγαλύτερη ομοιότητα, μαζί με την αντίστοιχη έκφραση ΚΛΠΤ για να καθοδηγήσει την διαδικασία φορμαλισμού της νέας πρότασης. Το αποτέλεσμα του κύκλου συλλογιστικής βασισμένης στις περιπτώσεις που ακολουθείται, είναι μια προτεινόμενη έκφραση ΚΛΠΤ για τη νέα πρόταση. Η έκφραση ΚΛΠΤ που καθορίζεται ως αποτέλεσμα, παρουσιάζεται στον εμπειρογνώμονα για να την αξιολογήσει και να καθορίσει την ορθότητά της σε σημασιολογικό επίπεδο.

Δεδομένου ότι ο πιο κρίσιμος παράγοντας στον φορμαλισμό μιας νέας πρότασης είναι ο προσδιορισμός όμοιων προτάσεων, μελετήθηκε η απόδοσης της μεθοδολογίας με βάση την ομοιότητα μεταξύ της πρότασης προς μετατροπή σε ΚΛΠΤ και των προτάσεων της βάσης γνώσης. Η μετρική ομοιότητας κανονικοποιήθηκε να παίρνει τιμές στο διάστημα από 0 μέχρι 1, όπου η τιμή 0 δείχνει ότι τα δέντρα εξαρτήσεων των δυο προτάσεων δεν μοιάζουν καθόλου ενώ η τιμή 1 δείχνει μια πλήρη ομοιότητα μεταξύ των δέντρων εξαρτήσεων των προτάσεων. Στα πλαίσια της πειραματικής μελέτης θεωρήσαμε τέσσερα επίπεδα ομοιότητας μεταξύ των προτάσεων (Πίνακας 3.6).

Πίνακας 3.6 Επίπεδα ομοιότητας μεταξύ προτάσεων φυσικής γλώσσας

Μετρική Similarity	Επίπεδο ομοιότητας
0 - 0.40	Ανεπαίσθητη
0.40 - 0.60	Χαμηλή
0.60 – 0.80	Μέτρια
0.80 – 1	Μεγάλη

και έγινε μελέτη της αποτελεσματικότητας της μεθοδολογίας ως προς την ομοιότητα μεταξύ των προτάσεων καθώς και ως προς το επίπεδο δυσκολίας φορμαλισμού τους. Πιο συγκεκριμένα, προσδιορίστηκε, με βάση το επίπεδο της μετρικής ομοιότητας, το ποσοστό των προτάσεων που μετατράπηκαν με επιτυχία σε ΚΛΠΤ. Τα αποτελέσματα παρουσιάζονται στον Πίνακα 3.7.

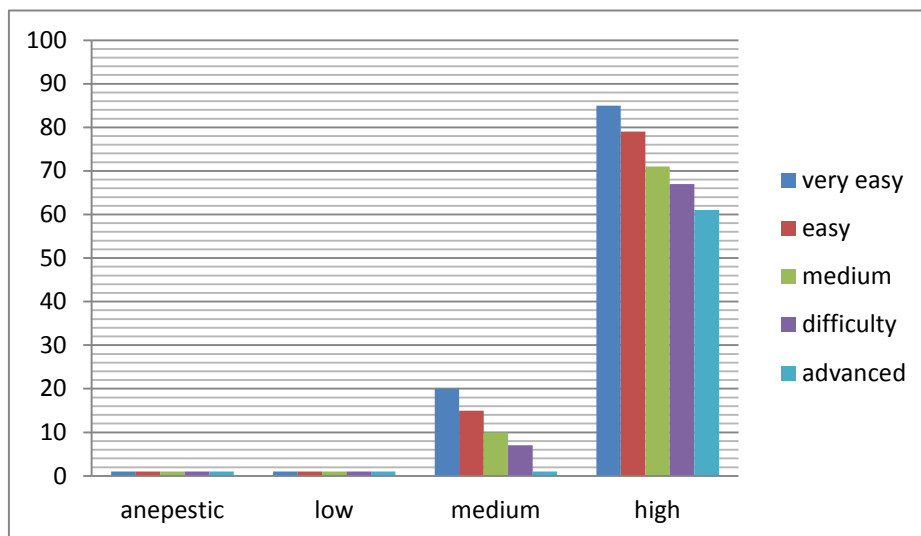
Πίνακας 3.7 Απόδοση με βάση την μετρική ομοιότητας

Επίπεδο ομοιότητας	Ποσοστό ορθών φορμαλισμών
Ανεπαίσθητη	0 %
Χαμηλή	0 %
Μέτρια	11 %
Μεγάλη	73 %

Από τα αποτελέσματα φαίνεται η μεγάλη σημασία που έχει η ομοιότητα στην αποτελεσματικότητα της μεθοδολογίας. Τα αποτελέσματα της μελέτης δείχνουν ότι η μεθοδολογία μετέτρεψε σωστά σε ΚΛΠΤ το 73% νέων προτάσεων, οι οποίες υπολογίστηκε ότι είχαν μεγάλη ομοιότητα με πρόταση (ή προτάσεις) της βάσης γνώσης, και το 11% των προτάσεων, που υπολογίστηκε ότι είχαν μέτρια ομοιότητα. Σε περιπτώσεις νέων προτάσεων προς φορμαλισμό, όπου η ομοιότητα τους με προτάσεις της βάσης γνώσης ήταν χαμηλή ή ανεπαίσθητη, δεν ήταν δυνατή η μετατροπή τους σε ΚΛΠΤ. Δεδομένου ότι η μεθοδολογία σχεδιάστηκε και βασίζεται στην αρχή ότι παρόμοιες προτάσεις ΦΓ θα έχουν παρόμοιες εκφράσεις ΚΛΠΤ, τα χαμηλά αποτελέσματα σε ανόμοιες προτάσεις ήταν αναμενόμενα. Όσον αφορά νέες προτάσεις, όπου υπολογίστηκε ότι είχαν μεγάλη ομοιότητα με πρόταση

από την βάση γνώσης, η μεθοδολογία παρουσιάζει μια πολύ ικανοποιητική απόδοση και μάλιστα, ακόμη και στην περίπτωση που δεν προσδιορίστηκε με πλήρη ακρίβεια η ορθή ΚΛΠΤ έκφραση, η παραγόμενη έκφραση ΚΛΠΤ ήταν πολύ κοντά στην ορθή.

Στην συνέχεια, έγινε περαιτέρω ανάλυση των αποτελεσμάτων της μεθοδολογίας με βάση και το επίπεδο δυσκολίας των προτάσεων. Στην Εικόνα 3.5, παρουσιάζονται τα αποτελέσματα της απόδοσης ως προς το επίπεδο δυσκολίας των προτάσεων προς μετατροπή σε ΚΛΠΤ.



Εικόνα 3.5 Η απόδοση ως προς το επίπεδο δυσκολίας των προτάσεων.

Από τα αποτελέσματα φαίνεται ο σημαντικός ρόλος που διαδραματίζει η ομοιότητα στον φορμαλισμό μιας νέας πρότασης και τον μεγάλο βαθμό στον οποίο επηρεάζει την απόδοση της μεθοδολογίας. Η απόδοση είναι πολύ καλή αναλογικά σε όλα τα επίπεδα δυσκολίας σε περιπτώσεις όπου η ομοιότητα μεταξύ της νέας πρότασης προς φορμαλισμό και προτάσεων της βάσης γνώσης είναι μεγάλη. Η μεθοδολογία παρουσιάζει την καλύτερή της απόδοση σε προτάσεις που χαρακτηρίστηκαν ως πολύ εύκολες και η ομοιότητα τους με πρόταση από την βάση γνώσης ήταν μεγάλη, όπου προσδιορίστηκαν σωστά οι εκφράσεις ΚΛΠΤ στο 85% των περιπτώσεων. Η απόδοση μειώνεται καθώς αυξάνεται το επίπεδο δυσκολίας των προτάσεων και μάλιστα σε προτάσεις που χαρακτηρίστηκαν ως προχωρημένης δυσκολίας η μεθοδολογία παρουσιάζει την χαμηλότερη απόδοσή της. Το γεγονός αυτό οφείλεται κυρίως την φύση των προτάσεων του προχωρημένου επιπέδου δυσκολίας, όπως το μεγάλο μήκος και οι πολλές έννοιες και πληροφορίες που μεταφέρουν. Σε πολλές περιπτώσεις, αν και η μεγάλη ομοιότητα μεταξύ της νέας πρότασης και της πρότασης στη βάση γνώσης βοηθάει τη μεθοδολογία στον προσδιορισμό της έκφρασης ΚΛΠΤ, εντούτοις μια ασυμφωνία μπορεί να επηρεάζει και να απαιτεί αλλαγές σε πολλά

μέρη και έννοιες της προτάσεις. Σε περιπτώσεις προτάσεων, που η ομοιότητα δεν είναι μεγάλη, η απόδοση της μεθοδολογίας μειώνεται αισθητά. Γενικά, από τα αποτελέσματα φαίνεται ότι στις περιπτώσεις όπου η ομοιότητα είναι μεγάλη η μεθοδολογία είναι ακριβής και μπορεί να αποτελέσει ένα αποτελεσματικό εργαλείο στην ανάλυση και κατανόηση φυσικής γλώσσας και την μετατροπή προτάσεων σε κατηγορηματική λογικής πρώτης τάξης.

Επιπροσθέτως, σχεδιάστηκε και διεξήχθη μια ακόμη πειραματική μελέτη για να γίνει πιο αναλυτική αποτίμηση της απόδοσης της μεθοδολογίας. Στα πλαίσια της μελέτης αυτής, το σύνολο των εκφράσεων ΚΛΠΤ, που προσδιορίστηκαν αυτόματα από το σύστημα για τις προτάσεις φυσικής γλώσσας του συνόλου δεδομένων, αξιολογήθηκαν και βαθμολογήθηκαν από δύο εμπειρογνώμονες ως προς την ορθότητά τους σε συντακτικό και σημασιολογικό επίπεδο. Πιο συγκεκριμένα, προκειμένου να γίνει με ακρίβεια η βαθμολόγηση των παραγόμενων εκφράσεων ΚΛΠΤ, οι εμπειρογνώμονες αξιολόγησαν τα δομικά μέρη των εκφράσεων ΚΛΠΤ και καθόρισαν τις ακόλουθες τρεις βαθμολογίες:

- βαθμολογία για την δημιουργία των ατομικών εκφράσεων
- βαθμολογία για τον προσδιορισμό των διασυνδεδετικών μεταξύ των ατομικών εκφράσεων
- βαθμολογία για τον προσδιορισμό των ποσοδεικτών.

Η συνολική βαθμολογία μιας έκφρασης ΚΛΠΤ, που προσδιορίζει αυτόματα η μεθοδολογία, προκύπτει από την σύνθεση των επιμέρους τριών βαθμολογιών. Ανάλογα με τα χαρακτηριστικά της κάθε έκφρασης ΚΛΠΤ, προκύπτουν διαφορετικές συνθέσεις των επί μέρους βαθμολογιών. Για παράδειγμα, ο καθορισμός των ατομικών εκφράσεων είναι απαραίτητο να πραγματοποιείται πάντα, ενώ ο καθορισμός ποσοδεικτών ή διασυνδεδετικών ενδέχεται να μην είναι πάντα αναγκαίος. Τα διάφορα σχήματα βαθμολόγησης, ανάλογα με τα χαρακτηριστικά των εκφράσεων ΚΛΠΤ, παρουσιάζονται στον Πίνακα 3.8, όπου προσδιορίζονται τέσσερα διαφορετικά σχήματα σύνθεσης για τον υπολογισμό της συνολικής βαθμολογίας αξιολόγησης.

Πίνακας 3.8 Σχήματα Βαθμολόγησης Ασκήσεων Μετατροπής ΦΓ σε ΚΛΠΤ

Σχήμα	Ατομικές εκφράσεις	Διασυνδεδετικά	Ποσοδείκτες	Σύνθεση Βαθμολογίας	Παράδειγμα ΚΛΠΤ έκφρασης
SS1	Yes	Yes	Yes	35-35-30	$(\forall x) (\text{dog}(x) \wedge \text{loves}(x, \text{master-of}(x))) \Rightarrow \text{loves}(\text{maria}, x)$

SS2	Yes	Yes	No	55-45-0	$\text{dog(pluto)} \wedge \text{dog(jack)}$
SS3	Yes	No	Yes	65-0-35	$(\exists x) \text{hope}(x)$
SS4	Yes	No	No	100-0-0	man(socrates)

Οι εκφράσεις αξιολογήθηκαν από τους εμπειρογνώμονες και καθορίστηκε η κάθε επιμέρους βαθμολογία τους στην κλίμακα από 0 έως 10. Στην συνέχεια, έγινε σύνθεσή των επιμέρους βαθμολογιών και προσδιορίστηκε η συνολική βαθμολογία με βάση το κατάλληλο σχήμα σύνθεσης. Τα αποτελέσματα της συνολικής βαθμολόγησης των εκφράσεων ΚΛΠΤ παρουσιάζονται στον Πίνακα 3.9.

Πίνακας 3.9 Συνολική βαθμολογία εκφράσεων ΚΛΠΤ

	Mean	SD
Expert 1	7.25	1.21
Expert 2	7.53	1.10

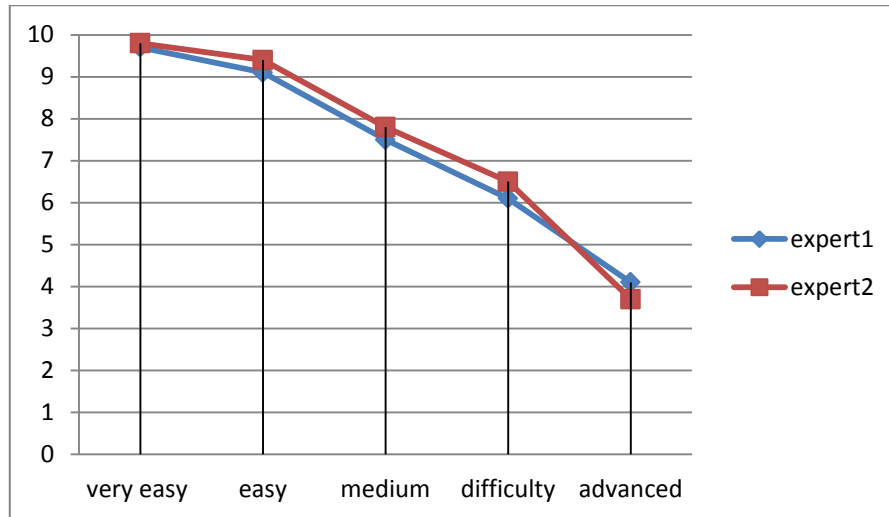
Τα αποτελέσματα δείχνουν μια καλή απόδοση του συστήματος, δεδομένου ότι ακόμη και σε περιπτώσεις που το σύστημα δεν κατάφερε να προσδιορίσει την σωστή έκφραση ΚΛΠΤ μιας πρότασης, η παραγόμενη ΚΛΠΤ βαθμολογήθηκε υψηλά και ήταν κοντά στην σωστή. Στην συνέχεια, η ανάλυση των αξιολογήσεων των εμπειρογνομόνων για κάθε δομικό μέρος των παραγόμενων εκφράσεων ΚΛΠΤ και οι βαθμολογίες τους παρουσιάζονται στον Πίνακα 3.10.

Πίνακας 3.10 Συνολική βαθμολογία Εκφράσεων ΚΛΠΤ

	Ατομικές εκφράσεις		Διασυνδετικά		Ποσοδείκτες	
	Mean	SD	Mean	SD	Mean	SD
Expert 1	9.03	0.66	5.82	1.32	7.17	1.24
Expert 2	8.81	0.65	6.21	1.25	7.48	0.95

Από τα αποτελέσματα προκύπτει ότι το σύστημα δημιούργησε εκφράσεις ΚΛΠΤ που ήταν ορθές σε συντακτικό και σημασιολογικό επίπεδο. Τα αποτελέσματα της αξιολόγησης ήταν πολύ ικανοποιητικά ($M=7.25$ $SD=1.21$, $M=7.53$ $SD=1.10$). Η αξιολόγηση των εμπειρογνομόνων ήταν πολύ υψηλή για τον καθορισμό των ατομικών εκφράσεων των εκφράσεων ΚΛΠΤ από την προσέγγιση, ενώ ήταν χαμηλότερη για τον καθορισμό των διασυνδετικών. Από τα αποτελέσματα προκύπτει ότι στις παραγόμενες εκφράσεις ΚΛΠΤ, στις περισσότερες περιπτώσεις, η μεθοδολογία προσδιορίζει και δημιουργεί σωστά τις κατάλληλες ατομικές εκφράσεις. Το γεγονός αυτό οφείλεται στην χρησιμοποίηση της δομής των ατομικών εκφράσεων παρόμοιας έκφρασης ΚΛΠΤ, που προσδιορίζεται στο βήμα της ανάκτησης, καθώς και στους κανόνες που εφαρμόζονται κατά το στάδιο της επαναχρησιμοποίησης. Από την άλλη πλευρά, η χαμηλότερη βαθμολογία των εμπειρογνομόνων στην κατηγορία των διασυνδετικών των παραγόμενων εκφράσεων ΚΛΠΤ οφείλεται στην πολυπλοκότητα της διασύνδεσης των ατομικών εκφράσεων στην τελική έκφραση ΚΛΠΤ, καθώς σε ορισμένες περιπτώσεις, μικρές αλλαγές μεταξύ της νέας πρότασης και της παρόμοιας που ανακτάται, μπορεί να απαιτούν μεγάλες αλλαγές στην διασύνδεση των ατομικών εκφράσεων των αντίστοιχων εκφράσεων ΚΛΠΤ. Επίσης, ο σχηματισμός των κατάλληλων ομάδων ατόμων απαιτεί εις βάθος κατανόηση της σημασιολογικής πληροφορίας της πρότασης καθώς και κοινότυπη γνώση, κάτι που είναι εξαιρετικά δύσκολο για κάθε αυτόματη διεργασία επεξεργασίας και κατανόησης της φυσικής γλώσσας.

Επιπροσθέτως, έγινε ανάλυση της αξιολόγησης των εμπειρογνομόνων όσον αφορά την απόδοση της μεθοδολογίας ανάλογα με το επίπεδο της δυσκολίας φορμαλισμού των προτάσεων. Τα αποτελέσματα που προέκυψαν παρουσιάζονται στην Εικόνα 3.6.



Εικόνα 3.6 Η αξιολόγηση ως προς το επίπεδο δυσκολίας των προτάσεων.

Από τα αποτελέσματα προκύπτει ότι, οι εκφράσεις που έλαβαν χαμηλότερη βαθμολόγηση ήταν κυρίως προτάσεις με επίπεδο δυσκολίας φορμαλισμού “difficulty” και “advanced”, ενώ η αξιολόγηση των εμπειρογνομόνων ήταν πολύ υψηλή για προτάσεις που ήταν επιπέδου “very easy” και “easy”. Επίσης, φαίνεται να επαληθεύεται και η απόδοση του συστήματος προσδιορισμού του επιπέδου δυσκολίας των προτάσεων, όπου σε προτάσεις δύσκολου και προχωρημένου επιπέδου, σε πολλές περιπτώσεις, δεν έγινε ακριβής προσδιορισμός της έκφρασης ΚΛΠΤ και πραγματοποιήθηκαν σφάλματα κυρίως σε σημασιολογικό επίπεδο.

4

Ανάλυση Κειμένου και Δημιουργία Παραφράσεων

4.1 Εισαγωγή

Η τεράστια ποσότητα κειμενικής πληροφορίας στο διαδίκτυο καθιστά αναγκαίο το σχεδιασμό μεθόδων για την αυτόματη ανάλυση και την εξαγωγή γνώσης από αυτό. Τα τελευταία χρόνια, η ανάλυση προτάσεων φυσικής γλώσσας και η δημιουργία σημασιολογικά ισοδύναμων παραφράσεων έχει προσελκύσει το ενδιαφέρον της ερευνητικής κοινότητας, κυρίως λόγω των εφαρμογών και των πεδίων στα οποία μπορεί να συνεισφέρει. Ως παραφράσεις θεωρούνται εναλλακτικές εκφράσεις μιας πρότασης που φέρουν την ίδια ή ισοδύναμη σημασιολογική πληροφορία σε μία συγκεκριμένη γλώσσα (Shapiro, 1992).

Ένας συμβολικός ορισμός που μπορεί να δοθεί για τις παραφράσεις και την δημιουργία σημασιολογικά ισοδύναμων προτάσεων είναι ο εξής: Θεωρούμε ότι η σημασία των προτάσεων μπορεί να αποδοθεί σε Κατηγορηματική Πρώτη Λογική Τάξης (ΚΛΠΤ) και ως Φ_1 είναι η έκφραση ΚΛΠΤ που σημασιολογικά αναπαριστά την πρόταση S_1 και κατά αντιστοιχία θεωρούμε ότι Φ_2 είναι η έκφραση ΚΛΠΤ που αντιπροσωπεύει σημασιολογικά την πρόταση S_2 . Ορίζουμε ότι η πρόταση S_2 είναι μια παράφραση της πρότασης S_1 αν και μόνο αν ισχύουν τα εξής:

$$(\Phi_1 \wedge B) \models \Phi_2$$

και

$$(\Phi_2 \wedge B) \models \Phi_1$$

όπου B αντιπροσωπεύει την κοινότυπη γνώση (common sense knowledge) και την αξιωματική θεμελίωση του πεδίου (Androutsopoulos & Malakasiotis, 2009).

Αξίζει να σημειωθεί ότι η αυτόματη παραγωγή συντακτικά ορθών παραφράσεων που διατηρούν τη σημασιολογία των προτάσεων φυσικής γλώσσας είναι ιδιαίτερα χρήσιμη και βρίσκει εφαρμογή σε ένα ευρύ φάσμα πεδίων και διαδικασιών (Fujita & Isabelle, 2015; Ganitkevitch, & Callison-Burch, 2014). Πιο συγκεκριμένα, οι μέθοδοι και τα συστήματα δημιουργίας παραφράσεων προτάσεων φυσικής γλώσσας είναι πολύ χρήσιμα σε εφαρμογές και σε πεδία όπως η αυτόματη μετάφραση κειμένου (machine translation), η εξαγωγή πληροφοριών από το κείμενο (information extraction), η απάντηση ερωτημάτων (question answering), η ανάκτηση πληροφορίας (information retrieval), η σημασιολογική ανάλυση και πολλές βασικές διαδικασίες της επεξεργασίας και της παραγωγής φυσικής γλώσσας (Berant & Liang, 2014). Για παράδειγμα, η παραγωγή παραφράσεων μπορεί να βοηθήσει στην ανάκτηση πληροφορίας και στην απάντηση ερωτημάτων μέσω της επέκτασης των ερωτημάτων με παραφράσεις τους και στην μηχανική μετάφραση με την επέκταση των προτάσεων προς μετάφραση με σημασιολογικά ισοδύναμες προτάσεις (Ganitkevitch & Callison-Burch, 2014). Επομένως, συστήματα που ενσωματώνουν μεθόδους παράφρασης μπορούν να βελτιώσουν τις επιδόσεις και την ακρίβειά τους (Kampeera & Cardey-Greenfield, 2012).

Ωστόσο, η ανάπτυξη προσεγγίσεων για την αυτόματη ανάλυση φυσικής γλώσσας και η παραγωγή σημασιολογικά ισοδύναμων προτάσεων είναι ένα πολύ σύνθετο πρόβλημα. Οι έρευνες έχουν δείξει ότι είναι από τα δυσκολότερα προβλήματα στον τομέα της Τεχνητής Νοημοσύνης, το οποίο θεωρείται ότι είναι AI-complete (Shapiro, 1992; Wong & Mooney, 2007; Yampolskiy, 2012) και απαιτεί πλήρη κατανόηση του περιεχομένου της φυσικής γλώσσας και τη μοντελοποίηση της σημασιολογίας των εννοιών που αναφέρονται σε αυτές (Zhao et al., 2009). Η διαδικασία παραγωγής σημασιολογικά ισοδύναμων προτάσεων είναι ιδιαίτερα δύσκολη διαδικασία, λόγω της ενδογενούς πολυπλοκότητας της φυσικής γλώσσας, όπου η εκφραστικότητα και η σύνταξή της μπορούν να προβάλουν πολλές παραλλαγές, καθώς υπάρχουν πολλοί διαφορετικοί τρόποι να εκφραστεί η ίδια πληροφορία.

Στην ενότητα αυτή, παρουσιάζεται μια μεθοδολογία για την ανάλυση προτάσεων φυσικής γλώσσας και την δημιουργία σημασιολογικά ισοδύναμων παραφράσεων τους. Η προσέγγιση, χρησιμοποιεί ένα σύνολο τεχνικών παράφρασης που μπορούν να τροποποιήσουν τη δομή της πρότασης και να αλλάξουν συγκεκριμένες λέξεις με στόχο να δημιουργήσουν σημασιολογικά ισοδύναμες παραφράσεις. Οι τεχνικές παράφρασης

μπορούν να μετατρέψουν δομικά μέρη του δέντρου εξάρτησης της πρότασης σε ισοδύναμη μορφή/ και επίσης να αλλάξουν λέξεις της πρότασης, με συνώνυμες και αντώνυμες. Επίσης, λεξικολογικοί πόροι, χρησιμοποιούνται για την παροχή συνώνυμων και αντώνυμων και δίνεται η σημασία των λέξεων που μπορούν να χρησιμοποιηθούν για να γίνει περιφραστική αντικατάστασή τους μέσα στην πρόταση. Στα πλαίσια της πειραματικής μελέτης αξιολογήθηκε η απόδοση της προσέγγισης στη δημιουργία συντακτικά ορθών παραφράσεων που διατηρούν το νόημα των προτάσεων και τα αποτελέσματα έδειξαν ότι η μεθοδολογία είναι ακριβής και λειτούργησε αποτελεσματικά ακόμη και σε πολύπλοκες προτάσεις.

4.2 Σχετικές Εργασίες

Τα τελευταία χρόνια, η παραγωγή και η αναγνώριση παραφράσεων έχουν προσελκύσει το ενδιαφέρον των ερευνητών στην επιστήμη των υπολογιστών και της επεξεργασίας φυσικής γλώσσας. Στη διεθνή βιβλιογραφία, υπάρχει μεγάλο ενδιαφέρον για σχεδιασμό προσεγγίσεων και ανάπτυξη συστημάτων για την αυτόματη ανάλυση και παράφραση κειμένου (Androutsopoulos & Malakasiotis, 2009; Ganitkevitch & Callison-Greenfield, 2010; Mandani & Dorr, 2010). Οι κύριες μέθοδοι και τεχνικές στο πεδίο της δημιουργίας παραφράσεων μπορούν να κατηγοριοποιηθούν σε τέσσερα βασικά είδη (Zhao et al., 2009), τα οποία είναι τα εξής: οι μέθοδοι που βασίζονται σε θησαυρούς (thesaurus-based), οι μέθοδοι στατιστικής μετάφρασης (statistic machine translation), οι μέθοδοι που βασίζονται σε κανόνες (rule based) και οι μέθοδοι που βασίζονται στην επεξεργασία της φυσικής γλώσσας (natural language processing).

Ο απλούστερος τύπος μεθόδων είναι οι μέθοδοι που βασίζονται σε θησαυρούς, που στηρίζονται στην αρχή της αντικατάστασης λέξεων μια δοθείσας πρόταση S_1 με συνώνυμους, χρησιμοποιώντας ένα σύνολο λεξικολογικών πόρων. Γενικά, οι μέθοδοι που βασίζονται σε θησαυρούς είναι αρκετά απλές, αλλά δεν λαμβάνουν υπόψη την δομή της πρότασης και δεν μπορούν να δημιουργήσουν πιο πολύπλοκες μορφές παραφράσεων.

Οι μέθοδοι στατιστικής μετάφρασης σε γενικές γραμμές αντιμετωπίζουν το πρόβλημα της παραγωγής παραφράσεων ως μονόγλωσση (monolingual) αυτόματη μετάφραση, όπου μια δεδομένη πρόταση S_1 μεταφράζεται σε μια πρόταση S_2 στην ίδια γλώσσα. Γενικά, αυτό το είδος μεθόδων χρειάζεται εκτεταμένα μεγάλο όγκο παράλληλων συνόλων κειμενικών δεδομένων για την εκπαίδευση τους, γεγονός που στις περισσότερες περιπτώσεις είναι δύσκολο να επιτευχθεί.

Οι μέθοδοι βασισμένοι σε κανόνες βασίζονται σε σύνολα κανόνων παράφρασης, τα οποία μπορούν να χρησιμοποιηθούν για να διαχειριστούν και να αλλάξουν μια δοθείσα πρόταση εισόδου S_1 , κυρίως με τροποποίηση των δομών που αναγνωρίζονται σε αυτές.

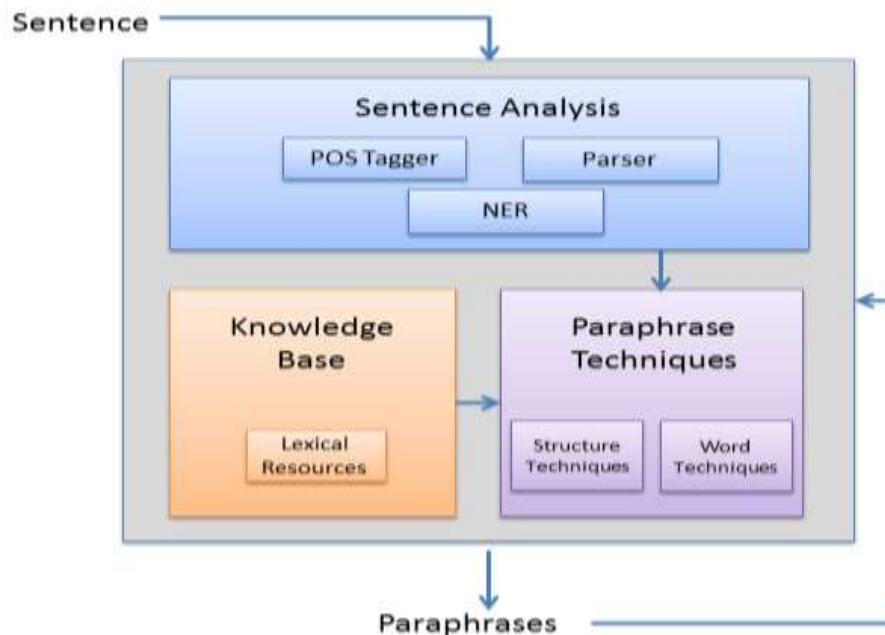
Οι μέθοδοι επεξεργασίας φυσικής γλώσσας αναλύουν μια δεδομένη πρόταση S_1 και καθορίζουν την σημασιολογική αναπαράστασή της. Στην συνέχεια, χρησιμοποιούν τεχνικές παραγωγής φυσικής γλώσσας για να παράγουν προτάσεις από την σημασιολογική αναπαράσταση της δεδομένης πρότασης. Ωστόσο, η διαμόρφωση ενός συστήματος παραγωγής φυσικής γλώσσας για μια δεδομένη πρόταση είναι μια αρκετά περίπλοκη διαδικασία.

Στην εργασία (Quirk et al., 2004), οι συγγραφείς παρουσιάζουν μια προσέγγιση στατιστικής μετάφρασης για να παράγουν παραφράσεις σε επίπεδο πρότασης, η οποία χρησιμοποιεί εκπαίδευση με ζεύγη σημασιολογικά ισοδύναμων προτάσεων που προέρχονται από άρθρα ειδήσεων. Τα αποτελέσματα της προσέγγισης αξιολογήθηκαν από εμπειρογνώμονα-ειδικό και έδειξαν ότι δημιουργούνται ποιητικές παραφράσεις. Στην εργασία (Malakasiotis, 2009), οι συγγραφείς παρουσιάζουν τεχνικές μηχανικής μάθησης για να αναγνωρίσουν παραφράσεις στο κείμενο. Οι συγγραφείς χρησιμοποιούν ταξινομητές μέγιστης εντροπίας (maximum entropy) που εκπαιδεύονται σε δεδομένα αποτελούμενα από ζεύγη σημασιολογικών ισοδύναμων παραφράσεων και μπορούν να αναγνωρίζουν νέες παραφράσεις με ακρίβεια έως και 76%. Στην εργασία (Narayan et al., 2016), οι συγγραφείς εισάγουν μια προσέγγιση που βασίζεται σε ένα μοντέλο γραμματικής για παραγωγή παραφράσεων, το οποίο βασίζεται σε γραμματικές χωρίς συμφραζόμενα και δεν απαιτείται εκπαίδευση σε σύνολο παραφράσεων. Στην εργασία (Zhao et al., 2010), οι συγγραφείς παρουσιάζουν μια προσέγγιση για την παραγωγή παραφράσεων σε επίπεδο πρότασης, η οποία βασίζεται στην αυτόματη μηχανική μετάφραση. Η προσέγγιση, δοθείσας μιας πρότασης στην αγγλική γλώσσα, χρησιμοποιεί τρία εργαλεία μετάφρασης, τα οποία είναι το Google Translate, το Microsoft Translator και το Systran Translator, για την αυτόματη μετάφρασή της σε 6 γλώσσες, οι οποίες είναι τα γερμανικά, γαλλικά, ισπανικά, ιταλικά, πορτογαλικά και τα κινέζικα και έτσι δημιουργούνται 18 (3x6) μεταφράσεις της αρχικής πρότασης σε αυτές τις γλώσσες. Στην συνέχεια, για κάθε μια πρόταση κάνει την αντίστροφη μετάφρασή ξανά στην αγγλική γλώσσα και δημιουργούνται με αυτό τον τρόπο 54 (18x3) παραφράσεις. Στην συνέχεια, οι συγγραφείς χρησιμοποιούν τεχνικές, όπως η Minimum Bayes Risk (MBR), για να αποτιμήσουν την ποιότητα κάθε μιας παραγόμενης παράφρασης και να κρατήσουν αυτές που είναι κάτω από ένα όριο. Τα αποτελέσματα

αξιολογήθηκαν από εμπειρογνώμονα ειδικών και έδειξαν ότι η προσέγγιση έχει πολύ καλή απόδοση. Στην εργασία (Mehay & White, 2012), οι συγγραφείς παρουσιάζουν μια προσέγγιση δημιουργίας παραφράσεων, που βασίζεται στην αντικατάσταση λέξεων καθώς και την χρησιμοποίηση εργαλείων κατάλληλης γραμματικής ανά-σύνταξης των προτάσεων. Στα πλαίσια της αντικατάστασης λέξεων για την δημιουργία παραφράσεων, οι συγγραφείς μελετούν την αντικατάσταση λέξεων σε συγκεκριμένα πρότυπα ακολουθιών λέξεων με στόχο να περιορίσουν την πιθανότητα λάθος αντικατάστασης λέξης. Επίσης, για την δημιουργία παραφράσεων μέσω της κατάλληλης γραμματικής ανά-σύνταξης μιας πρότασης, χρησιμοποιείται το εργαλείο OpenCCG (White, et al., 2008), που μετατρέπει με χρήση γραμματικών την πρόταση σε ορθή γραμματική μορφή.

4.3 Μηχανισμός Παραγωγής Παραφράσεων

Σε αυτή την ενότητα, παρουσιάζουμε την προσέγγιση και το σύστημα που αναπτύχθηκε για την παραγωγή σημασιολογικά ισοδύναμων προτάσεων φυσικής γλώσσας. Πιο συγκεκριμένα, γίνεται περιγραφή των αρχών που βασίστηκε ο σχεδιασμός και των τεχνικών που αναπτύχθηκαν για την ανάλυση και τον ακριβή προσδιορισμό παραφράσεων προτάσεων φυσικής γλώσσας. Η αρχιτεκτονική του συστήματος αποτελείται από τρεις κύριες μονάδες Sentence Analysis, Knowledge Base και Paraphrase Techniques. Η βασική αρχιτεκτονική της προσέγγισης παρουσιάζεται στην Εικόνα 4.1.



Εικόνα 4.1 Αρχιτεκτονική Συστήματος

Αρχικά, η δοθείσα πρόταση προς παράφραση υποβάλλεται σε εις βάθος ανάλυση της δομής της, καθορίζεται ο γραμματικός ρόλος των λέξεων και σχηματίζεται το δέντρο εξαρτήσεων της. Στην συνέχεια, χρησιμοποιείται ένα σύνολο τεχνικών παράφρασης που τροποποιούν τη δομή της πρότασης και αλλάζουν συγκεκριμένες λέξεις, με στόχο να δημιουργήσουν σημασιολογικά ισοδύναμες παραφράσεις. Οι τεχνικές παράφρασης μετατρέπουν δομικά μέρη του δέντρου εξάρτησης της πρότασης σε ισοδύναμη μορφή και επίσης αλλάζουν λέξεις της πρότασης με συνώνυμες και αντώνυμες. Η βάση γνώσης περιέχει λεξικολογικούς πόρους, όπως το WordNet, που χρησιμοποιούνται για την παροχή πληροφοριών για λέξεις, όπως είναι τα συνώνυμα και αντώνυμα τους, καθώς επίσης περιέχει την σημασία των λέξεων και που μπορεί να χρησιμοποιηθεί για να γίνει περιφραστική αντικατάστασή τους μέσα στην πρόταση.

4.3.1 Ανάλυση της Φυσικής Γλώσσας

Αρχικά, δοθείσας μιας πρότασης, πραγματοποιείται μορφοσυντακτική ανάλυσή της χρησιμοποιώντας το εργαλείο Tree Tagger (Schmid, 2013), το οποίο προσδιορίζει για κάθε λέξη της πρότασης τι μέρος του λόγου είναι, καθώς και τη βασική της μορφή (λήμμα). Η μορφοσυντακτική ανάλυση της δομής μια πρότασης είναι το πρώτο στάδιο για την επεξεργασία της. Μετά την μορφοσυντακτική ανάλυσή της, από το εργαλείο Tree Tagger, γίνεται μια βαθύτερη ανάλυση της δομής της πρότασης και προσδιορίζονται οι σχέσεις μεταξύ των λέξεων και δημιουργείται το *δέντρο εξαρτήσεων*. Το δέντρο εξαρτήσεων της πρότασης αναπαριστά τις γραμματικές σχέσεις μεταξύ των λέξεων, οι οποίες αναπαρίστανται ως τριπλέτες (triples). Κάθε τριπλέτα αποτελείται από το όνομα της σχέσης, τον κυβερνήτη (governor), που εκτελεί την σχέση, και το εξαρτώμενο (dependent), που δέχεται τη σχέση αντίστοιχα. Οι εξαρτήσεις αυτές δείχνουν τον ακριβή τρόπο που οι λέξεις της πρότασης συνδέονται και αλληλεπιδρούν μεταξύ τους και μπορούν να προσφέρουν σημαντικές πληροφορίες για τον τρόπο που προκύπτει το νόημά της. Έπειτα, γίνεται προσδιορισμός των κυρίων ονομάτων, που ενδεχομένως υπάρχουν στην πρόταση. Ο προσδιορισμός των κυρίων ονομάτων πραγματοποιείται από το εργαλείο Named Entity Recognizer (NER), το οποίο αναγνωρίζει λέξεις σε μια πρόταση οι οποίες είναι ονόματα πραγμάτων (π.χ. προσώπων) και προσδιορίζει το είδος της οντότητας τους.

Η ανάλυση της πρότασης γίνεται με στόχο να αναγνωριστούν συγκεκριμένες δομές στο δέντρο εξάρτησης, για να εφαρμοστούν σ' αυτό τεχνικές παράφρασης μετατρέποντάς το δέντρο σε ισοδύναμη μορφή και στη συνέχεια να αντιστοιχηθούν οι αλλαγές στην αρχική πρόταση, έτσι ώστε να δημιουργηθεί μια παράφραση της αρχικής πρότασης. Για το σκοπό

αυτό, σχεδιάστηκαν και δημιουργήθηκαν τεχνικές δημιουργίας παραφράσεων που εφαρμόζόμενες σε προτάσεις φυσικής γλώσσας δημιουργούν σημασιολογικά ισοδύναμες προτάσεις.

4.3.2 Δημιουργία Τεχνικών Παράφρασης

Για την δημιουργία παραφράσεων σχεδιάσαμε και αναπτύξαμε τεχνικές παράφρασης, που τις κατηγοριοποιούμε σε δυο βασικές κατηγορίες: τις *δομικές τεχνικές* (*structure techniques*) και τις *λεξικές τεχνικές* (*word techniques*).

Οι *δομικές τεχνικές* βασίζονται στην ανάλυση της δομής της πρότασης και μπορούν να εφαρμοστούν σε συγκεκριμένα πρότυπα (patterns) του δέντρου εξαρτήσεων της πρότασης, προκειμένου να την μετασχηματίσουν σε ισοδύναμη μορφή. Αξίζει να σημειωθεί ότι σημαντικό ρόλο στην εφαρμογή των δομικών τεχνικών αποτελεί ο προσδιορισμός της κεντρικής δομής της πρότασης «subject-verb-object», η οποία αποτελεί τη βασική δομή της και καθορίζει το βασικό νόήμά της πρότασης. Κάθε τεχνική έχει ως στόχο να πραγματοποιήσει κατάλληλο μετασχηματισμό μέρους ή μερών της δομής αυτής. Πιο συγκεκριμένα, οι δομικές τεχνικές που σχεδιάστηκαν και αναπτύχθηκαν στα πλαίσια της προσέγγισης μας είναι οι εξής:

T1: subject -verb- object $\leftrightarrow$ object-passive_voice(verb)-subject

T2: subject -verb- object $\leftrightarrow$ subject-noun(verb)-object

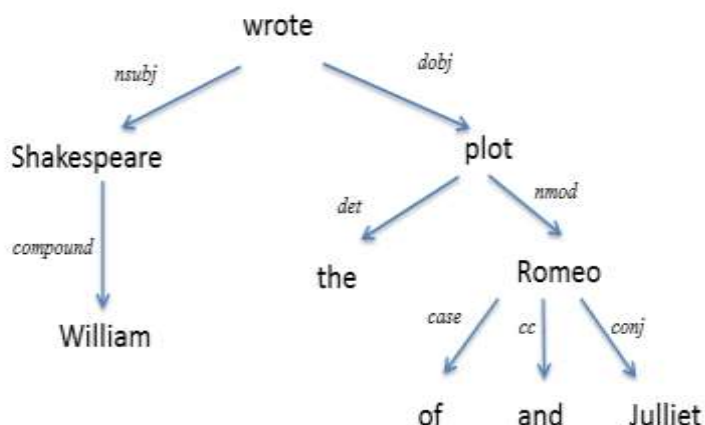
T3: compound dependency between noun1 + noun2 $\leftrightarrow$ noun2 + “of the” +noun1

T4: Αναδιάταξη χρονικών προσδιορισμών (tmod)

Η πρώτη τεχνική (T1) έχει ως στόχο να μετατρέψει μια φράση από την ενεργητική φωνή στην παθητική φωνή με την αλλαγή της δομής στην μορφή: subject(obj)-passive_voice (verb)-object(sub). Για παράδειγμα, ας θεωρήσουμε ότι έχουμε την ακόλουθη πρόταση:

Π1: William Shakespeare wrote the plot of Romeo and Juliet.

Η ανάλυση της πρότασης προσδιορίζει για κάθε λέξη τον γραμματικό της ρόλο, αναλύει την δομή της πρότασης και δημιουργείται το δέντρο εξαρτήσεων της, το οποίο παρουσιάζεται στην Εικόνα 4.2.



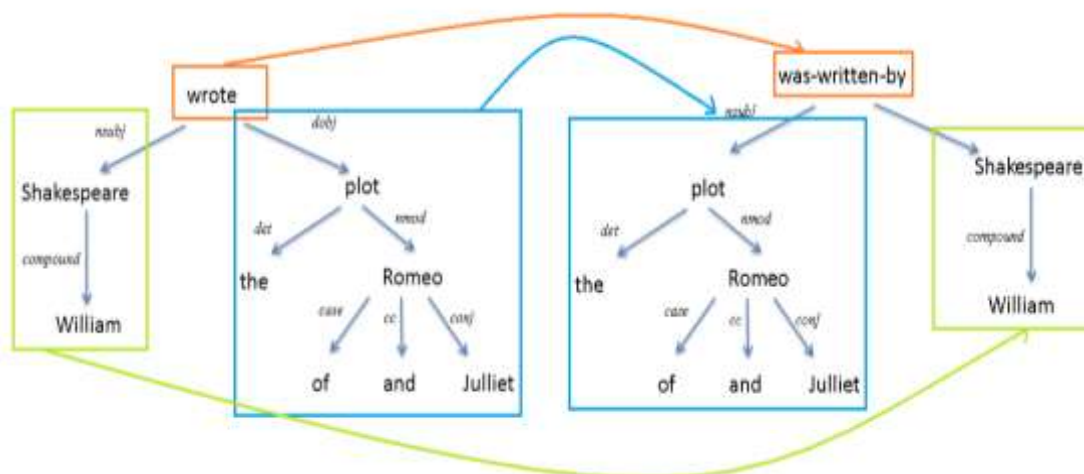
Εικόνα 4.2 Το δέντρο εξαρτήσεων της πρότασης Π1

Επίσης, γίνεται προσδιορισμός των κύριων οντοτήτων της πρότασης και η ανάλυση προσδιορίζει:

<PERSON>William Shakespeare</PERSON> wrote the plot of <PERSON> Romeo </PERSON> and <PERSON>Juliet</PERSON>

αναγνωρίζονται τρία κύρια ονόματα τα “ William Shakespeare”, “Romeo” και “Juliet”.

Αρχικά, χρησιμοποιώντας πληροφορίες της βάσης γνώσης γίνεται προσδιορισμός της παθητικής μορφής της ρηματικής φράσης “wrote” και στην συνέχεια, μετασχηματίζεται το δέντρο εξαρτήσεων αντιμεταθέτοντας μεταξύ τους τα υποδέντρα του υποκειμένου και του αντικειμένου όπως παρουσιάζεται στην Εικόνα 4.3.



Εικόνα 4.3 Οι μετατροπές στο δέντρο εξαρτήσεων της Π1.

Το αποτέλεσμα της εφαρμογής της τεχνικής στην πρόταση είναι η παραγωγή της ακόλουθης ισοδύναμης πρότασης:

ΠΠ1-T1: The plot of Romeo and Juliet was written by William Shakespeare.

Η δεύτερη τεχνική (T2), που βασίζεται επίσης στην κεντρική δομή της πρότασης, μετατρέπει τη δομή subject-verb-object στην μορφή subject-noun(verb)-object, όπου το noun(verb) αναπαριστά την μετατροπή του ρήματος της πρότασης στο αντίστοιχο περιφραστικό ουσιαστικό (noun) του. Για παράδειγμα, ας θεωρήσουμε και πάλι την πρόταση:

Π1: William Shakespeare wrote the plot of Romeo and Juliet.

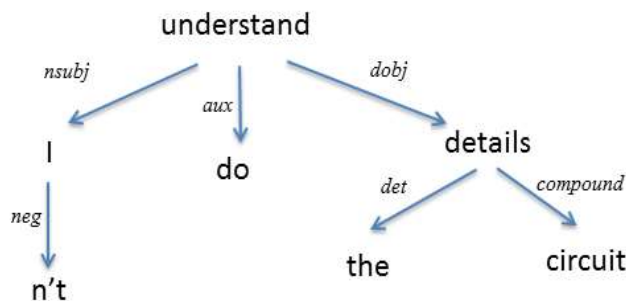
Από την ανάλυση προσδιορίζεται το κεντρικό ρήμα “write” και το οποίο μεταφράζεται στη μορφή του περιφραστικού ουσιαστικού με την αντικατάστασή του από την μορφή “be the + verb + of”. Το αποτέλεσμα είναι η μετατροπή της πρότασης στην μορφή:

ΠΠ1-T2: William Shakespeare is the writer of the plot of Romeo and Juliet.

Μια άλλη τεχνική που αναπτύχθηκε (T3) βασίζεται στην ανάλυση προτάσεων που τα δέντρα τους έχουν ως πρότυπα (patterns) δύο διαδοχικά ουσιαστικά. Ο μετασχηματισμός που γίνεται αναλύει φράσεις ουσιαστικών (noun phrases -NPs) της μορφής “noun1–noun2”, οι οποίες συνδέονται με την εξάρτηση “compound”, και πραγματοποιεί αντικατάστασή τους με τη δομή “noun2 of the noun1”. Για παράδειγμα, ας θεωρήσουμε την ακόλουθη πρόταση:

Π2 : I don’t understand the circuit details.

Η ανάλυση της πρότασης προσδιορίζει για κάθε λέξη τον γραμματικό της ρόλο, αναλύει την δομή της πρότασης και δημιουργεί το δέντρο εξαρτήσεων της, το οποίο παρουσιάζεται στην Εικόνα 4.4.



Εικόνα 4.4 Το δέντρο εξαρτήσεων της πρότασης Π2

Η αναγνώριση και ανάλυση της σχέσης “compound” μεταξύ των ουσιαστικών “circuit” και “details” οδηγεί στην αναδιαμόρφωση τους στην έκφραση ως “details of the circuit”.

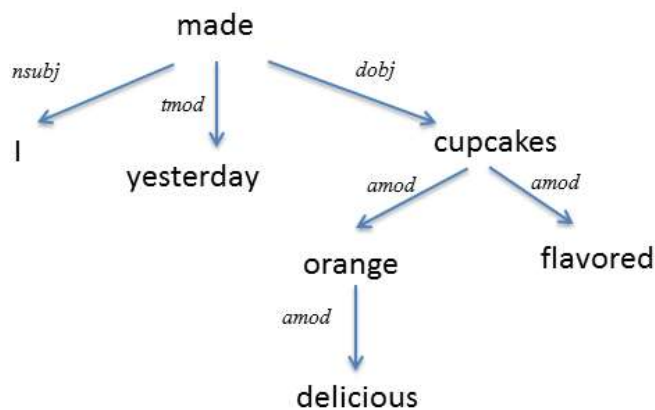
Επομένως, η εφαρμογή της τεχνικής αυτής έχει ως αποτέλεσμα την παραγωγή της ακόλουθης παράφρασης της πρότασης:

ΠΠ2-T3: I don’t understand the details of the circuit.

Μια τελευταία τεχνική (T4) που αναπτύχθηκε αφορά την αναδιάταξη των δομών των προτάσεων που περιέχουν χρονικούς προσδιορισμούς. Οι χρονικοί προσδιορισμοί αναγνωρίζονται μέσω της εξάρτησης “tmod” στο δέντρο εξαρτήσεων της πρότασης. Στόχος της τεχνικής είναι η μεταφορά των χρονικών προσδιορισμών σε άλλα σημεία-δομές της πρότασης, όπως στο τέλος της πρότασης ή μετά την ρηματική φράση. Για παράδειγμα, ας θεωρήσουμε την ακόλουθη πρόταση:

Π3: Yesterday, I made delicious, orange flavored cupcakes.

Η ανάλυση της πρότασης προσδιορίζει για κάθε λέξη τον γραμματικό της ρόλο, αναλύει την δομή της πρότασης και δημιουργεί το δέντρο εξαρτήσεών της, το οποίο παρουσιάζεται στην Εικόνα 4.5.



Εικόνα 4.5 Το δέντρο εξαρτήσεων της πρότασης Π3

Η θέση της χρονικής αναφοράς “yesterday” αλληλεπιδρά με την λέξη “made” και η σχέση τους αναγνωρίζεται ότι είναι σχέση χρονικού προσδιορισμού “tmod”. Η τεχνική που σχεδιάστηκε, ενεργοποιείται όταν αναγνωριστούν στο δέντρο εξαρτήσεων της πρότασης χρονικές αναφορές “tmod” μεταξύ λέξεων. Η εφαρμογή της τεχνικής αλλάζει τη θέση της χρονικής αναφοράς μεταφέροντάς την αμέσως μετά τη ρηματική φράση με την οποία

αλληλεπιδρά ή στο τέλος της πρότασης. Επομένως, η εφαρμογή της τεχνικής αυτής έχει ως αποτέλεσμα την παραγωγή των ακόλουθων δυο παραφράσεων της πρότασης:

ΠΠ3-T4α: I made yesterday delicious, orange flavored cupcakes.

ΠΠ3-T4β: I made delicious, orange flavored cupcakes yesterday.

Από την άλλη πλευρά, οι λεξικές τεχνικές βασίζονται στην αντικατάσταση συγκεκριμένων συνόλων από λέξεις της αρχικής πρότασης με τις συνώνυμές τους και, σε ορισμένες περιπτώσεις, με τις αντώνυμές τους. Για να γίνει αυτό, χρησιμοποιείται η λεξιλογική βάση δεδομένων WordNet. Το WordNet ομαδοποιεί τις λέξεις σε σύνολα συνωνύμων και παρέχει σύντομους ορισμούς και παραδείγματα χρήσης, και επίσης καταγράφει μια σειρά σχέσεων μεταξύ των συνόλων συνωνύμων ή των μελών τους. Το εργαλείο WordNet που χρησιμοποιείται έχει διττό ρόλο. Πρώτον, χρησιμοποιείται με σκοπό να παρέχει τη σημασία των λέξεων, η οποία μπορεί να χρησιμοποιηθεί για να γίνει περιφραστική αντικατάστασή τους μέσα στην πρόταση. Δεύτερον, χρησιμοποιείται για να καθορίσει συνώνυμες και αντώνυμες λέξεις συγκεκριμένων λέξεων της πρότασης.

Μία τεχνική αυτού του τύπου βασίζεται στην αντικατάσταση λέξεων της πρότασης με τους ορισμούς τους, όπως αυτοί καταγράφονται στη βάση γνώσης του συστήματος. Για παράδειγμα, ας θεωρήσουμε την πρόταση:

Π4: Yesterday, I made delicious, orange flavored cupcakes.

Από την ανάλυση της πρότασης, προσδιορίζεται η σύνταξή της, το μέρος του λόγου για κάθε λέξη, καθώς και το δέντρο εξαρτήσεών της. Η τεχνική εφαρμόζεται σε λέξεις της πρότασης με το να παρέχει τον περιφραστικό ορισμό τους, όπως αυτός καταγράφεται στην βάση γνώσης, με στόχο να γίνει αντικατάσταση κάθε λέξης με τον περιφραστικό ορισμό της. Επομένως, η εφαρμογή της τεχνικής αυτής έχει ως αποτέλεσμα την παραγωγή των ακόλουθων παραφράσεων της πρότασης:

ΠΠ4-T5α: Day immediately before today, I made delicious, orange flavored cupcakes.

ΠΠ4-T5β: Day immediately before today, I made extremely pleasing to the sense of taste, orange flavored cupcakes.

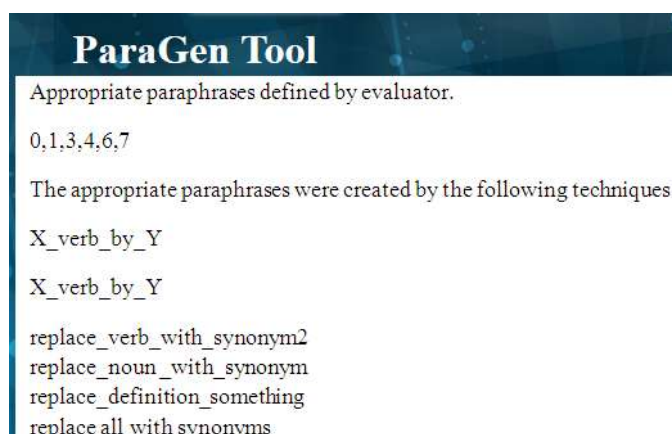
ΠΠ4-T5γ: Day immediately before today, I made delicious, orange flavored small cake baked in a muffin tin.

Η διεπαφή του εργαλείου που αναπτύχθηκε και που υλοποιεί την προσέγγιση παραγωγής παραφράσεων παρουσιάζεται στην Εικόνα 4.6.



Εικόνα 4.6 Η διεπαφή του εργαλείου παραγωγής παραφράσεων.

Οι παραφράσεις που παράγονται από το εργαλείο παρουσιάζονται στον ειδικό-εμπειρογνώμονα για να τις αξιολογήσει. Ο εμπειρογνώμονας καλείται να αξιολογήσει την ποιότητα της κάθε παράφρασης που δημιουργείται από το εργαλείο, με βάση την συντακτική ορθότητα και την διατήρηση του αρχικού νοήματος. Στην περίπτωση όπου και οι δύο προϋποθέσεις ικανοποιούνται, η παράφραση χαρακτηρίζεται ως “κατάλληλη”. Σε διαφορετική περίπτωση, αν κάποια από τις προϋποθέσεις δεν πληρείται, τότε χαρακτηρίζεται ως “μη κατάλληλη”.



Εικόνα 4.7 Προσδιορισμός καταλληλότητας των παραγόμενων παραφράσεων.

Μετά την εφαρμογή των τεχνικών παράφρασης, οι τεχνικές που παρήγαγαν τις κατάλληλες παραφράσεις, καθώς και η πρόταση μαζί με τις κατάλληλες παραφράσεις αποθηκεύονται στην βάση του εργαλείου.

4.4 Πειραματική Αξιολόγηση

Για την αξιολόγηση της απόδοσης της προσέγγισής και του εργαλείου που αναπτύχθηκε διενεργήθηκε μια πειραματική μελέτη. Για τις ανάγκες της μελέτης, χρησιμοποιήθηκαν προτάσεις από το σύνολο δεδομένων MSR (Dolan et al., 2005) καθώς και προτάσεις φυσικής γλώσσας που συλλέχθηκαν από διάφορες πηγές, όπως τίτλους ειδήσεων στο διαδίκτυο και από βιβλία λογικής και δημιουργήθηκε ένα σύνολο από 140 προτάσεις. Στην συνέχεια παρουσιάζονται μερικά παραδείγματα προτάσεων:

1. The exercise that we solved yesterday was very difficult.
2. Yesterday, I made delicious, orange flavored cupcakes.
3. All purple mushrooms are poisonous.
4. John is happy as long as Maria is happy.
5. He agreed to buy that car.
6. Learn the details of circuit.
7. Yesterday was a very cloudy day.
8. I don't understand the circuit details.
9. All planets revolve around a star.
10. Today is the most rainy day.
11. The binary code was invented by Leibniz.
12. That gift had been bought for me.
13. Cats are happy as long as someone feeds them.
14. At the carnival, Helen was helping Anna, but I wrote the letters.
15. Every city has a dog catcher of it that has been bitten by every dog living in the city.

Στα πλαίσια της μελέτης, το σύνολο των παραφράσεων που παράχθηκαν από το εργαλείο για τις 140 αρχικές προτάσεις, χρησιμοποιήθηκαν για να προσδιοριστεί η απόδοσή του. Για κάθε πρόταση, όλες οι παραφράσεις της παρουσιάστηκαν και αξιολογήθηκαν από τον εμπειρογνώμονα-ειδικό για να γίνει προσδιορισμός της απόδοσης του εργαλείου. Η απόδοση σε κάθε πρόταση υπολογίζεται ως ο λόγος των παραφράσεων που παρήγαγε το εργαλείο και που προσδιορίστηκαν από τον ειδικό ως κατάλληλες προς όλες τις

παραφράσεις που παρήγαγε το εργαλείο. Πιο συγκεκριμένα, η απόδοση σε επίπεδο πρότασης υπολογίζεται ως εξής:

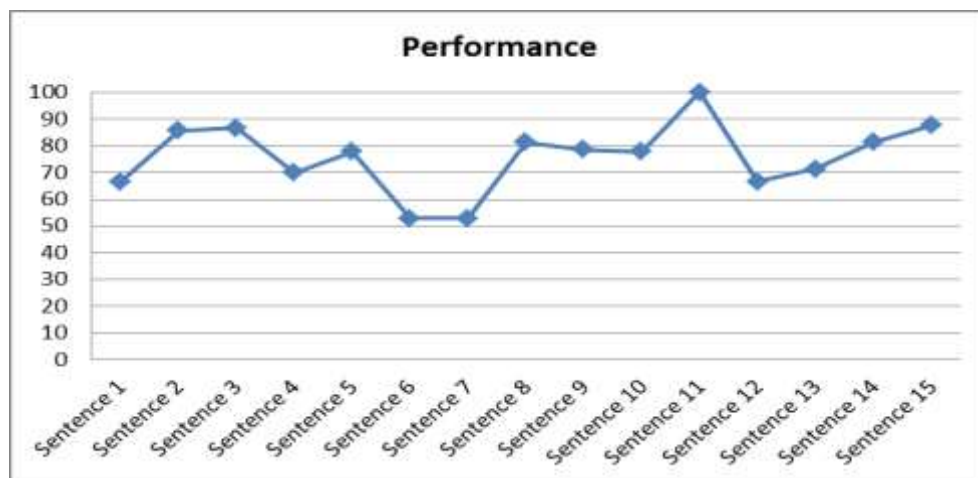
$$performance = \frac{appropriate_paraphrases}{all_paraphrases_generated} \quad (1)$$

Ενώ η συνολική απόδοση του εργαλείου σε αυτό το μέρος του πειράματος ορίστηκε ως εξής:

$$total_performance = \frac{\sum_{i=0}^n performance_i}{n} \quad (2)$$

Όπου n αναπαριστά τον αριθμό των προτάσεων και $performance_i$ είναι η απόδοση του εργαλείου για την πρόταση i .

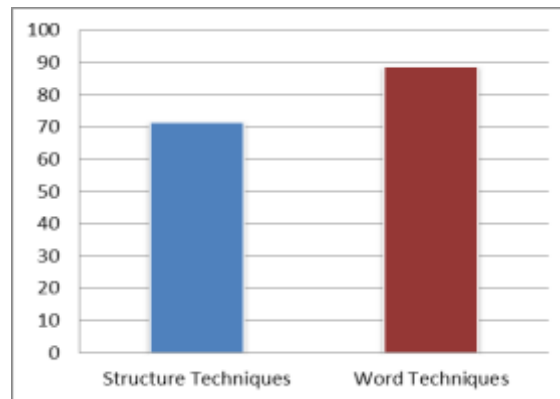
Στα πλαίσια της πειραματικής μελέτης, η συνολική απόδοση του συστήματος υπολογίστηκε ότι ήταν 76.4%. Αυτό σημαίνει ότι το 76.4% των παραφράσεων που δημιουργήθηκαν αξιολογήθηκαν από το ειδικό ως κατάλληλες. Στην Εικόνα 4.8, παρουσιάζεται η απόδοση του εργαλείου σε κάθε πρόταση όπως αξιολογήθηκε από τον εμπειρογνώμονα.



Εικόνα 4.8 Απόδοση του μηχανισμού σε δείγμα 15 προτάσεων

Στην συνέχεια, έγινε προσδιορισμός της απόδοσης των ειδών των τεχνικών παράφρασης που δημιουργήθηκαν. Πιο συγκεκριμένα, για το σύνολο των παραφράσεων που δημιουργούνται από το εργαλείο για κάθε πρόταση, αξιολογήθηκε η απόδοσή των ειδών των τεχνικών παράφρασης μέσω του τύπου (1), και η απόδοσή τους υπολογίστηκε ως το ποσοστό των κατάλληλων παραφράσεων που δημιουργήθηκαν προς τις συνολικές

παραγόμενες παραφράσεις. Η συνολική απόδοση των τεχνικών απεικονίζεται στην Εικόνα 4.9.



Εικόνα 4.9 Η απόδοση για κάθε είδος τεχνικής

Από τα αποτελέσματα προκύπτει ότι λεξικές τεχνικές παρουσιάζουν καλύτερη επίδοση και μεγαλύτερη ακρίβεια στην δημιουργία κατάλληλων παραφράσεων. Η απόδοση αυτή οφείλεται κυρίως στο γεγονός ότι οι λεξικές τεχνικές μεταβάλλουν μικρά μέρη της πρότασης, όπως είναι η πραγματοποίηση αντικατάστασης μιας λέξης με μια συνώνυμη ή ανώνυμη ή ακόμη και με τον ορισμό της. Στις περισσότερες περιπτώσεις, οι λεκτικές αντικαταστάσεις που πραγματοποιούν διατηρούν την συντακτική ορθότητα των προτάσεων καθώς και τη σημασιολογία τους. Από την άλλη πλευρά, οι δομικές τεχνικές επηρεάζουν περισσότερα μέρη της δομής μιας πρότασης, και οι τροποποιήσεις που πραγματοποιούν είναι πιο πιθανό να οδηγήσουν σε συντακτικά και σημασιολογικά λάθη. Τα συνολικά αποτελέσματα της μελέτης έδειξαν ότι η μεθοδολογία και το εργαλείο που αναπτύχθηκε μπορούν να παρέχουν έναν άμεσο και ακριβή τρόπο ανάλυσης προτάσεων φυσικής γλώσσας και δημιουργίας σημασιολογικά ισοδύναμων προτάσεων.

Ένα βασικό πλεονέκτημα της προσέγγισής μας σε σχέση με τις προσεγγίσεις της διεθνούς βιβλιογραφίας αποτελεί η δυνατότητα ανασχηματισμού των δομών των προτάσεων και η παραγωγή πιο εκφραστικών και αναδομημένων παραφράσεων. Πιο συγκεκριμένα, οι περισσότερες προσεγγίσεις της διεθνούς βιβλιογραφίας βασίζονται σε τεχνικές μηχανικής μετάφρασης κειμένου, κάτι που δημιουργεί περιορισμούς ως προς την εκτεταμένη αναδόμηση της συντακτικής δομής των προτάσεων και οι παραφράσεις που παράγονται συνήθως, αν και ακριβείς, εμφανίζουν περιορισμένη συντακτική διαφοροποίηση μεταξύ τους. Ένα βασικό πλεονέκτημα της προσέγγισής μας αποτελεί, η λειτουργία της εις βάθος ανάλυσης των προτάσεων που πραγματοποιεί και η δημιουργία

παραφράσεων μέσω της τροποποίησης συγκεκριμένων εξαρτήσεων των δέντρων εξαρτήσεων των προτάσεων, γεγονός που συμβάλει στην παραγωγή πιο εκφραστικών παραφράσεων, οι οποίες διαφοροποιούνται συντακτικά σε μεγάλο βαθμό μεταξύ τους.

5

Επεξεργασία Κειμένου και Προσδιορισμός Συναισθηματικού Περιεχομένου

5.1 Εισαγωγή

Τα συναισθήματα αποτελούν ένα βασικό χαρακτηριστικό της ανθρώπινης νοημοσύνης που χρωματίζουν τον τρόπο της επικοινωνίας και διαδραματίζουν σημαντικό ρόλο στην αλληλεπίδραση μεταξύ ανθρώπου-υπολογιστή. Ο ρόλος των συναισθημάτων αρχικά ερευνήθηκε από την Rosalind Picard, που εισήγαγε την έννοια της συναισθηματικής υπολογιστικής (affective computing) (Picard, 1997), υποδεικνύοντας τη σημασία των συναισθημάτων στην αλληλεπίδραση ανθρώπου-υπολογιστή και ορίζοντας μια κατεύθυνση για την διεπιστημονική έρευνα σε τομείς όπως η τεχνητή νοημοσύνη, η γνωστική επιστήμη και η επεξεργασία φυσικής γλώσσας. Ο κύριος στόχος της συναισθηματικής υπολογιστικής είναι να αναπτύξει μηχανισμούς που να επιτρέπουν σε υπολογιστικά συστήματα να αναλύουν, να αναγνωρίζουν και να προσαρμόζονται στις συναισθηματικές καταστάσεις του χρήστη (Calvo & D' Mello, 2010).

Η αυτόματη αναγνώριση των συναισθηματικών καταστάσεων του χρήστη μπορεί να συμβάλει στην βελτίωση της ποιότητας της αλληλεπίδρασης ανθρώπου-υπολογιστή και να βοηθήσει ένα υπολογιστικό σύστημα να κατανοήσει και να προσαρμοστεί καλύτερα στην συναισθηματική κατάσταση του χρήστη. Η βελτίωση της αλληλεπίδρασης μπορεί να αυξήσει την αποτελεσματικότητα της λειτουργίας συστημάτων και εφαρμογών. Στα πλαίσια των ευφών συστημάτων διδασκαλίας (ΕΣΔ), η αναγνώριση της συναισθηματικής κατάσταση των εκπαιδευόμενων και η προσαρμογή των ΕΣΔ στις ανάγκες του κάθε εκπαιδευόμενου μπορούν να βελτιώσουν σημαντικά την διαδικασία της διδασκαλίας ώστε

να γίνει πιο εξατομικευμένη και αποτελεσματική (Shen et al., 2009). Επίσης, οι διαλογικοί πράκτορες (Embodied Conversational Agents- ECAs) μπορούν να επωφεληθούν σημαντικά από την αναγνώριση των συναισθηματικών καταστάσεων των χρηστών, επιτυγχάνοντας πιο ρεαλιστικές αλληλεπιδράσεις σε συναισθηματικό επίπεδο, όπως για παράδειγμα να προσαρμόζουν τη φωνή τους με βάση το συναισθηματικό περιεχόμενο του κειμένου επιτυγχάνοντας έτσι πιο ομαλή και πιο ζωντανή αλληλεπίδραση.

Επίσης, μια σημαντική συνεισφορά αποτελεί η ανάλυση δεδομένων κειμένων χρηστών στα κοινωνικά δίκτυα συμβάλλοντας στον προσδιορισμό της συμπεριφοράς κάθε χρήστη και επίσης στη δημόσια στάση σε διάφορα θέματα. Η ανάλυση των κειμένων και ο καθορισμός του συναισθηματικού περιεχομένου είναι ένα πολύ ενδιαφέρον και υψηλής πρόκλησης πεδίο έρευνας στην περιοχή του microblogging (De Choudhury et al., 2012), όπου από τον προσδιορισμό της συναισθηματικής κατάστασης των χρηστών μπορεί να εξαχθούν ουσιαστικές πληροφορίες σχετικά με την συμπεριφορά τους (Qiu et al., 2012).

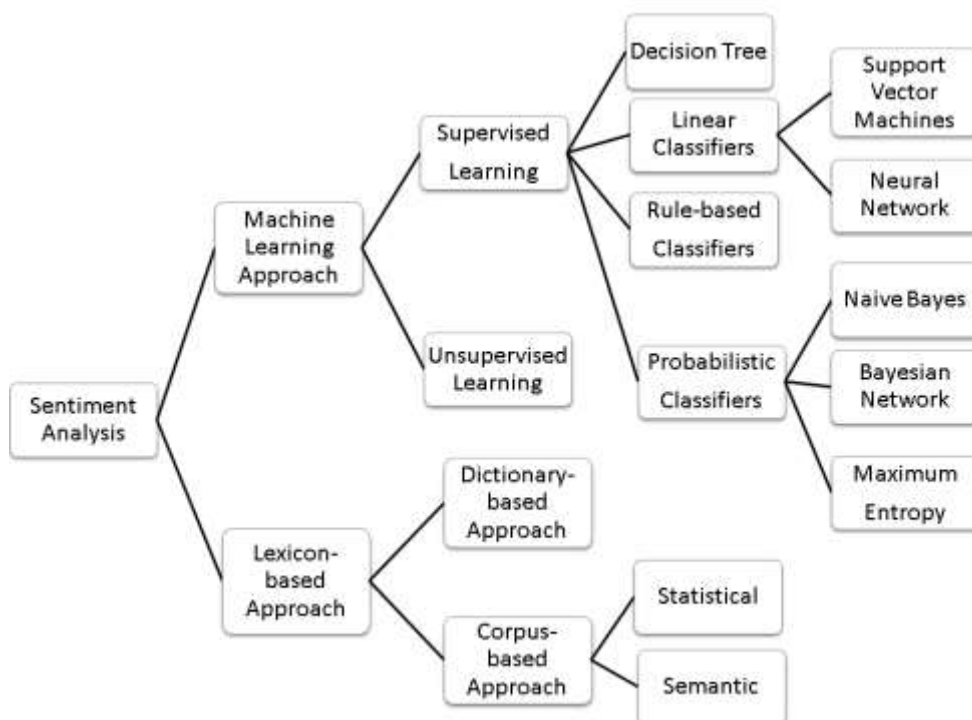
Τα συναισθήματα μπορούν να εκφραστούν μέσω διαφόρων μέσων, όπως η ομιλία, οι εκφράσεις του προσώπου, οι χειρονομίες και το γραπτό κείμενο. Ο πιο συνηθισμένος τρόπος είναι μέσω γραπτού κείμενου, ο οποίος αποτελεί το κύριο μέσο επικοινωνίας και δόμησης της πληροφορίας στο διαδίκτυο. Τα τελευταία χρόνια η εξέλιξη της τεχνολογίας και κυρίως η εξάπλωση του παγκόσμιου ιστού και των κοινωνικών δικτύων έχουν αλλάξει τον τρόπο της επικοινωνίας και παρέχουν νέα μέσα που συνδέουν ανθρώπους σε όλο τον κόσμο με πληροφορίες, ειδήσεις και γεγονότα σε πραγματικό χρόνο. Επίσης, έχουν αλλάξει εντελώς το ρόλο των χρηστών και τους έχουν μετατρέψει από απλούς, παθητικούς καταναλωτές πληροφορίας σε ενεργούς παραγωγούς. Καθημερινά, ένας τεράστιος αριθμός άρθρων και προσωπικών μηνυμάτων κειμένου αναρτώνται σε ειδησεογραφικά portals και κοινωνικά δίκτυα και η τεράστια αυτή ποσότητα δεδομένων κειμένου απαιτεί αυτοματοποιημένες μεθόδους για την ανάλυση τους και την εξαγωγή γνώσης (Anusha & Sandhya, 2015; Shaheen et al., 2014).

Στα πλαίσια αυτού του κεφαλαίου, παρουσιάζουμε νέες προσεγγίσεις για την αυτόματη ανάλυση προτάσεων φυσικής γλώσσας και την εξαγωγή των συναισθηματικών πληροφοριών που φέρουν. Πιο συγκεκριμένα, εισάγονται μέθοδοι που βασίζονται σε σχήματα ομαδοποιημένων ταξινομητών (ensemble classifiers), οι οποίες συνδυάζουν μεθόδους ανάλυσης των μορφοσυντακτικών δομών των προτάσεων φυσικής γλώσσας με τεχνικές μηχανικής μάθησης. Αρχικά, παρουσιάζεται ο σχεδιασμός και η ανάπτυξη μεθόδου για την εις βάθος ανάλυση προτάσεων φυσικής γλώσσας, όπου γίνεται προσδιορισμός των συντακτικών δομών και των εξαρτήσεων μεταξύ των λέξεων, χρήση λεξικολογικών πόρων

για τον καθορισμό του συναισθηματικού περιεχομένου των λέξεων, ποσοτικοποίηση τους με βάση τις εξαρτήσεις της πρότασης και καθορισμός του συνολικού περιεχομένου της πρότασης με βάση τις συντακτικές δομές της. Στην συνέχεια, παρουσιάζεται η χρήση μεθόδων μηχανικής μάθησης και η εκπαίδευση στατιστικών ταξινομητών, όπως οι *απλοϊκός bayes* (naïve Bayes) και *μηχανές διανυσμάτων υποστήριξης* (support vector machines-SVMs) και ο συνδυασμός τους σε σχήματα ομαδοποιημένων ταξινομητών. Τα σχήματα ομαδοποιημένων ταξινομητών που σχεδιάστηκαν και αναπτύχθηκαν συνδυάζουν τους βασικούς ταξινομητές κάτω από διάφορες μεθόδους ομαδοποίησης όπως η αρχή της πλειοψηφίας καθώς και μεθόδους ομαδοποίησης των ταξινομητών μηχανικής μάθησης όπως οι bagging, boosting, με στόχο να γίνει αποτελεσματικός συνδυασμός των επιμέρους ταξινομητών, να αυξηθεί η ακρίβεια των προβλέψεων των ομαδοποιημένων σχημάτων, καθώς επίσης να βελτιωθεί η απόδοση και η κλιμάκωσή τους σε νέα δεδομένα. Τέλος, γίνεται μελέτη της απόδοσης τόσο των βασικών ταξινομητών όσο και των ομαδοποιημένων σχημάτων ταξινόμησης που αναπτύχθηκαν με διάφορες μεθόδους ομαδοποίησης και παρουσιάζονται τα αποτελέσματα τους σε διάφορα είδη δεδομένων κειμένων.

5.2 Μέθοδοι, Λεξικολογικοί Πόροι και Σχετικές Εργασίες

Οι βασικές μέθοδοι και προσεγγίσεις για την αναγνώριση του συναισθηματικού περιεχομένου κείμενου μπορούν να διαχωριστούν σε δύο κύρια είδη: τις προσεγγίσεις που βασίζονται σε τεχνικές μηχανικής μάθησης (machine-learning) και τις προσεγγίσεις που βασίζονται στην γνώση και τα λεξικά (knowledge-based/lexicon-based) (Chaffar & Inkpen, 2011).



Εικόνα 5.1 Μέθοδοι και τεχνικές ανάλυσης και συναισθηματικής ταξινόμησης κειμένου.

Οι προσεγγίσεις μηχανικής μάθησης βασίζονται σε αλγορίθμους και τεχνικές της μηχανικής μάθησης για να εκπαιδεύσουν ταξινομητές χρησιμοποιώντας συλλογές κειμένων και χαρακτηριστικά που έχουν εξαχθεί ώστε να μπορούν οι ταξινομητές να αναγνωρίζουν το συναισθηματικό περιεχόμενο νέων κειμένων. Στις περισσότερες περιπτώσεις, οι ταξινομητές εκπαιδεύονται σε σχολιασμένες (annotated) συλλογές κειμένων και προτάσεων όπου για κάθε πρόταση έχει καθοριστεί το συναισθηματικό περιεχόμενό της και στην συνέχεια γίνεται εξαγωγή χαρακτηριστικών και αναπαράστασης της κάθε πρότασης σε κατάλληλη μορφή ώστε να γίνει εκπαίδευση των ταξινομητών. Από την άλλη πλευρά, οι προσεγγίσεις βασισμένες στην γνώση (lexicon based/knowledge based) προσδιορίζουν το συναισθηματικό περιεχόμενο νέων προτάσεων με βάση τεχνικές επεξεργασίας φυσικής γλώσσας και λεξικών που καθορίζουν το συναισθηματικό περιεχόμενο των λέξεων (Medhat et al., 2014).

5.2.1 Λεξικολογικοί Πόροι

Στην διεθνή βιβλιογραφία έχουν αναπτυχθεί γενικού και ειδικού σκοπού λεξικολογικοί πόροι (lexicon resources) με στόχο να ενισχυθεί η αποτελεσματικότητα των συστημάτων ανάλυσης κειμένου και αναγνώρισης συναισθηματικού περιεχομένου. Ένας από τους πρώτους γενικής χρήσης λεξικολογικούς πόρους ανάλυσης κειμένου είναι ο

General Inquirer (Stone et al., 1966), ο οποίος αναπτύχθηκε από την IBM και περιέχει 11.788 λέξεις που επισημαίνονται με 182 ετικέτες κατηγοριών, συμπεριλαμβανομένων λέξεων θετικής και αρνητικής πολικότητας. Ο λεξικολογικός πόρος General Inquirer στηρίζεται σε ψυχολογικά λεξικά όρων του Harvard, τα οποία συσχετίζονται με καταστάσεις, κίνητρα, κοινωνικές και πολιτιστικές παραμέτρους καθώς και διάφορες πτυχές συναισθηματικής δυσφορίας.

Ο λεξικολογικός πόρος ANEW (Affective Norms for English Words) (Bradley & Lang, 1999) παρέχει ένα σύνολο κοινωνικοποιημένων συναισθηματικών αξιολογήσεων για ένα μεγάλο πλήθος λέξεων της αγγλικής γλώσσας, τις οποίες χαρακτηρίζει ως προς την ένταση της ευχαρίστησης (pleasure), της διέγερσης (arousal) και της κυριαρχίας (dominance). Πιο συγκεκριμένα, οι λέξεις έχουν χαρακτηριστεί ως προς τον βαθμό της ευχαρίστησης (από ευχάριστες έως δυσάρεστες), της διέγερσης (από ηρεμία έως τον ενθουσιασμό) και της κυριαρχίας (από τον έλεγχο έως την κατάσταση εκτός ελέγχου). Με αυτό τον τρόπο μπορεί να συνεισφέρει και να βοηθήσει συστήματα και προσεγγίσεις ανάλυσης κειμένου και προσδιορισμού συναισθηματικών καταστάσεων σε αυτό.

Ένα πολύ δημοφιλές και ευρέως χρησιμοποιούμενο λεξικό είναι το WordNet Affect (Strapparava & Valitutti, 2004), το οποίο βασίζεται στο γνωστό λεξικό WordNet και το οποίο επεκτείνεται με την προσθήκη νέου υποσυνόλου από σύνολα συνωνύμων (synsets), τα οποία είναι κατάλληλα για την αναπαράσταση συναισθηματικών εννοιών. Αυτά τα σύνολα συνωνύμων (synsets) συνδέονται με μια ή περισσότερες συναισθηματικές ετικέτες. Στο WordNet βασίζεται επίσης και ο λεξικολογικός πόρος SentiWordNet 3.0 (Baccianella et al., 2010), ο οποίος συνδέει κάθε synset με τρεις αριθμητικές βαθμολογίες, όπου ο καθένας δείχνει το βαθμό που το synset είναι αντικειμενικό (objective), θετικό (positive) ή αρνητικό (negative). Το SentiWordNet 3.0 περιλαμβάνει περίπου 200.000 καταχωρήσεις και χρησιμοποιεί μια ημι-επιβλεπόμενη (semi-supervised) προσέγγιση για να χαρακτηρίσει τον θετικό - αρνητικό βαθμό κάθε λέξης.

5.2.2 Σχετικές Εργασίες

Τα τελευταία χρόνια η ανάλυση κειμένου, η αναγνώριση της ανθρώπινης συμπεριφοράς και η αναγνώριση συναισθημάτων σε κείμενα έχουν προσελκύσει την προσοχή των ερευνητών. Στη βιβλιογραφία, υπάρχουν πολλές εργασίες που αφορούν στο σχεδιασμό μεθόδων και στην ανάπτυξη συστημάτων για την ανάλυση κειμένων και τον προσδιορισμό του συναισθηματικού περιεχομένου τους (Medhat et al., 2014; Liu & Zhang, 2012; Vinodhini & Chandrasekaran, 2012).

Ένα μεγάλο πλήθος εργασιών βασίζεται στην εκπαίδευση αλγορίθμων μηχανικής μάθησης για την επίλυση της αναγνώρισης συναισθημάτων ως ένα πρόβλημα ταξινόμησης κειμένων. Σε μια αρχική εργασία (Alm et al., 2005), οι συγγραφείς αρχικά σχεδίασαν μια συλλογή κειμένων από παραμύθια, τα οποία εμπλουτίστηκαν από ειδικούς με συναισθηματικές πληροφορίες. Στη συνέχεια, βασιζόμενοι σε προσεγγίσεις μηχανικής μάθησης πέτυχαν ακρίβεια 64% στην αναγνώριση της ύπαρξης συναισθηματικού περιεχομένου στο κείμενο. Στην εργασία (Brilis et al., 2012), οι συγγραφείς εξετάζουν προσεγγίσεις μηχανικής μάθησης για την αναγνώριση στίχων τραγουδιών. Πιο συγκεκριμένα, δεδομένα κειμένου αρχικά υποβάλλονται σε βήματα προ-επεξεργασίας όπου γίνεται απομάκρυνση των λέξεων τερματισμού (stop words) και μετατροπή των λέξεων στο κύριο λήμμα τους. Τα κειμενικά δεδομένα αναπαρίστανται χρησιμοποιώντας την προσέγγιση Bag-of-Words (BoW) και επίσης κάθε λέξη συνδέεται και με την τιμή της TF-IDF (Term Frequency-Indirect Document Frequency) και στην συνέχεια ταξινομητές μηχανικής μάθησης εκπαιδεύονται και χρησιμοποιούνται για τον προσδιορισμό της διάθεσης (θετική, ουδέτερη, αρνητική διάθεση) των στίχων. Από τα αποτελέσματα της εργασίας προκύπτει ότι ο αλγόριθμος Random Forest παρουσίασε τα καλύτερα αποτελέσματα με 71.5 % ακρίβεια (accuracy). Στην εργασία (Danisman & Alpkoçak, 2008) οι συγγραφείς παρουσιάζουν μια προσέγγιση που βασίζεται στο μοντέλο vector space για την αναγνώριση του συναισθηματικού περιεχομένου του κειμένου. Πιο συγκεκριμένα, οι συγγραφείς κάνουν χρήση των συνόλων δεδομένων του ISEAR (International Survey on Emotion Antecedents and Reaction) και του SemEval και η προσέγγιση τους επικεντρώνεται στην αναγνώριση πέντε βασικών συναισθημάτων. Η εφαρμογή της σε συνολικά 801 τίτλους ειδήσεων επιτυγχάνει F-measure 32.22%. Στην εργασία (Ho & Cao, 2012), οι συγγραφείς παρουσιάζουν μια προσέγγιση που βασίζεται στη λογική Hidden Markov με στόχο να καθοριστεί το πιο πιθανό συναίσθημα για ένα συγκεκριμένο κείμενο. Τα αποτελέσματα της εργασίας αναφέρουν 35% ως τιμή του F-measure στο σύνολο δεδομένων του ISEAR. χαμηλή ακρίβεια που προέκυψε οφείλεται κυρίως στο γεγονός ότι αγνοήθηκαν σημασιολογικά και συντακτικά χαρακτηριστικά της ανάλυσης της πρότασης.

Αρκετές εργασίες για την αναγνώριση συναισθημάτων σε κείμενα χρησιμοποιούν τις πληροφορίες που περιλαμβάνονται σε συναισθηματικά λεξικά είτε αποκλειστικά είτε σε συνδυασμό με άλλα χαρακτηριστικά των προτάσεων. Στην εργασία (Osherenko & Andre, 2007), οι συγγραφείς αντιμετωπίζουν την αναγνώριση συναισθημάτων σε κείμενο βασιζόμενοι στη συναισθηματική πληροφορία των συναισθηματικών λέξεων τις οποίες αναγνωρίζουν χρησιμοποιώντας λεξικολογικούς πόρους. Στην συνέχεια, η ταξινόμηση του

κειμένου γίνεται με βάση το πλήθος των συναισθηματικών λέξεων που αναγνωρίστηκαν. Ένα κύριο συμπέρασμα είναι ότι η χρήση συναισθηματικά σχολιασμένων λεξικολογικών πόρων μπορεί να αποτελέσει ένα καλό τρόπο για τη μείωση του αριθμού των παραμέτρων για επιτυχημένη ταξινόμηση. Στην εργασία (Chaumartin, 2007), παρουσιάζεται μια προσέγγιση που βασίζεται στο λεξιολογικό πόρο SentiWordNet, ο οποίος είναι ένα υποσύνολο του WordNet-Affect, με στόχο να ανιχνευθούν θετικές και αρνητικές λέξεις, οι οποίες μπορούν να τροποποιήσουν την ένταση του συναισθηματικού περιεχομένου των προτάσεων. Από τα αποτελέσματα της εργασίας προκύπτει ότι η προσέγγιση έχει ακρίβεια 90% στην αναγνώριση της ύπαρξης συναισθημάτων σε τίτλους ειδήσεων. Οι συγγραφείς της εργασίας (Ptaszynski et al., 2013) παρουσιάζουν το σύστημα ML-Ask (eMotiveElement and Expression Analysis system), που βασίζεται στην αναγνώριση λέξεων-κλειδιών με χρήση κατάλληλων λεξικολογικών πόρων και είναι σε θέση να προσδιορίσει ύπαρξη των συναισθημάτων της χαράς, της αγάπης, της ανακούφισης, του φόβου, της θλίψης και της οργής σε κείμενα με ακρίβεια 60%. Επίσης, η προσέγγιση των συγγραφέων είναι σε θέση να προσδιορίσει αν μια πρόταση είναι συναισθηματική ή όχι με ακρίβεια της τάξης του 90%. Οι συγγραφείς της εργασίας (Neviarouskaya et al., 2007) έχουν ως στόχο να αναγνωρίσουν τα έξι βασικά συναισθήματα του Ekman (Ekman, 1999) σε online ιστολόγια. Για το σκοπό αυτό χρησιμοποιούν το συντακτικό αναλυτή Machine, προκειμένου να αναλύσουν συντακτικά τις προτάσεις και να εξάγουν συντακτικές πληροφορίες και χρησιμοποιούν κανόνες για τον προσδιορισμό του συναισθηματικού περιεχομένου μιας πρότασης. Το σύστημα που αναπτύχθηκε έχει ακρίβεια 70% στην αναγνώριση του συναισθηματικού περιεχομένου προτάσεων.

Πρόσφατες προσεγγίσεις ομαδοποιημένων ταξινομητών (ensemble classifiers) έχουν προσελκύσει το ενδιαφέρον της έρευνας και έχει μελετηθεί η απόδοσή τους σε διάφορα προβλήματα ανάλυσης και ταξινόμησης κειμένου. Στην εργασία (da Silva et al., 2014), γίνεται διερεύνηση της ανάλυσης του συναισθηματικού περιεχομένου tweets των χρηστών στο κοινωνικό δίκτυο twitter χρησιμοποιώντας ομαδοποιημένο ταξινομητή. Πιο συγκεκριμένα, ο ομαδοποιημένος ταξινομητής σχηματίζεται από τους εξής βασικούς ταξινομητές μηχανικής μάθησης: random forest, support vector machines, multinomial naïve Bayes και logistic regression. Τα πειραματικά αποτελέσματα σε μια ποικιλία από σύνολα δεδομένων tweets έδειξαν ότι οι ομαδοποιημένοι ταξινομητές μπορούν να βελτιώσουν την ακρίβεια ταξινόμησης των απλών ταξινομητών. Επίσης, έγινε σύγκριση προσεγγίσεων για την αναπαράσταση των κειμένων των tweets, όπως είναι οι αναπαράστασεις bag-of-words και η feature hashing, και προκύπτει ότι η αναπαράσταση

bag-of-words μπορεί να επιτύχει την καλύτερη ακρίβεια. Στην εργασία (Xia et al, 2011), οι συγγραφείς μελετούν ένα ομαδοποιημένο σχήμα ταξινόμησης για την ταξινόμηση συναισθημάτων, το οποίο συνδυάζει τρεις αλγόριθμους: Naïve Bayes, μέγιστης εντροπίας (Maximum Entropy) και Support Vector Machine. Το ομαδοποιημένο σχήμα αναγνωρίζει την πολικότητα (θετική ή αρνητική) του συναισθηματικού περιεχομένου του κειμένου με βάση τα μέρη του λόγου (part-of-speech features) και χαρακτηριστικά που βασίζονται στην σχέση μεταξύ λέξεων (word-relation features). Βασικό συμπέρασμα είναι ότι η χρήση ομαδοποιημένου ταξινομητή για το ίδιο σύνολο χαρακτηριστικών αποδίδει καλύτερα από τους μεμονωμένους ταξινομητές. Στην εργασία (Wang et al., 2014), οι συγγραφείς μελετούν την απόδοση ομαδοποιημένου ταξινομητή για την αναγνώριση συναισθημάτων σε κείμενα. Ο ομαδοποιημένος ταξινομητής αποτελείται από πέντε βασικούς ταξινομητές, οι οποίοι είναι: Naïve Bayes, maximum entropy, decision tree, k-nearest neighbor, support vector machine και τους οποίους συνδυάζει με την μέθοδο random subspace. Τα αποτελέσματα δείχνουν ότι ο ομαδοποιημένος ταξινομητής βελτιώνει σημαντικά την απόδοση των επιμέρους ταξινομητών και πετυχαίνει καλύτερα αποτελέσματα από ότι χρησιμοποιώντας μόνο τους επιμέρους βασικούς ταξινομητές. Από τα αποτελέσματα προκύπτει ότι οι ομαδοποιημένοι ταξινομητές έχουν την προοπτική και μπορούν να χρησιμοποιηθούν ως μια πολύ αποδοτική προσέγγιση για την αναγνώριση συναισθημάτων σε κείμενα φυσικής γλώσσας.

Οι προσεγγίσεις σχεδιασμού και ανάπτυξης ομαδοποιημένων ταξινομητών στη διεθνή βιβλιογραφία βασίζονται σε αποκλειστική χρήση ταξινομητών μηχανικής μάθησης. Ωστόσο, οι προσεγγίσεις μηχανικής μάθησης γενικά δεν λαμβάνουν υπόψη τους σημασιολογικά και συντακτικά χαρακτηριστικά της ανάλυσης των προτάσεων, γεγονός που τις κάνει ευαίσθητες σε σημασιολογικό επίπεδο. Από την άλλη πλευρά, οι μέθοδοι ταξινόμησης που βασίζονται μόνο σε λέξεις-κλειδιά και λεξικολογικούς πόρους έχουν το πρόβλημα της ασάφειας σε ορισμούς, με την έννοια ότι μια λέξη μπορεί να έχει διαφορετικές σημασίες ανάλογα με τη χρήση και τα συμφραζόμενα στην πρόταση, καθώς και αδυναμία αναγνώρισης των συναισθημάτων μέσα σε προτάσεις που δεν περιέχουν συναισθηματικές λέξεις-κλειδιά (Shaheen et al., 2014). Επομένως, με βάση τα παραπάνω, μια προσέγγιση που θα συνδυάζει τόσο προσεγγίσεις μηχανικής μάθησης όσο και προσεγγίσεις που βασίζονται στη γνώση θα μπορούσε να παρουσιάσει καλά αποτελέσματα και θα έχει μεγάλο ενδιαφέρον. Η έρευνα που έγινε στην παρούσα εργασία αποτελεί μια από τις πρώτες προσεγγίσεις στον τομέα της ανάλυσης συναισθημάτων σε κείμενο που εξετάζει αυτήν την κατεύθυνση.

5.3 Μοντελοποίηση Συναισθημάτων

Ο ορισμός του τι είναι συναίσθημα καθώς και ο τρόπος που προσδιορίζονται τα συναισθήματα είναι ένα πολυδιάστατο ερώτημα που παραμένει ανοικτό για πάνω από έναν αιώνα. Γενικά, ένα συναίσθημα θεωρείται ότι είναι μια συγκίνηση που απορρέει από την κατάσταση ενός ατόμου, την διάθεσή του και την αλληλεπίδραση με ανθρώπους και γεγονότα (Oxford Dictionary, 2008). Ο τρόπος με τον οποίο γίνεται η αναπαράσταση των συναισθημάτων, αποτελεί μια βασική παράμετρο για την κατασκευή εργαλείων και συστημάτων αναγνώρισης συναισθηματικών καταστάσεων (Reisenzein et al., 2013). Τα βασικά μοντέλα για την αναπαράσταση των συναισθημάτων διακρίνονται σε δυο βασικές κατηγορίες, οι οποίες είναι τα *κατηγορικά (categorical)* και τα *διανυσματικά (dimensional)* μοντέλα.

Τα κατηγορικά μοντέλα βασίζονται στην προσέγγιση ότι υπάρχει ένας πεπερασμένος αριθμός διακριτών συναισθημάτων, όπου το κάθε ένα είναι αυτοπροσδιοριζόμενο και διαφοροποιείται από τα υπόλοιπα. Από την άλλη πλευρά, τα διανυσματικά μοντέλα ακολουθούν μια διαφορετική προσέγγιση και αναπαριστούν τις συναισθηματικές καταστάσεις σε ένα διανυσματικό χώρο. Με βάση αυτήν την προσέγγιση, δημιουργείται ένας διανυσματικός συναισθηματικός χώρος και κάθε συναίσθημα έγκειται σε αυτό το χώρο.

Ένα πολύ δημοφιλές και ευρέως χρησιμοποιούμενο κατηγορικό μοντέλο είναι το μοντέλο που ψυχολόγου Paul Ekman (Ekman, 1999), το οποίο καθορίζει έξι βασικά ανθρώπινα συναισθήματα, τα οποία είναι ο θυμός (anger), η αηδία (disgust), ο φόβος (fear), η χαρά (joy), η θλίψη (sadness) και η έκπληξη (surprise). Αυτά τα συναισθήματα χαρακτηρίζονται ως καθολικά (universal) συναισθήματα της ανθρώπινης φύσης, καθώς εκφράζονται με τον ίδιο τρόπο από διαφορετικούς πολιτισμούς και εποχές. Το μοντέλο του Ekman έχει χρησιμοποιηθεί σε πολλές ερευνητικές μελέτες και σε διάφορα συστήματα, τα οποία αναγνωρίζουν το συναισθηματικό περιεχόμενο που φέρουν προτάσεις κειμένου καθώς και την συναισθηματική κατάσταση από εκφράσεις προσώπου. Ένα επίσης κατηγορικό μοντέλο που έχει υιοθετηθεί σε πολλές μελέτες για την αναγνώριση συναισθημάτων είναι το μοντέλο Ortony-Clore-Collins (OCC) (Ortony et al., 1988). Το μοντέλο αυτό καθορίζει 22 κατηγορίες συναισθημάτων και βασίζεται σε συναισθηματικές αντιδράσεις σε διάφορες καταστάσεις και γεγονότα. Έχει σχεδιαστεί για να μοντελοποιήσει τις συναισθητικές αντιδράσεις που προκύπτουν σε καταστάσεις και γεγονότα, έχει

καθιερωθεί ως το πρότυπο μοντέλο για τη σύνθεση συναισθημάτων και χρησιμοποιείται κυρίως σε συστήματα που ενθέτουν συναισθηματικές καταστάσεις σε ευφυής πράκτορες.

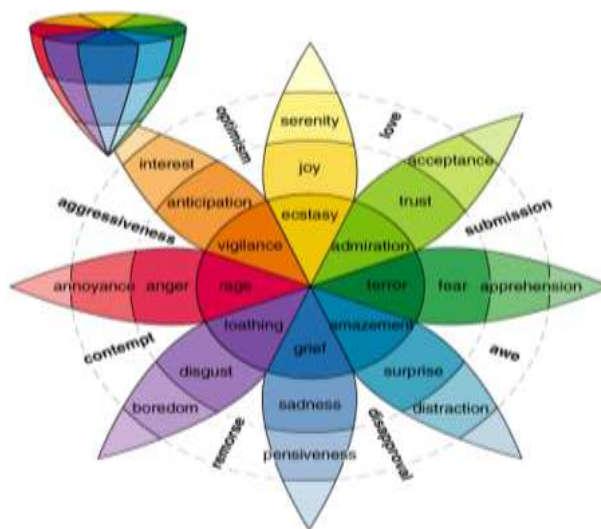
Το μοντέλο αναπαράστασης των συναισθημάτων Parrot, το οποίο εισήχθη από τον ψυχολόγο Gerrod Parrott (Parrot, 2001), ακολουθεί μια δενδρική κατηγορική αναπαράσταση των συναισθημάτων και αποτελείται από τρία βασικά επίπεδα. Το πρώτο επίπεδο του μοντέλου αποτελείται από έξι συναισθήματα τα οποία είναι: αγάπη (love), χαρά (joy), έκπληξη (surprise), θυμός (anger), θλίψη (sadness) και φόβος (fear) και κάθε επόμενο επίπεδο βελτιώνει την ευαισθησία του προηγούμενου επιπέδου κάνοντας τις βασικές αυτές συναισθηματικές καταστάσεις πιο λεπτομερείς. Το μοντέλο του Parrot είναι πολύ αναλυτικό και προσδιορίζει περισσότερα από 100 συναισθήματα και θεωρείται ότι είναι η πιο λεπτομερής ταξινόμηση συναισθημάτων όπως παρουσιάζεται στην Εικόνα 5.2

Primary emotion	Secondary emotion	Tertiary emotions
Love	Affection	Adoration, affection, love, fondness, liking, attraction, caring, tenderness, compassion, sentimentally
	Lust/Longing	Arousal, desire, lust, passion, infatuation, longing
Joy	Cheerfulness	Amusement, bliss, cheerfulness, gaiety, glee, jolliness, joviality, joy, delight, enjoyment, gladness, happiness, jubilation, elation, satisfaction, ecstasy, euphoria
	Zest/Contentment	Enthusiasm, zeal, zest, excitement, thrill, exhilaration, Contentment, pleasure
	Pride	Pride, triumph
	Optimism	Eagerness, hope, optimism
	Enthrallment	Enthrallment, rapture
	Relief	Relief
Surprise	Surprise	Amusement, surprise, astonishment
Anger	Irritation	Aggravation, irritation, agitation, annoyance, grouches, grumpiness
	Exasperation	Exasperation, frustration
	Rage	Anger, rage, outrage, fury, wrath, hostility, ferocity, bitterness, hate, loathing, scorn, spite, vengefulness, dislike, resentment
	Disgust	Disgust, revulsion, contempt
	Envy	Envy, jealousy
	Torment	Torment
Sadness	Suffering	Agony, suffering, hurt, anguish
	Sadness	Depression, despair, hopelessness, gloom, gloominess, sadness, unhappiness, grief, sorrow, woe, misery, melancholy
	Disappointment	Disarray, disappointment, displeasure
	Guilt	Guilt, shame, regret, remorse
	Neglect	Alienation, isolation, neglect, loneliness, rejection, homesickness, defeat, dejection, insecurity, embarrassment, humiliation, insult
	Sympathy	Pity, sympathy
Fear	Horror	Alarm, shock, fear, fright, horror, terror, panic, hysteria, mortification
	Nervousness	Anxiety, nervousness, timidity, uneasiness, apprehension, worry, distress, dread

Εικόνα 5.2 Επίπεδα και συναισθηματικές καταστάσεις του μοντέλου Parrot.

Το μοντέλο συναισθημάτων Plutchik (Plutchik, 2001) είναι ένα διανυσματικό μοντέλο το οποίο βασίζεται στην εξελικτική αρχή και ορίζει οκτώ βασικά συναισθήματα, όπως παρουσιάζεται στην Εικόνα 5.3. Αυτά τα οκτώ συναισθήματα οργανώνονται σε τέσσερις διπολικές ομάδες : χαρά - θλίψης, θυμός -φόβος, εμπιστοσύνη -αηδία, έκπληξη -αναμονή. Κάθε συναίσθημα μπορεί να διαιρεθεί περαιτέρω σε τρεις βαθμούς. Για παράδειγμα, η γαλήνη είναι ένα μικρότερου βαθμού συναίσθημα της χαράς και η έκσταση

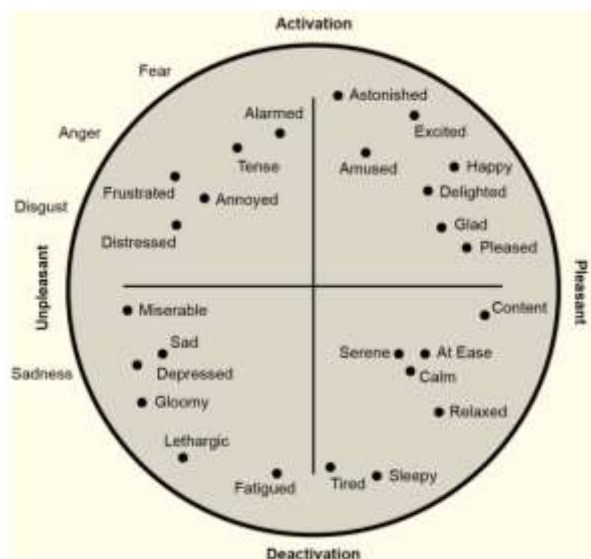
είναι ένα πιο έντονο συναίσθημα της χαράς. Επίσης, τα οκτώ βασικά συναισθήματα μπορούν να συνδυαστούν για να σχηματίσουν άλλες συναισθηματικές καταστάσεις. Για παράδειγμα, η χαρά και η εμπιστοσύνη μπορούν να συνδυαστούν για να σχηματίσουν την αγάπη. Το μοντέλο ονομάζεται και τροχός των συναισθημάτων διότι απέδειξε ότι διαφορετικά συναισθήματα μπορούν να συνδυαστούν μεταξύ και να δημιουργήσουν νέες συναισθηματικές καταστάσεις.



Εικόνα 5.3 Το συναισθηματικό μοντέλου του Plutchik¹.

Ένα ακόμη διανυσματικό μοντέλο είναι το συναισθηματικό μοντέλο circumplex (Russell, 1980), που εισήχθη από τον ψυχολόγο James Russell και με βάση το οποίο τα συναισθήματα αναπαρίστανται σε ένα δισδιάστατο κυκλικό χώρο όπως παρουσιάζεται στην Εικόνα 5.4. Η μία διάσταση του χώρου αντιπροσωπεύει την διέγερση (activation) και η άλλη διάσταση την πολικότητα (polarity). Η διάσταση της πολικότητας χαρακτηρίζει ένα συναίσθημα ως θετικό ή αρνητικό, ενώ η διέγερση χαρακτηρίζει ένα συναίσθημα ως ενεργό ή ανενεργό. Το κέντρο του κύκλου αντιπροσωπεύει ένα ουδέτερου σθένους με μέσο επίπεδο διέγερσης συναίσθημα. Σε αυτό το μοντέλο, οι συναισθηματικές καταστάσεις μπορούν να αναπαρίστανται σε οποιοδήποτε επίπεδο του σθένους και της πολικότητας, ή σε ουδέτερο επίπεδο του ενός ή και των δύο από τα δυο επίπεδα. Το μοντέλο circumplex του Russell χρησιμοποιείται συχνά για τη αποτίμηση του συναισθήματος των λέξεων, καθώς και της ανάλυσης του συναισθηματικού περιεχομένου εκφράσεων του προσώπου.

¹ Plutchik 2001



Εικόνα 5.4 Το Circumflex μοντέλου του Russell.

5.4 Σχεδίαση και ανάπτυξη ομαδοποιημένου ταξινομητή αναγνώρισης συναισθηματικού περιεχομένου κειμένου

5.4.1 Βασικές αρχές σχεδιασμού

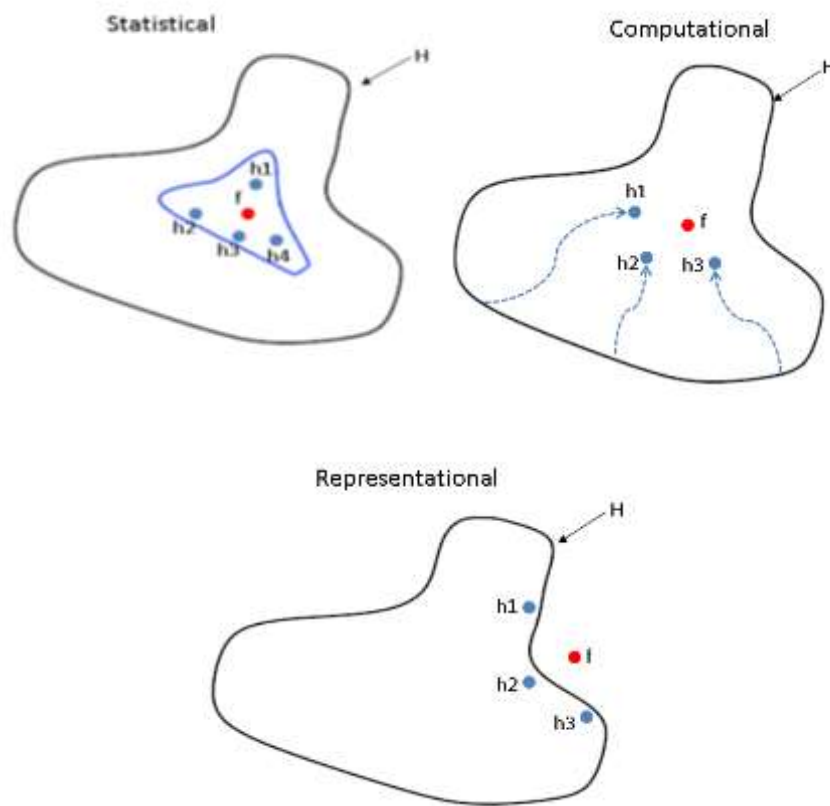
Οι ομαδοποιημένοι ταξινομητές (ensemble classifiers) είναι μια αποτελεσματική μέθοδος για τη βελτίωση της απόδοσης συστημάτων ταξινόμησης (Li et al., 2007). Ο σκοπός της ομαδοποίησης ταξινομητών είναι το σχήμα ταξινόμησης που προκύπτει, να επωφεληθεί και να εκμεταλλευτεί τα πλεονεκτήματα των επιμέρους εκπαιδευμένων ταξινομητών καθώς και να περιορίζει τα μειονεκτήματά τους. Ο σχεδιασμός και η ανάπτυξη αποτελεσματικών σχημάτων ομαδοποιημένων ταξινομητών προϋποθέτει οι επιμέρους ταξινομητές να έχουν κάποιο επίπεδο διαφορετικότητας. Στη διεθνή βιβλιογραφία, σχήματα ομαδοποιημένων ταξινομητών έχουν εφαρμοστεί με επιτυχία σε διάφορους τομείς της εξόρυξης γνώσης από κείμενα, όπως η ταξινόμηση κειμένου, η αποσαφήνιση της σημασίας λέξεων σε προτάσεις, η αναγνώριση του είδους των κύριων ονομάτων και των οντοτήτων σε προτάσεις κ.α. (Xia et al., 2011). Σε γενικές γραμμές, οι μέθοδοι ομαδοποίησης ταξινομητών βασίζονται σε ένα σύνολο βασικών ταξινομητών τους οποίους συνδυάζουν για να λάβουν μια απόφαση. Σε αντίθεση με τις κλασικές μεθόδους μηχανικής μάθησης, οι οποίες εκπαιδεύουν και χρησιμοποιούν μια απλή μέθοδο ταξινόμησης, τα σχήματα ομαδοποίησης ταξινομητών εκπαιδεύουν ποικιλοτρόπως και χρησιμοποιούν πολλαπλούς διαφορετικούς ταξινομητές.

Υπάρχουν πολλοί λόγοι για το σχεδιασμό, την ανάπτυξη και τη χρήση ομαδοποιημένων ταξινομητών. Γενικά, ένας αλγόριθμος μάθησης μπορεί να θεωρηθεί σαν ένα πρόβλημα αναζήτησης στο χώρο υποθέσεων H με στόχο να βρει την καλύτερη υπόθεση μέσα στον χώρο. Από την *στατιστική* σκοπιά (statistical scope), όταν το σύνολο δεδομένων είναι σχετικά μικρό σε σχέση με τον χώρο υποθέσεων, ο αλγόριθμος μάθησης μπορεί να βρει διαφορετικές υποθέσεις στον χώρο H που να έχουν την ίδια ακρίβεια για το σύνολο δεδομένων εκπαίδευσης (Dietterich, 2000). Επομένως, κατασκευάζοντας ένα ομαδοποιημένο ταξινομητή από βασικούς ταξινομητές, ο αλγόριθμος προσδιορισμού της τελικής απόφασης λαμβάνει υπόψη του τα δεδομένα απόφασης όλων των ταξινομητών μειώνοντας έτσι την πιθανότητα λάθους επιλογής και χαμηλής επίδοσης σε νέα δεδομένα. Στην Εικόνα 5.5 παρουσιάζεται η εσωτερική περιοχή που δείχνει τον υποχώρο, ο οποίος περικλείει όλες τις υποθέσεις που έχουν πολύ καλή ακρίβεια, και η f αναπαριστά την καλύτερη πραγματική συνάρτηση, η οποία μπορεί να προσεγγιστεί καλύτερα μέσω τις ομαδοποίησης των υποθέσεων που έγιναν (h_1 - h_4). Όταν διαφορετικοί ταξινομητές εκπαιδεύονται, ακόμη και στην περίπτωση που ο καθένας τους έχει πολύ καλή απόδοση, η επιλογή ενός μόνο ταξινομητή μπορεί να μην επιφέρει την καλύτερη απόδοση και γενίκευση σε νέα δεδομένα.

Από την *υπολογιστική* σκοπιά (computational scope), υπάρχουν αρκετοί αλγόριθμοι μάθησης που λειτουργούν εκτελώντας κάποιο είδος τοπικής αναζήτησης και είναι πολύ πιθανό η αναζήτηση να εγκλωβιστεί σε ένα τοπικό βέλτιστο (Dietterich, 2000). Για παράδειγμα, τα νευρωνικά δίκτυα μπορούν να εφαρμόζουν μεθόδους, όπως η βαθμιαία κάθοδος (gradient descent), για να ελαχιστοποιήσουν την συνάρτηση λάθους στο σύνολο δεδομένων και τα δέντρα αποφάσεων να ακολουθούν άπληστες (greedy) στρατηγικές για να δημιουργήσουν το δέντρο απόφασης παίρνοντας μια σειρά από τοπικά βέλτιστες αποφάσεις. Η αναζήτηση και η ανεύρεση της καλύτερης υπόθεσης μπορεί να είναι εξαιρετικά πολύπλοκη διαδικασία για τον αλγόριθμο και η βέλτιστη εκπαίδευση των νευρωνικών δικτύων και των δέντρων αποφάσεων αποτελεί ένα πρόβλημα NP-hard (Hyafil & Rivest, 1976; Blum & Rivest, 1992). Επομένως, ένας ομαδοποιημένος ταξινομητής, ο οποίος εκτελεί την λειτουργία της τοπικής αναζήτησης με διαφορετικά σημεία εκκίνησης, μπορεί να παρέχει μια καλύτερη προσέγγιση ως προς την καλύτερη πραγματική συνάρτηση f σε σχέση με τους επιμέρους ταξινομητές.

Τέλος, από την σκοπιά της αναπαράστασης (representational scope), σε ορισμένες περιπτώσεις η πραγματική συνάρτηση f ενδέχεται να μην μπορεί να αναπαρασταθεί από τις υποθέσεις του χώρου υποθέσεων H . Επομένως, συνθέτοντας τις υποθέσεις που έγιναν στον

χώρο H μπορεί να είναι εφικτή η επέκταση του χώρου, όπως απεικονίζεται στην Εικόνα 5.5 (Dietterich, 2000).



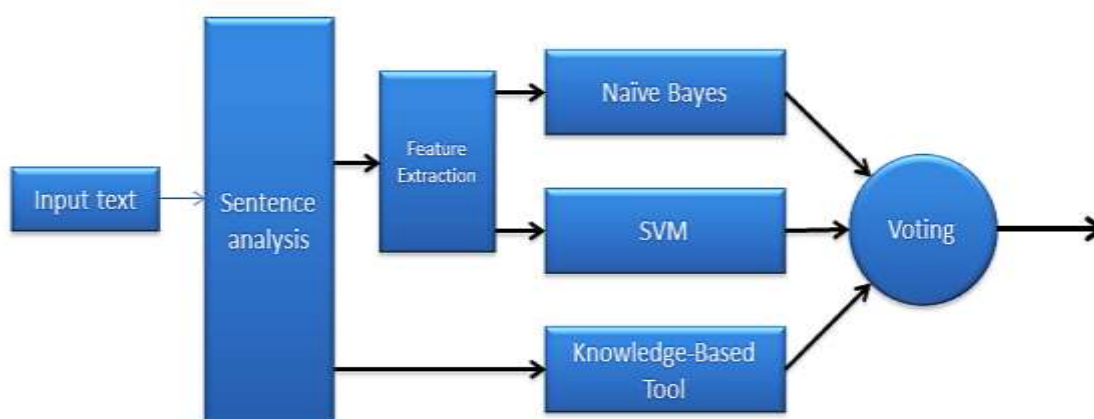
Εικόνα 5.5 Βασικές σκοπιές ομαδοποίησης ταξινομητών.

Επίσης, λαμβάνοντας υπόψη τις αρχές λειτουργίας και τα πλεονεκτήματα των ομαδοποιημένων ταξινομητών και με βάση την ιδιαιτερότητα του πεδίου της ταξινόμησης κειμένων και των χαρακτηριστικών των προτάσεων φυσικής γλώσσας, η χρησιμοποίηση ομαδοποιημένων ταξινομητών είναι μια κατάλληλη και ενδιαφέρουσα προσέγγιση. Η εργασία μας στην δημιουργία ομαδοποιημένου ταξινομητή που συνδυάζει μεθόδους που βασίζονται στην επεξεργασία και ανάλυση της φυσικής γλώσσας καθώς και ταξινομητές μηχανικής μάθησης αποτελεί μια πρώτη καινοτόμα εργασία στην διεθνή βιβλιογραφία.

5.4.2 Ομαδοποιημένος Ταξινομητής Αναγνώρισης Συναισθημάτων

Ο ομαδοποιημένος ταξινομητής που σχεδιάστηκε και αναπτύχθηκε βασίζεται σε τρεις επιμέρους ταξινομητές, δύο εκ των οποίων βασίζονται σε τεχνικές μηχανικής μάθησης και ο τρίτος ακολουθεί μια προσέγγιση βασισμένη στην γνώση. Πιο συγκεκριμένα, οι δυο ταξινομητές μηχανικής μάθησης που χρησιμοποιήθηκαν είναι ένας Naïve Bayes και ένας

Support Vector Machine, οι οποίοι εκπαιδεύτηκαν χρησιμοποιώντας συλλογές σχολιασμένων κειμένων και προτάσεων όπως η International Survey on Emotion Antecedents and Reaction (ISEAR) (Scherer & Wallbott, 1994) και η Affective text (Strapparava & Mihalcea, 2007). Ο τρίτος ταξινομητής, που ακολουθεί μια προσέγγιση βασισμένη στην γνώση, πραγματοποιεί μια εις βάθος ανάλυση των προτάσεων φυσικής γλώσσας προκειμένου να αναγνωρίσει λέξεις που φέρουν συναισθηματικό περιεχόμενο, να αναλύσει συντακτικές δομές και να προσδιορίσει την συναισθηματική πληροφορία της κάθε πρότασης. Για την μορφοσυντακτική ανάλυση των προτάσεων χρησιμοποιούνται εργαλεία επεξεργασίας φυσικής γλώσσας όπως ο Stanford parser (De Marneffe et al., 2006) και λεξικολογικοί πόροι όπως το WordNet Affect (Strapparava & Valitutti, 2004) για τον προσδιορισμό λέξεων που φέρουν συναισθηματική πληροφορία. Στην συνέχεια, με βάση τις συντακτικές δομές της πρότασης και τις συναισθηματικές λέξεις γίνεται αρχικά προσδιορισμός του βαθμού της έντασης των συναισθημάτων των λέξεων και στην συνέχεια ο προσδιορισμός των συναισθημάτων που περιέχονται στην κάθε πρόταση. Η βασική αρχιτεκτονική του ένθετου ταξινομητή παρουσιάζεται στην Εικόνα 5.6.



Εικόνα 5.6 Βασική αρχιτεκτονική του ομαδοποιημένου ταξινομητή.

Ο ομαδοποιημένος ταξινομητής πραγματοποιεί αναγνώριση συναισθηματικής πληροφορίας σε επίπεδο πρότασης. Με αυτό τον τρόπο, δοθέντος ενός νέου κειμένου, το κείμενο αρχικά χωρίζεται σε προτάσεις και στην συνέχεια κάθε πρόταση αναλύεται αυτόνομα και εξάγονται κατάλληλα χαρακτηριστικά. Ο ομαδοποιημένος ταξινομητής καθορίζει αν η πρόταση περιέχει συναισθηματική πληροφορία ή είναι συναισθηματικά ουδέτερη, και σε περίπτωση που είναι συναισθηματική, καθορίζεται η πολικότητα του συναισθηματικού περιεχομένου της. Μελετήθηκε η απόδοση στην περίπτωση που οι ταξινομητές ομαδοποιούνται σύμφωνα με την αρχή της πλειοψηφίας καθώς και στις

περιπτώσει που οι ταξινομητές μηχανικής μάθησης έχουν εκπαιδευτεί και ομαδοποιηθεί με την μέθοδο bagging και boosting.

5.4.3 Εξαγωγή Χαρακτηριστικών από Κείμενα

Ο τρόπος με τον οποίο αναπαρίσταται ένα κείμενο και το είδος των χαρακτηριστικών που εξάγονται από αυτό αποτελεί ένα βασικό παράγοντα της λειτουργίας των ταξινομητών μηχανικής μάθησης. Στα πλαίσια του τρόπου λειτουργίας του ομαδοποιημένου ταξινομητή, μελετήθηκαν και χρησιμοποιήθηκαν η τεχνική αναπαράστασης Bag of Words (BOW) και η bigrams. Οι τεχνικές αυτές αποτελούν τις πιο βασικές τεχνικές αναπαράστασης κειμένου και χρησιμοποιούνται ευρέως σε διαδικασίες ταξινόμησης κειμένων. Μια νέα πρόταση στο σύστημα αρχικά σπάει σε λέξεις και κάθε λέξη μετατρέπεται στην βασική της μορφή (λήμμα). Επίσης, συνδετικές λέξεις (stop-words), οι οποίες δεν φέρουν κάποια πληροφορία και προσθέτουν μόνο θόρυβο, δεν λαμβάνονται υπόψη. Τα χαρακτηριστικά που εξάγονται από την κάθε πρόταση οδηγούνται στους ταξινομητές μηχανικής μάθησης.

5.4.4 Ταξινομητής Naïve Bayes

Ο ταξινομητής Naive Bayes είναι ένας απλός πιθανοτικός ταξινομητής που βασίζεται στην εφαρμογή του θεωρήματος Bayes και στόχος του είναι να ταξινομήσει ένα δείγμα X σε μια από τις προκαθορισμένες κατηγορίες. Ο στόχος αυτός είναι ισοδύναμος με το να αναθέσει, ανάλογα με την κατηγορία στην οποία ταξινομείται, μια ετικέτα (label) στο δείγμα X . Ο όρος Naive αναφέρεται στο γεγονός ότι η λειτουργία του ταξινομητή στηρίζεται σε μια ισχυρή παραδοχή ανεξαρτησίας. Η ισχυρή παραδοχή είναι ότι η παρουσία (ή μη) ενός χαρακτηριστικού σε μια κλάση είναι ανεξάρτητη από την παρουσία (ή μη) ενός άλλου. Γενικά, θεωρεί ότι οι λέξεις είναι ανεξάρτητες μεταξύ τους και κάθε μια είναι ένας παράγοντας που συμβάλει στην ανάθεση σε μια συναισθηματική κατηγορία. Με βάση τα χαρακτηριστικά της πρότασης, γίνεται υπολογισμός της πιθανότητας να ανήκει σε μια κατηγορία ως εξής:

$$P(c|s) = \frac{P(c)P(s|c)}{P(s)}$$

όπου $P(c)$ είναι η πιθανότητα να ανήκει μια πρόταση στην κατηγορία c και εκφράζει την ισχύ της κατηγορίας c χωρίς να έχει προηγηθεί παρατήρηση των δεδομένων της πρότασης. Η πιθανότητα αυτή ονομάζεται εκ των προτέρων πιθανότητα της c (a-priori probability). Η $P(s)$ είναι η πιθανότητα ύπαρξης της πρότασης s , δηλαδή είναι η εκ των προτέρων γνωστή πιθανότητα παρατήρησης των δεδομένων της πρότασης. Η $P(s|c)$ είναι η δεσμευμένη πιθανότητα που εκφράζει το ενδεχόμενο παρατήρησης της πρότασης s δεδομένης της

κατηγορίας c και η $P(c|s)$ είναι η πιθανότητα δεδομένου ότι παρατηρήθηκε η πρόταση s αυτή να ανήκει στην κατηγορία c . Η $P(s|c)$ μπορεί να υπολογιστεί λαμβάνοντας υπόψη τις δεσμευμένες πιθανότητες παρατήρησης των όρων της πρότασης δεδομένης της κατηγορίας c ως εξής:

$$P(s|c) = \prod_{1 \leq k \leq n} P(s_k | c)$$

όπου $P(s_k|c)$ αναπαριστά την πιθανότητα ο όρος (λέξη) s_k να συμβεί δεδομένης της κατηγορίας c και το n αναπαριστά το μήκος της πρότασης s .

5.4.5 Μηχανές Διανυσμάτων Υποστήριξης – SVM ταξινομητής

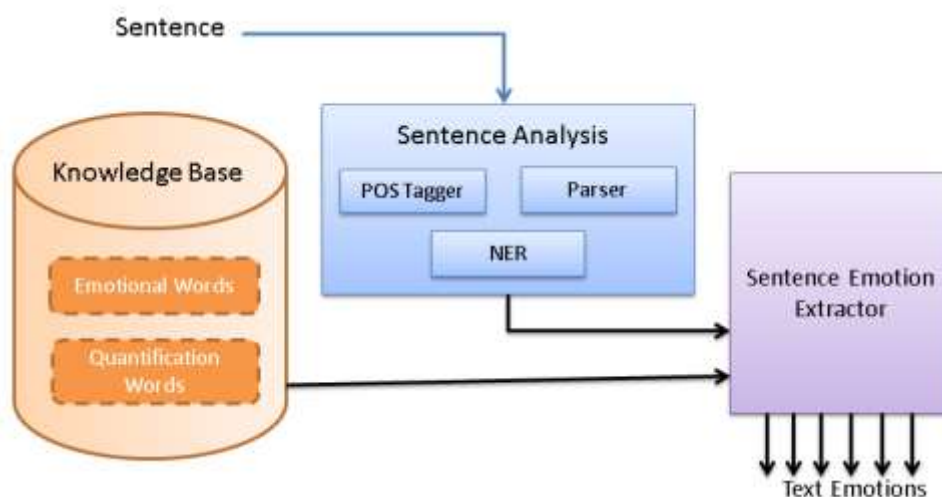
Οι μηχανές διανυσμάτων υποστήριξης (Support Vector Machines- SVMs) (Cortes & Vapnik, 1990) στηρίζονται στη θεωρία της στατιστικής μάθησης και στα νευρωνικά δίκτυα και αποτελούν μια από τις πιο διαδεδομένες μεθόδους (γραμμικής και μη) ταξινόμησης, με πλήθος εφαρμογών σε πεδία όπως, η αναγνώριση γραφής (handwriting recognition), η ταξινόμηση δεδομένων, η ταξινόμηση κειμένων (text categorization) κ.α. Οι μηχανές διανυσμάτων υποστήριξης για την δημιουργία ενός μοντέλου ταξινόμησης προσπαθούν να βρουν μια υπερεπιφάνεια (hypersurface), που να διαχωρίζει στο χώρο τα στιγμιότυπα των δύο κλάσεων. Η υπερεπιφάνεια αυτή επιλέγεται ώστε να απέχει όσο το δυνατόν περισσότερο από τα κοντινότερα στιγμιότυπα των δύο κλάσεων (maximum margin hypersurface). Από γεωμετρική σκοπιά, η προσέγγιση των μηχανών διανυσμάτων υποστήριξης προσπαθεί να βρει στον $|T|$ -διάστατο χώρο, ανάμεσα από τα διαχωριστικά υπερεπίπεδα $h_1, h_2, h_3 \dots h_n$, εκείνο το υπερεπίπεδο που διαχωρίζει τα στιγμιότυπα των κλάσεων και έχει το μεγαλύτερο εύρος ανάμεσά τους.

Στην περίπτωση που τα δεδομένα είναι γραμμικά διαχωρίσιμα, δηλαδή υπάρχει ένα τέτοιο διαχωριστικό υπερεπίπεδο h_i , η μέθοδος των μηχανών διανυσμάτων υποστήριξης επιλέγει το ζεύγος παράλληλων διαχωριστικών επιπέδων μεταξύ των κλάσεων που έχει το μεγαλύτερο εύρος ανάμεσα τους και στην συνέχεια επιλέγει ως τελικό διαχωριστικό υπερεπίπεδο τη διχοτόμο τους. Στην περίπτωση που τα δεδομένα δεν είναι γραμμικώς διαχωρίσιμα, δηλαδή δεν υπάρχει διαχωριστικό υπερεπίπεδο που να διαχωρίζει πλήρως τα δεδομένα στον διανυσματικό χώρο $|T|$ -διαστάσεων, η μέθοδος των μηχανών διανυσμάτων υποστήριξης αναζητά ενδεχόμενο διαχωριστικό υπερεπίπεδο σε έναν άλλο διανυσματικό χώρο υψηλότερων διαστάσεων $|T'|$, με $|T'| > |T|$ στον οποίο μεταβαίνει με την τεχνική του τεχνάσματος πυρήνων (kernel trick).

Ένα βασικό πλεονέκτημα των μηχανών διανυσμάτων υποστήριξης είναι η ικανότητά τους να χειρίζονται πολύ μεγάλους χώρους χαρακτηριστικών, καθώς και η ανεκτικότητα που παρουσιάζουν όσον αφορά το πλήθος των στιγμιότυπων εκπαίδευσης, ιδιαίτερα μάλιστα όταν υπάρχει διαφορά μεταξύ του πλήθους των στιγμιότυπων των δύο κλάσεων, κάτι που οφείλεται στο γεγονός ότι οι μηχανές διανυσμάτων υποστήριξης δεν επιδιώκουν να ελαχιστοποιήσουν το σφάλμα των δεδομένων εκπαίδευσης, αλλά να τα διαχωρίσουν αποτελεσματικά σε ένα χώρο μεγάλης διάστασης.

5.4.6 Μηχανισμός Βασισμένος στην Γνώση (Knowledge-Based Mechanism)

Η προσέγγιση που βασίζεται στη γνώση και ο μηχανισμός που αναπτύχθηκε, σε αντίθεση με τις στατιστικές προσεγγίσεις, προσπαθεί να αναλύσει τις συντακτικές δομές των προτάσεων προκειμένου να καθορίσει το συναισθηματικό περιεχόμενο των προτάσεων (Perikos & Hatzilygeroudis, 2013). Η αρχιτεκτονική του μηχανισμού απεικονίζεται στην Εικόνα 5.7. Ο μηχανισμός πραγματοποιεί επεξεργασία και εις βάθος ανάλυση των προτάσεων χρησιμοποιώντας εργαλεία επεξεργασίας φυσικής γλώσσας και είναι σε θέση να αναγνωρίσει σε μια πρόταση την ύπαρξη των βασικών συναισθημάτων του μοντέλου Ekman.



Εικόνα 5.7. Η αρχιτεκτονική του βασισμένου σε γνώση μηχανισμού

Αρχικά, δοθείσας μιας πρότασης, πραγματοποιείται μορφοσυντακτική ανάλυσή της πρότασης χρησιμοποιώντας το εργαλείο Tree Tagger (Schmid, 2013), το οποίο προσδιορίζει για κάθε λέξη της πρότασης τι μέρος του λόγου είναι καθώς και τη βασική της μορφή (λήμμα). Στην συνέχεια γίνεται ανάλυση της δομής της πρότασης, προσδιορίζεται ο τρόπος που συνδέονται οι λέξεις και δημιουργείται το αντίστοιχο δέντρο εξαρτήσεων (dependency

tree). Έπειτα, γίνεται ανίχνευση των κύριων ονομάτων και των οντοτήτων που ενδεχομένως εμφανίζονται στην πρόταση χρησιμοποιώντας το εργαλείο Named Entity Recognizer (NER) (Finkel et al., 2005), με στόχο να καθοριστεί ο τρόπος που συναισθηματικές λέξεις συνδέονται κυρίως με οντότητες όπως άτομα.

Η βάση γνώσης (knowledge base) του μηχανισμού χρησιμοποιείται για τον εντοπισμό των λέξεων του φέρουν συναισθηματική πληροφορία και βασίζεται σε λεξικολογικούς πόρους όπως είναι το Wordnet Affect. Κάθε συναισθηματική λέξη που ανιχνεύτηκε στην πρόταση αναλύεται περαιτέρω και γίνεται επεξεργασία του ακριβούς τρόπου που συσχετίζεται και αλληλεπιδρά με τις λέξεις της πρότασης. Στόχος είναι να προσδιοριστούν οι σχέσεις των συναισθηματικών λέξεων με λέξεις ποσοτικοποίησης, ώστε να καθοριστεί η συναισθηματική τους ένταση. Τέλος, με βάση τα χαρακτηριστικά και τις πληροφορίες που εξήχθησαν από την πρόταση, καθορίζεται η συναισθηματική πληροφορία της. Πιο συγκεκριμένα, για μια νέα πρόταση:

1. Προσδιορίζεται για κάθε λέξη της πρότασης τι μέρος του λόγου είναι καθώς και η βασική μορφή της όπως προσδιορίστηκε από το εργαλείο Tree Tagger.
2. Προσδιορίζεται ο τρόπος που συνδέονται οι λέξεις και δημιουργείται το δέντρο εξαρτήσεων από το εργαλείο Stanford Parser.
3. Προσδιορίζονται οι κύριες οντότητες και τα άτομα από το εργαλείο Named Entity Recognizer.
4. Για κάθε λέξη, καθορίζεται αν περιέχει ή όχι συναισθηματικό περιεχόμενο από την βάση γνώσης. Εάν περιέχει, τότε:
 - 4.1 Γίνεται ανάλυση των συσχετίσεών της με άλλες λέξεις της πρότασης.
 - 4.2 Αναγνωρίζονται συσχετίσεις της με λέξεις ποσοτικοποίησης, καθορίζεται ο βαθμός της συναισθηματικής έντασης της λέξης και δημιουργείται η αντίστοιχη συναισθηματική δομή.
5. Γίνεται ανάλυση του δέντρου εξαρτήσεων, καθορίζονται οι συναισθηματικές δομές που υπάρχουν σε αυτή και προσδιορίζεται το συναισθηματικό περιεχόμενό της με βάση τις συναισθηματικές δομές.

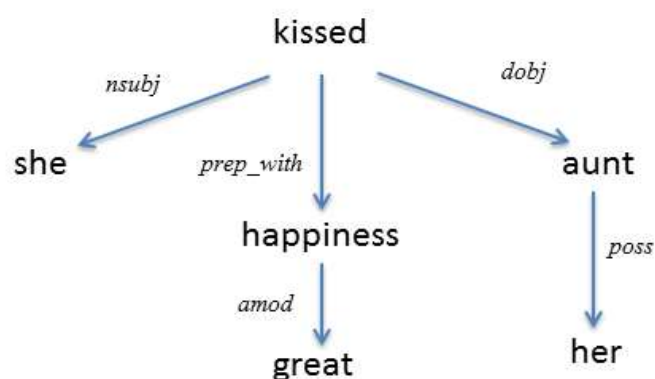
Η βάση γνώσης του μηχανισμού περιέχει δυο ειδών πληροφορίες, οι οποίες αφορούν τις συναισθηματικές λέξεις που είναι γνωστό ότι φέρουν συναισθηματικό περιεχόμενο καθώς και πληροφορίες για λέξεις ποσοτικοποίησης (quantification words). Η βάση γνώσης βασίζεται και επεκτείνει τον λεξικολογικό πόρο Wordnet Affect, που είναι μια ευρέως

χρησιμοποιούμενη επέκταση του Wordnet και η οποία επεκτάθηκε και εμπλουτίστηκε με επιπρόσθετες λέξεις, την κατηγορία τους και τον γραμματικό τους ρόλο. Επιπλέον, η βάση γνώσης διατηρεί πληροφορίες σχετικά με λέξεις ποσοτικοποίησης, οι οποίες αποτελούν έναν ειδικό τύπο λέξεων που μπορούν να ποσοτικοποιήσουν και να καθορίσουν τον βαθμό και το περιεχόμενο άλλων λέξεων καθώς και των συναισθηματικών λέξεων. Μερικά παραδείγματα λέξεων ποσοτικοποίησης παρουσιάζονται στον Πίνακα 5.1.

5.4.6.1 Μορφοσυντακτική ανάλυση προτάσεων

Ο προσδιορισμός και η μορφοσυντακτική ανάλυση της δομής μια πρότασης είναι το πρώτο στάδιο για την επεξεργασία της. Μετά την μορφοσυντακτική ανάλυσή της από το εργαλείο Tree Tagger γίνεται μια βαθύτερη ανάλυση της δομής της πρότασης και προσδιορίζονται οι σχέσεις μεταξύ των λέξεων και δημιουργείται το δέντρο εξαρτήσεων. Το δέντρο εξαρτήσεων της πρότασης αναπαριστά τις γραμματικές σχέσεις μεταξύ των λέξεων, οι οποίες αναπαρίστανται ως τριπλέτες (triples). Κάθε τριπλέτα αποτελείται από το *όνομα* της σχέσης, τον *κυβερνήτη* (governor), που εκτελεί την σχέση, και το *εξαρτώμενο* (dependent), που δέχεται την σχέση αντίστοιχα. Οι εξαρτήσεις αυτές δείχνουν τον ακριβή τρόπο που οι λέξεις της πρότασης συνδέονται και αλληλεπιδρούν μεταξύ τους και μπορούν να προσφέρουν σημαντικές πληροφορίες για τον τρόπο που προκύπτει το νόημα μιας πρότασης.

Στην Εικόνα 5.8 παρουσιάζεται ένα παράδειγμα της ανάλυσης της πρότασης “She kissed her aunt with great happiness” και το αντίστοιχο δέντρο εξαρτήσεων που προέκυψε με βάση τις σχέσεις μεταξύ των λέξεων της πρότασης.



Εικόνα 5.8 Το δέντρο εξαρτήσεων της πρότασης “She kissed her aunt with great happiness”

Οι κόμβοι του δέντρου εξαρτήσεων είναι οι λέξεις της πρότασης και οι ακμές του δέντρου αναπαριστούν τις σχέσεις μεταξύ των λέξεων. Για κάθε λέξη καθορίζεται αν αυτή συσχετίζεται με κάθε μια από τις υπόλοιπες λέξεις της πρότασης και στην περίπτωση που

συσχετίζεται, γίνεται καθορισμός του τρόπου συσχέτισής τους. Η αλληλεπίδραση μεταξύ δύο λέξεων της πρότασης συμβολίζεται με την ύπαρξη μια ακμής μεταξύ τους και το είδος της σχέσης συσχέτισης και αλληλεπίδρασης συμβολίζεται με το όνομα της ακμής αυτής. Για παράδειγμα, στην παραπάνω πρόταση η συσχέτιση *nsubj* (*she*, *kissed*) ορίζει ότι η λέξη “*she*” είναι το υποκείμενο της πρότασης και του ρήματος “*kissed*”.

5.4.6.2 Δημιουργία συναισθηματικών δομών

Το σύστημα με βάση την ανάλυση της πρότασης αναγνωρίζει τις λέξεις που φέρουν συναισθηματικό περιεχόμενο βασιζόμενο στην βάση γνώσης. Για κάθε μια λέξη της πρότασης, η βάση γνώσης αναγνωρίζει αν η λέξη φέρει συναισθηματική πληροφορία και επιστρέφει την συναισθηματική της κατηγορία. Αν δεν αναγνωριστούν συναισθηματικές λέξεις στην πρόταση, τότε εκτελείται μια βαθύτερη ανάλυση των λέξεων της πρότασης, όπου γίνεται για κάθε λέξη της πρότασης ανάκτηση των συνώνυμων της και στην συνέχεια γίνεται προσδιορισμός αν κάποια από αυτές φέρει συναισθηματική πληροφορία. Στην συνέχεια ο μηχανισμός, μετά την αναγνώριση των συναισθηματικών λέξεων, δημιουργεί τις συναισθηματικές δομές που υπάρχουν στην πρόταση. Ως συναισθηματική δομή ορίζεται κάθε ποσοτικοποιημένη συναισθηματική λέξη της πρότασης. Για να γίνει αυτό, ο μηχανισμός βασίζεται στην βάση γνώσης και ανακτά πληροφορίες σχετικά με τις λέξεις ποσοτικοποίησης. Οι λέξεις ποσοτικοποίησης μπορούν να καθορίσουν τον βαθμό και το περιεχόμενο συναισθηματικών λέξεων μέσω της αλληλεπίδρασης με αυτές. Για τον σκοπό αυτό, αποθηκεύονται στην βάση γνώσης λέξεις ποσοτικοποίησης καθώς και ο βαθμός της επίδρασής τους. Μερικά παράδειγμα τέτοιων λέξεων είναι: *very*, *great*, *huge*, *quite*, *little* κ.α. Επίσης, μια ειδική κατηγορία των λέξεων ποσοτικοποίησης αποτελούν οι λέξεις που δηλώνουν άρνηση. Οι λέξεις άρνησης, όταν εμφανίζονται σε μια πρόταση μπορούν να αντιστρέψουν την πολικότητα των λέξεων με τις οποίες αλληλεπιδρούν, και παραδείγματα τέτοιων λέξεων αποτελούν οι λέξεις: *no*, *not*, *no one* κ.α. Οι λέξεις ποσοτικοποίησης, όπως για παράδειγμα οι λέξεις “*great*” και “*very*”, έχουν ένα ισχυρό, θετικό αντίκτυπο στις λέξεις που σχετίζονται μαζί τους, αυξάνοντας την ένταση των συναισθηματικών λέξεων, ενώ λέξεις άρνησης μπορούν να αναστρέψουν το συναισθηματικό περιεχόμενο των λέξεων. Στον Πίνακα 5.1, παρουσιάζονται παραδείγματα λέξεων ποσοτικοποίησης καθώς και ο βαθμός επίδρασης τους.

Πίνακας 5.1 Λέξεις ποσοτικοποίησης και βαθμός επίδρασης τους

Λέξεις	Βαθμός επίδρασης
<i>very</i> , <i>great</i> , <i>huge</i> , <i>extensive</i>	<i>high</i>

hardly, quite	average
little, less	low
no, not	flip

Ο βαθμός επίδρασης μπορεί να κατηγοριοποιηθεί σε τέσσερις βασικές κατηγορίες high, average, low και flip. Η αποτύπωση των τιμών του βαθμού επίδρασης κυμαίνεται από 0 μέχρι 100, όπου το 100 αναπαριστά την μέγιστη δυνατή ένταση. Ο καθορισμός των τιμών έγινε με βάση εμπειρικές και πειραματικές μελέτες και καθορίστηκαν, ο βαθμός επίδρασης “low” να έχει τιμή 20, ο βαθμός επίδρασης “average” να έχει τιμή 50 και ο βαθμός επίδρασης “high” να έχει τιμή 100. Επομένως, λέξεις όπως η λέξη “great”, όταν συσχετίζεται με συναισθηματική λέξη, θέτει την ένταση της στο 100, ενώ λέξεις, όπως για παράδειγμα η λέξη “quite”, θέτει την ένταση της συναισθηματικής λέξης στο 20. Η αλληλεπίδραση των λέξεων ποσοτικοποίησης με άλλες λέξεις αποτυπώνεται στο δέντρο εξαρτήσεων και αναγνωρίζονται ως εξαρτήσεις τύπου «mod».

Μετά τον καθορισμό και την αναγνώριση των λέξεων που φέρουν συναισθηματική πληροφορία, γίνεται ανάλυση του ρόλου τους στην πρόταση και του τρόπου που συνδέονται και αλληλεπιδρούν με τις λέξεις της πρότασης. Πιο συγκεκριμένα, ο μηχανισμός εντοπίζει και αναλύει ειδικούς τύπους σχέσεων που μπορεί να υπάρχουν με λέξεις ποσοτικοποίησης. Οι σχέσεις αυτές αναγνωρίζονται ως εξαρτήσεις τύπου «mod» στο δέντρο εξαρτήσεων. Επομένως, οι εξαρτήσεις που συνδέουν συναισθηματικές λέξεις με λέξεις ποσοτικοποίησης αναλύονται προκειμένου να καθοριστεί ο βαθμός της έντασης των συναισθηματικών λέξεων.

Για παράδειγμα, ας θεωρήσουμε την πρόταση “is not very furious”. Από την ανάλυση της πρότασης η βάση γνώσης αναγνωρίζει την λέξη “furious” ως συναισθηματική λέξη, η οποία υποδηλώνει το συναίσθημα του θυμού (anger). Επίσης, αναγνωρίζεται ότι η λέξη “furious” συνδέεται με την λέξη ποσοτικοποίησης “very” μέσω εξάρτησης τύπου «mod» και επομένως η λέξη “very” έχει υψηλή επίδραση στο περιεχόμενο της λέξης “furious”, το οποίο η βάση γνώσης αποτιμά στο 80. Στην συνέχεια, αναγνωρίζεται ότι η λέξη “not” αντιστρέφει την επίδραση αυτή, από υψηλή σε χαμηλή και έτσι το συναισθηματικό περιεχόμενο της φράσης αυτής καθορίζεται να δηλώνει λίγο θυμό με βαθμό 20.

Μετά την δημιουργία των συναισθηματικών δομών, ο μηχανισμός προσδιορίζει την γενική συναισθηματική πληροφορία της πρότασης με βάση τη μορφοσυντακτική δομή της πρότασης και του τρόπου συσχέτισης των επιμέρους συναισθηματικών δομών της. Πιο συγκεκριμένα, αναγνωρίζει και αναλύει την βασική δομή της πρότασης που αποτελείται από το κύριο ρήμα, το αντικείμενο και το υποκείμενο. Με βάση το δέντρο εξαρτήσεων της πρότασης, γίνεται προσδιορισμός του προτύπου υποκείμενο-ρήμα-αντικείμενο (subject-verb-object), το οποίο αποτελεί τη βασική δομή της πρότασης και το οποίο καθορίζει το βασικό νόημά της. Πιο συγκεκριμένα, ο μηχανισμός επεξεργάζεται την δομή των προτάσεων ως εξής:

1. Γίνεται ανάλυση του δέντρου εξαρτήσεων της πρότασης και προσδιορισμός του προτύπου υποκείμενο-ρήμα-αντικείμενο (subject-verb-object).
2. Για κάθε ένα από τα μέρη του προτύπου:
 - 2.1 Προσδιορίζεται αν αποτελεί μια συναισθηματική δομή.
 - 2.2 Αναλύεται η σχέση της με άλλες συναισθηματικές δομές της πρότασης (αν υπάρχουν),
 - 2.3 Προσδιορίζεται το συναισθηματικό περιεχόμενο του μέρους του προτύπου.
3. Γίνεται συνδυασμός των συναισθηματικών δομών της πρότασης και προσδιορισμός του περιεχομένου της πρότασης.

5.4.7 Ομαδοποίηση Ταξινομητών

5.4.7.1 Αρχή Πλειοψηφίας

Στην προσέγγιση της ομαδοποίησης των ταξινομητών, σε ένα σχήμα που ακολουθεί την αρχή της πλειοψηφίας (majority voting), κάθε ταξινομητής δηλώνει την κλάση που ταξινομεί το νέο στιγμιότυπο και ο μηχανισμός κατατάσσει το στιγμιότυπο στην κλάση που συγκέντρωσε τις περισσότερες ψήφους των επιμέρους ταξινομητών. Η αρχή της πλειοψηφίας αποτελεί την πιο απλή προσέγγιση συνδυασμού επιμέρους ταξινομητών σε ένα ομαδοποιημένο σχήμα (Kuncheva, 2004). Ας θεωρήσουμε ένα παράδειγμα x προς ταξινόμηση, δεδομένου ενός συνόλου C από R κλάσεις $\{c_1, c_2, \dots, c_R\}$ και επίσης ενός ομαδοποιημένου ταξινομητή αποτελούμενου από M επιμέρους ταξινομητές $\{A_1, A_2, \dots, A_M\}$. Αν $c_{xl} \in C$ δηλώνει την κλάση του παραδείγματος x όπως ταξινομήθηκε από τον ταξινομητή A_l , $l = 1, 2, \dots, R$, τότε ορίζεται μια *συνάρτηση ψηφοφορίας* F_k ως εξής:

$$F_k(c_i, c_{xl}) = \begin{cases} 1, & c_i = c_{xl} \\ 0, & c_i \neq c_{xl} \end{cases}$$

όπου c_i είναι οι κλάσεις του C . Ο συνολικός αριθμός ψήφων για μια κλάση c_i για το παράδειγμα x υπολογίζεται ως:

$$T_i = \sum_{l=1}^M F_k(c_i, c_{xl})$$

Η κλάση c που θα προσδιοριστεί για το παράδειγμα x ορίζεται ως :

$$c = \operatorname{argmax} T_i$$

$$\text{όπου } i \in \{1, \dots, R\}$$

Στην περίπτωση που οι επιμέρους ταξινομητές κάνουν ανεξάρτητα λάθη, η προσέγγιση της αρχής της πλειοψηφίας (majority voting) είναι κατάλληλη και υπερτερεί του καλύτερου επιμέρους ταξινομητή (Orrite et al., 2008).

5.4.7.2 Bagging

Η μέθοδος Bagging αποτελεί μια από τις πρώτες μεθόδους ομαδοποίησης ταξινομητών. Η μέθοδος βασίζεται στην αρχή της εκπαίδευσης καθενός βασικού ταξινομητή με βάση ένα τυχαίο υποσύνολο δεδομένων εκπαίδευσης όπου παράγεται έτσι ένας αριθμός μοντέλων ταξινομητών, χρησιμοποιώντας όμως διαφορετική διαμέριση του σώματος εκπαίδευσης (επιλογή με επανατοποθέτηση) για κάθε ένα εξ αυτών (Breiman, 1996). Η μέθοδος Bagging υποθέτει ένα σύνολο χ από N παραδείγματα και εκπαιδεύει κάθε ταξινομητή σε ένα υποσύνολο δεδομένων (bags). Αρχικά, το σύνολο δεδομένων μετατρέπεται σε επιμέρους υποσύνολα εκπαίδευσης χρησιμοποιώντας τεχνικές επιλογής (sampling) και επανατοποθέτησης (replacement) και στην συνέχεια κάθε υποσύνολο ανατίθεται για την εκπαίδευση ενός ταξινομητή. Με αυτόν τον τρόπο επιτυγχάνεται η ποικιλομορφία χρησιμοποιώντας διαφορετικά μέρη του συνόλου δεδομένων. Για τη λήψη της τελικής απόφασης ακολουθείται η πλειοψηφική λογική όπου η τελική απόφαση του ένθετου συστήματος ταξινόμησης συμπίπτει με την απόφαση της πλειοψηφίας των βασικών ταξινομητών (Εικόνα 5.9).

Input: Data set $D = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\}$;

Base Learning algorithm L ;

Number of learning rounds T .

Process:

```

For t=1,2,...,T:

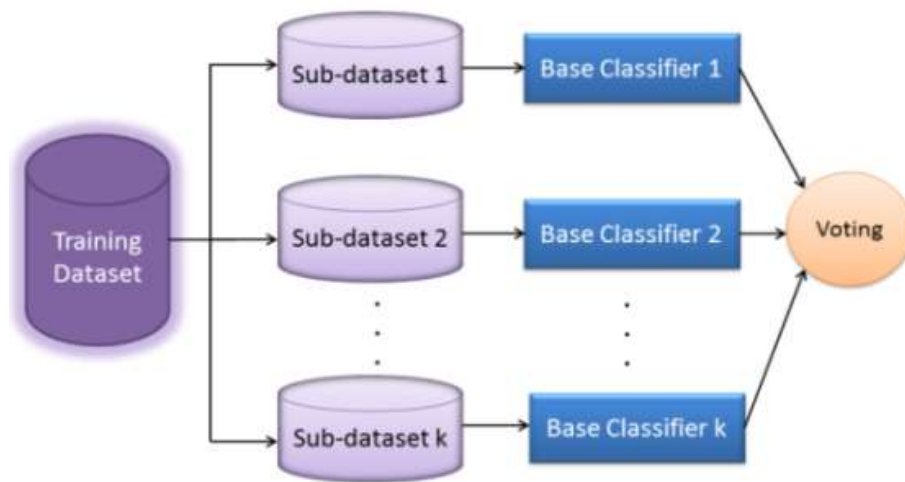
     $D_t = Bootstrap(D)$ ; %Generate a bootstrap sample from D

     $h_t = L(D_t)$           % Train a base learner  $h_t$  from the bootstrap sample

end;

Output:  $H(x) = \arg \max_{y \in Y} \sum_{t=1}^T l(y = h_t(x))$ 
% the value of l(a) is 1 if a is true;
% and 0 otherwise

```



Εικόνα 5.9 Η μέθοδος Bagging.

Η μέθοδος Bagging μπορεί να είναι ιδιαίτερα αποδοτική όταν χρησιμοποιούνται βασικοί ταξινομητές που μπορεί να είναι ασταθείς και να μεταβάλλονται αρκετά σε μικρές διαταραχές του συνόλου εκπαίδευσης. Επίσης, η μέθοδος Bagging δεν είναι επιρρεπής σε φαινόμενα υπερεκπαίδευσης (overfitting) καθώς αυξάνει το πλήθος των παραγόμενων υποθέσεων.

5.4.7.3 Boosting

Η μέθοδος Boosting έχει ως βασική αρχή την στάθμιση των παραδειγμάτων εκπαίδευσης σύμφωνα με τη δυσκολία της ταξινόμησης τους. Η μέθοδος σχηματίζει ομαδοποιημένο ταξινομητή με βάση την εκπαίδευση νέων μοντέλων σε παραδείγματα/στιγμιότυπα τα οποία κατηγοριοποιήθηκαν λανθασμένα. Η μέθοδος Boosting αναθέτει στα στιγμιότυπα εκπαίδευσης ένα βάρος, ενδεικτικό της δυσκολίας που παρουσιάζει το υπό εκμάθηση τρέχον μοντέλο στην ταξινόμησή του και κατ' επέκταση της

βαρύτητας που θα πρέπει να δοθεί σε αυτό κατά την παραγωγή του επόμενου μοντέλου, προκειμένου να αναγνωρισθεί σωστά. Αρχικά τα στιγμιότυπα εκπαίδευσης έχουν το ίδιο βάρος, $w = 1/N$, όπου N το σύνολο των στιγμιότυπων του συνόλου δεδομένων. Στην συνέχεια, το πρώτο μοντέλο εκπαιδεύεται και η απόδοσή του αξιολογείται, υπολογίζοντας το ζυγισμένο σφάλμα του e_1 επί του σώματος εκπαίδευσης ως το άθροισμα των βαρών των εσφαλμένα ταξινομηθέντων στιγμιότυπων προς το συνολικό άθροισμα των βαρών. Έπειτα, όλα τα βάρη προσαρμόζονται, και στα παραδείγματα που ταξινομήθηκαν σωστά τα βάρη μειώνονται, ενώ αυξάνονται στα παραδείγματα με εσφαλμένες προβλέψεις. Αυτή η ακολουθία βημάτων επαναλαμβάνεται μέχρι τη δημιουργία ενός καθορισμένου πλήθους μοντέλων. Πιο συγκεκριμένα, κατά την ταξινόμηση ενός άγνωστου στιγμιότυπου, κάθε ένα από τα παραγόμενα μοντέλα προβαίνει στην εκτίμησή του, η οποία μπορεί να συμμετέχει με διαφορετική βαρύτητα στην τελική απόφαση. Σε κάθε μοντέλο ανατίθεται ένας συντελεστής βαρύτητας, ανάλογα με το σφάλμα του μοντέλου. Ο συντελεστής αυτός πολλαπλασιάζεται με την εκτίμηση του μοντέλου, ώστε να αναδειχθεί η τελική απόφαση. Έστω e το κανονικοποιημένο σφάλμα ενός μοντέλου για ένα συγκεκριμένο παράδειγμα. Το βάρος του συγκεκριμένου παραδείγματος αναπροσαρμόζεται πολλαπλασιαζόμενο με $e/(1-e)$. Επομένως, μικρότερα λάθη οδηγούν σε μικρότερα βάρη και οι επόμενες επαναλήψεις εστιάζουν στα παραδείγματα που ταξινομούνται δύσκολα. Ο γενικός τρόπος λειτουργίας της μεθόδου Boosting απεικονίζεται στην Εικόνα 5.10.

Input: Data set $D = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\}$;

Base Learning algorithm L ;

Number of learning rounds T .

Process:

$$D_1(i) = 1/m;$$

For $t=1,2,\dots,T$:

$h_t = L(D, D_t)$ % Train a base learner h_t from D using distribution D_t

$\varepsilon_t = Pr_{i \sim D_t}[h_t(x_i) \neq y_i]$; % Measure the error of h_t

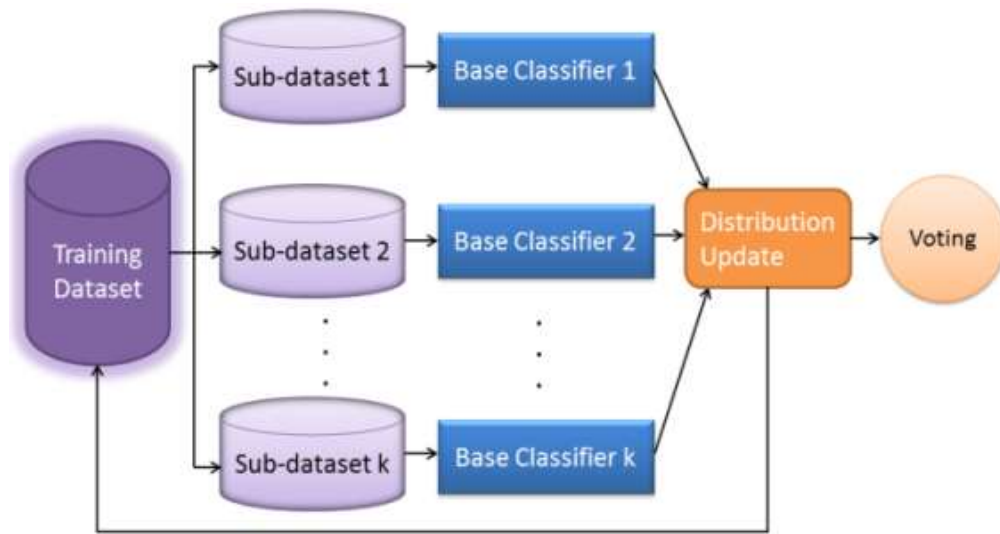
$$\alpha_t = \frac{1}{2} \ln \frac{1-\varepsilon_t}{\varepsilon_t}; \text{ % Determine the weight of } h_t$$

$$D_{t+1} = \frac{D_t(i)}{Z_t} \times \begin{cases} \exp(-a_t) & \text{if } h_t = y_i \\ \exp(a_t) & \text{if } h_t \neq y_i \end{cases} \quad \% \text{ Update the distribution where } Z_t \text{ is a}$$

$$= \frac{D_t(i) \exp(-a_t y_i h_t(x_i))}{Z_t} \quad \% \text{ normalization factor which enables to be a distribution}$$

end.

Output: $H(x) = \text{sign}(f(x)) = \text{sign} \sum_{t=1}^T a_t h_t(x)$



Εικόνα 5.10 Η μέθοδος Boosting.

Στην συνέχεια, γίνεται μελέτη της απόδοσης κάθε προσέγγισης δημιουργίας ομαδοποιημένων ταξινομητών στην αναγνώριση της συναισθηματικής πληροφορίας που φέρουν προτάσεις φυσικής γλώσσας. Πιο συγκεκριμένα, γίνεται μελέτη της συμπεριφοράς κάθε επιμέρους ταξινομητή καθώς και κάθε μεθόδου ομαδοποίησης, τόσο στον προσδιορισμό της ύπαρξης ή όχι συναισθηματικού περιεχομένου σε προτάσεις φυσικής γλώσσας όσο και στον προσδιορισμό της συναισθηματικής πολικότητας τους.

5.5 Προσδιορισμός Συναισθηματικής πολικότητας

Στο δισδιάστατο μοντέλο του Russell (Russell, 1980), τα συναισθήματα μπορούν να αναπαρασταθούν σε ένα χώρο δύο διαστάσεων, όπου η μία διάσταση αντιπροσωπεύει την πολικότητα (polarity) του συναισθήματος και η άλλη διάσταση την ενεργοποίηση (activation). Η πολικότητα χαρακτηρίζει ένα συναισθήμα ως θετικό (positive) ή αρνητικό (negative), ενώ η ενεργοποίηση χαρακτηρίζει ένα συναισθήμα σαν ενεργό ή ανενεργό

(activated ή deactivated). Με αυτή την προσέγγιση, διαμορφώνονται τέσσερις τομείς συναισθημάτων: ενεργό-θετικό, ενεργό-αρνητικό, ανενεργό-θετικό και ανενεργό-αρνητικό (Ptaszynski et al., 2009). Στην Εικόνα 5.11, απεικονίζονται ο συναισθηματικός χώρος και η αναπαράσταση των συναισθημάτων στον χώρο αυτό.



Εικόνα 5.11. Αποτύπωση των συναισθημάτων στον χώρο του μοντέλου του Russell

Η αποτύπωση επιτρέπει στο σύστημα να καθορίζει την πολικότητα των συναισθηματικών προτάσεων. Πιο συγκεκριμένα, σε περίπτωση που μια πρόταση αναγνωρίζεται από το σύστημα ότι περιέχει συναισθήματα, γίνεται καθορισμός της συναισθηματικής πολικότητά της με βάση τον χώρο του Russell. Για παράδειγμα, η χαρά συνδέεται με την θετική πολικότητα, ενώ τα συναισθήματα του θυμού, της αηδίας, του φόβου και της θλίψης χαρακτηρίζουν μια πρόταση ως αρνητική (Chaumartin, 2007). Σε αυτή την κατεύθυνση, η έκπληξη μπορεί να χαρακτηρίσει μια πρόταση είτε ως θετική σε περιπτώσεις που συνοδεύεται με το συναίσθημα της χαράς (χαρούμενη έκπληξη), είτε ως αρνητική στις περιπτώσεις όπου συνδέεται με τα συναισθήματα του φόβου, θύμου, θλίψης και αηδίας. Το σύστημα υιοθετεί αυτή την συναισθηματική αναπαράσταση για να καθορίζει τη συναισθηματική πολικότητα μιας πρότασης με βάση το συναισθηματικό της περιεχόμενο.

5.6 Πειραματική Αξιολόγηση

Για την αξιολόγηση της απόδοσης κάθε ενός ταξινομητή καθώς και των σχημάτων ομαδοποιημένων ταξινομητών σχεδιάστηκε και πραγματοποιήθηκε μια πειραματική μελέτη. Η πειραματική μελέτη που πραγματοποιήθηκε είχε τρία κύρια μέρη όπου κάθε ένα αξιολόγησε διαφορετικές συνιστώσες των συστημάτων. Πιο συγκεκριμένα, στα πλαίσια του πρώτου μέρους, έγινε μελέτη και αξιολόγηση της απόδοσης των βασικών και των

ομαδοποιημένων ταξινομητών στην αναγνώριση της ύπαρξης συναισθηματικού περιεχομένου σε προτάσεις. Στην συνέχεια, στο δεύτερο μέρος της πειραματικής μελέτης, αξιολογήθηκε η απόδοσή των ταξινομητών στον προσδιορισμό της συναισθηματικής πολικότητας των προτάσεων που φέρουν συναισθηματικό περιεχόμενο. Τέλος, στο τρίτο μέρος, έγινε αξιολόγηση της απόδοσης των ταξινομητών στην αναγνώριση του συναισθηματικού περιεχομένου των προτάσεων σύμφωνα με το μοντέλο συναισθημάτων του Ekman. Επίσης, στα πλαίσια της μελέτης, έγινε συλλογή διάφορων ειδών κειμενικών δεδομένων και δημιουργήθηκε ένα σύνολο δεδομένων με στόχο να γίνει πολύπλευρη αξιολόγηση της απόδοσης του συστήματος σε διαφορετικούς τύπους κειμένου.

5.6.1 Δημιουργία Συνόλου Δεδομένων

Για τους σκοπούς της πειραματικής μελέτης έγινε συλλογή κειμενικών δεδομένων από διάφορες πηγές και σχηματίστηκε ένα σύνολο δεδομένων το οποίο περιείχε συνολικά 1200 προτάσεις φυσικής γλώσσας. Οι προτάσεις επιλέχθηκαν από τρία είδη πηγών κειμένων τα οποία ήταν τίτλοι ειδήσεων, άρθρα ειδήσεων και μηνύματα χρηστών στο Twitter. Τα άρθρα ειδήσεων και οι τίτλοι συλλέχθηκαν από τις ειδησεογραφικές ιστοσελίδες BBC, CNN και Euronews. Μετά την δημιουργία του συνόλου, έγινε καθορισμός του συναισθηματικού περιεχομένου των προτάσεων από εμπειρογνώμονα-ειδικό. Πιο συγκεκριμένα, προσδιορίστηκαν για κάθε πρόταση από τον εμπειρογνώμονα τα εξής: (α) η ύπαρξη συναισθημάτων σύμφωνα με το μοντέλο του Ekman και (β) ο βαθμός της έντασης για κάθε ένα από τα συναισθήματα. Ο βαθμός της έντασης καθορίστηκε από τον ειδικό με βαθμολογία στην κλίμακα 0 έως 100, όπου το 0 χρησιμοποιείται για να υποδηλώσει την απουσία ενός συγκεκριμένου συναισθήματος και το 100 υποδηλώνει ότι το συναίσθημα είναι πολύ ισχυρό. Στα πλαίσια της πειραματικής μελέτης οι προσδιορισμοί του εμπειρογνώμονα χρησιμοποιήθηκαν ως αναφορά (gold standard) για την αξιολόγηση των επιδόσεων των ταξινομητών.

5.6.2 Αξιολόγηση Προσδιορισμού Ύπαρξης Συναισθηματικού Περιεχομένου

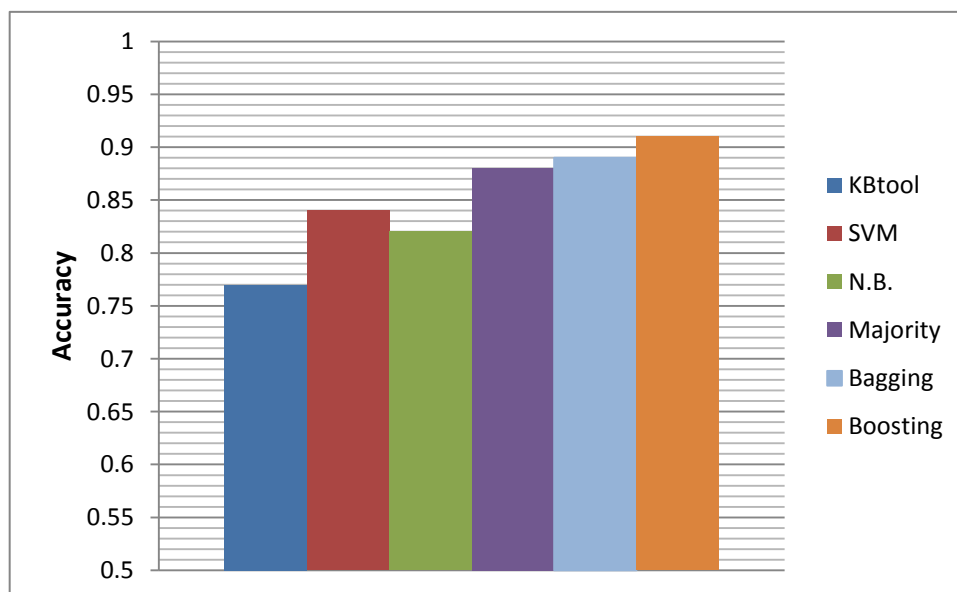
Το πρώτο μέρος της πειραματικής μελέτης αφορούσε την αξιολόγηση των επιδόσεων των βασικών ταξινομητών και των ομαδοποιημένων σχημάτων ταξινόμησης ως προς τον καθορισμό της ύπαρξης συναισθηματικού περιεχομένου σε αυτές. Για την αξιολόγηση της απόδοσης των μηχανισμών, χρησιμοποιήθηκαν οι μετρικές: ορθότητα (accuracy), ακρίβεια (precision), ευαισθησία (sensitivity) και εξειδίκευση (specificity), δεδομένου ότι η κλάση

εξόδου είναι μια δυαδική μεταβλητή (binary class output variable) (Sokolova & Lapalme, 2009). Τα αποτελέσματα που προέκυψαν παρουσιάζονται στον Πίνακα 5.2.

Πίνακας 5.2 Συνολική απόδοση των ταξινομητών

	Knowledge Based Tool	SVM	Naïve Bayes	Majority	Bagging	Boosting
Accuracy	0.77	0.84	0.82	0.88	0.89	0.91
Precision	0.73	0.90	0.87	0.91	0.93	0.93
Sensitivity	0.94	0.87	0.88	0.91	0.92	0.94
Specificity	0.58	0.76	0.70	0.80	0.83	0.83

Τα αποτελέσματα δείχνουν μια πολύ καλή απόδοση για τους βασικούς ταξινομητές καθώς και για τους ομαδοποιημένους ταξινομητές. Αξίζει να σημειωθεί ότι οι ομαδοποιημένοι ταξινομητές παρουσίασαν σε όλα τα πειράματα καθολικά καλύτερη απόδοση σε σχέση με τους βασικούς ταξινομητές. Η απόδοση αυτή οφείλεται κυρίως στο γεγονός ότι η κατηγοριοποίηση πραγματοποιείται με πολύ καλή απόδοση και από τους τρεις βασικούς ταξινομητές και επομένως στα ομαδοποιημένα σχήματα σε πολλές προτάσεις που ένας από τους βασικούς ταξινομητές αποτυγχάνει να πραγματοποιήσει μια σωστή πρόβλεψη η τελική πρόβλεψη γίνεται σωστά με βάση τους υπόλοιπους δύο ταξινομητές. Η ακρίβεια των ταξινομητών στο σύνολο δεδομένων της πειραματικής μελέτης παρουσιάζεται στην Εικόνα 5.12.



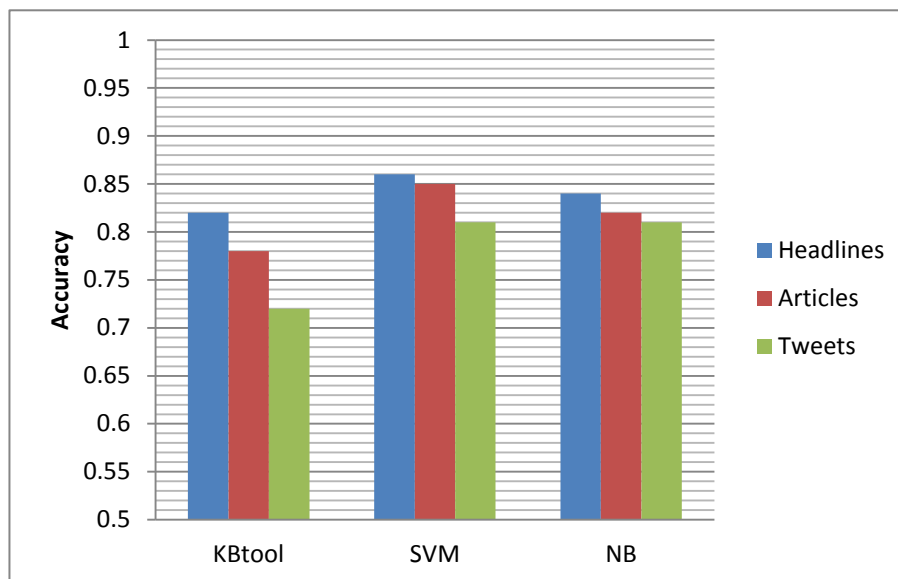
Εικόνα 5.12 Η ορθότητα των ταξινομητών σε όλο το σύνολο δεδομένων

Τα αναλυτικά αποτελέσματα των επιδόσεων των βασικών ταξινομητών για κάθε τύπο κειμενικών δεδομένων παρουσιάζονται στον Πίνακα 5.3.

Πίνακας 5.3 Ανάλυση της απόδοσης των βασικών ταξινομητών σε κάθε είδος κειμενικών δεδομένων

	Headlines			Articles			Tweets		
	KBtool	SVM	N.B.	KBtool	SVM	N.B.	KBtool	SVM	N.B.
Accuracy	0.82	0.86	0.84	0.78	0.85	0.82	0.72	0.81	0.81
Precision	0.80	0.92	0.91	0.72	0.91	0.89	0.66	0.89	0.81
Sensitivity	0.93	0.88	0.87	0.95	0.88	0.86	0.92	0.85	0.90
Specificity	0.65	0.80	0.76	0.59	0.77	0.71	0.52	0.70	0.64

Όσον αφορά τους βασικούς ταξινομητές, η καλύτερη απόδοση επιτυγχάνεται από τον SVM, όπου η ορθότητα και η ακρίβεια του είναι καλύτερες από εκείνες του μηχανισμού βασισμένου στην γνώση (knowledge based) καθώς και του ταξινομητή Naïve Bayes. Η ορθότητα των βασικών ταξινομητών στα διάφορα είδη κειμενικών δεδομένων παρουσιάζονται στην Εικόνα 5.13.



Εικόνα 5.13 Η ορθότητα των βασικών ταξινομητών σε κάθε είδος δεδομένων κειμένου.

Ο ταξινομητής SVM παρουσιάζει την καλύτερη απόδοση του σε τίτλους άρθρων και σε άρθρα, αλλά η απόδοση του μειώνεται σε δεδομένα από tweets σχεδόν στον ίδιο βαθμό που μειώνεται και η απόδοση του Naive Bayes. Όσον αφορά τον knowledge based μηχανισμό, τα αποτελέσματα δείχνουν ότι έχει πολύ καλή απόδοση σε τίτλους ειδήσεων. Το γεγονός αυτό οφείλεται κυρίως στ' ότι οι τίτλοι των ειδήσεων είναι συντακτικά ορθά δομημένοι ακολουθώντας γραμματικούς κανόνες, κάτι που διευκολύνει την αποτελεσματική και εις βάθος ανάλυσή τους από τον μηχανισμό και τον σχηματισμό ορθών συντακτικών δομών. Επιπροσθέτως, στους τίτλους ειδήσεων στις περισσότερες περιπτώσεις τα συναισθήματα εκφράζονται με έντονες συναισθηματικές λέξεις γεγονός που σε συνδυασμό με την ορθή δόμηση τους συμβάλει στην αναγνώριση και ποσοτικοποίηση των λέξεων και την δημιουργία συναισθηματικών δομών. Η απόδοση του μηχανισμού είναι χαμηλότερη σε κειμενικά δεδομένα από μηνύματα χρηστών στο Twitter τα οποία, σε αντίθεση με τις προτάσεις των άρθρων, δεν έχουν αυστηρή και κατάλληλη συντακτική δομή και συχνά περιέχουν αναφορικές εκφράσεις και αμφισημίες. Επίσης, αξίζει να σημειωθεί ότι ο μηχανισμός παρουσιάζει πολύ υψηλή ευαισθησία (sensitivity), κάτι που δείχνει ότι προτάσεις που ήταν συναισθηματικά ουδέτερες αναγνωρίστηκαν αποτελεσματικά από τον μηχανισμό και ταξινομήθηκαν σωστά στην ουδέτερη κατηγορία. Επιπλέον, η χαμηλή εξειδίκευση (specificity) είναι κυρίως αποτέλεσμα της ταξινόμησης συναισθηματικών προτάσεων στην ουδέτερη κατηγορία. Στις περισσότερες από αυτές τις περιπτώσεις το συναισθηματικό περιεχόμενο των προτάσεων εκφράζεται περιφραστικά, χωρίς την ύπαρξη συναισθηματικών λέξεων.

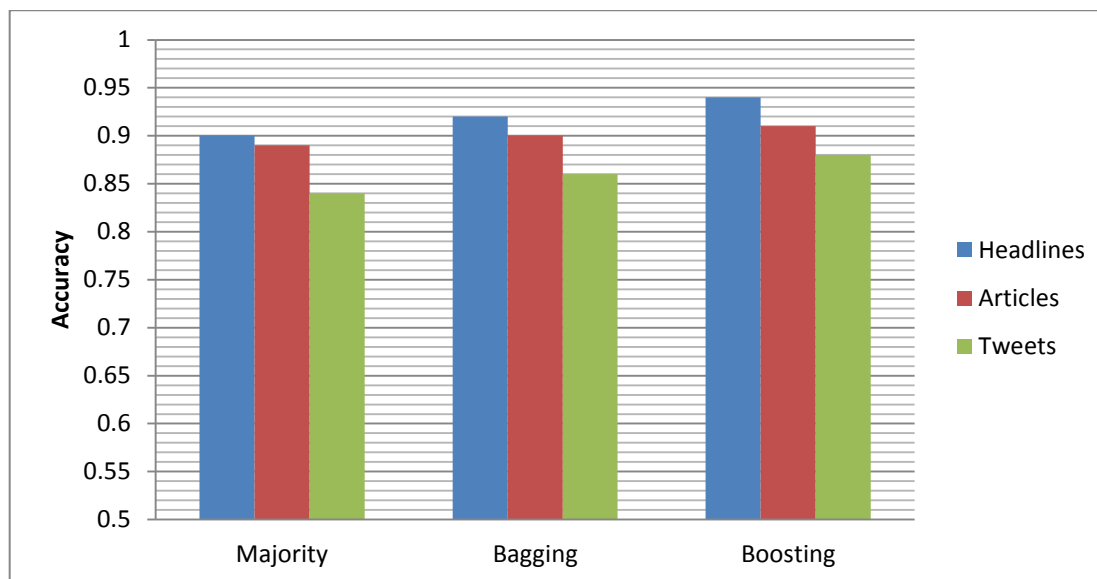
Οι βασικοί ταξινομητές παρουσιάζουν καλύτερη απόδοση σε τίτλους και κείμενα άρθρων και χαμηλότερη σε tweets. Η καλύτερη απόδοση στους τίτλους και τα κείμενα των άρθρων είναι λόγω του ότι ακολουθούν γραμματικούς κανόνες και εκφράζουν τα συναισθήματα χρησιμοποιώντας έντονα συναισθηματικές λέξεις. Από την άλλη πλευρά, τα δεδομένα των χρηστών στο Twitter δεν παρουσιάζουν μια καλή γραμματική σύνταξη και σε πολλές περιπτώσεις φέρουν συναισθηματικό περιεχόμενο με έμμεσο τρόπο, χωρίς την ύπαρξη συναισθηματικών λέξεων. Μάλιστα, στις περισσότερες περιπτώσεις είναι πολύ σύντομα, και αποτελούνται από λιγότερο από 10 λέξεις και επίσης έχουν αυθαίρετη και ασταθή δομή. Επιπλέον, τα συναισθήματα σε tweets εκφράζονται σε πολλές περιπτώσεις, με μια ειρωνική μορφή, με βάση τα σχόλια ή γεγονότα προηγούμενων posts και η συναισθηματική αναγνώριση αυτών των περιπτώσεων απαιτεί την ανάλυση και την βαθύτερη κατανόηση του θέματος και της συζήτησης. Όσον αφορά τους ομαδοποιημένους ταξινομητές, και οι τρεις παρουσιάζουν καθολικά καλύτερη απόδοση από τους βασικούς σε

όλα τα είδη κειμενικών δεδομένων και η ανάλυση της απόδοσής τους παρουσιάζεται στον Πίνακα 5.4.

Πίνακας 5.4 Ανάλυση της απόδοση των ομαδοποιημένων ταξινομητών

	Headlines			Articles			Tweets		
	Majority	Bagging	Boosting	Majority	Bagging	Boosting	Majority	Bagging	Boosting
Accuracy	0.90	0.92	0.94	0.89	0.90	0.91	0.84	0.86	0.88
Precision	0.96	0.94	0.95	0.92	0.95	0.93	0.86	0.90	0.90
Sensitivity	0.91	0.94	0.96	0.92	0.90	0.94	0.90	0.90	0.93
Specificity	0.88	0.87	0.89	0.82	0.87	0.84	0.71	0.77	0.78

Τα αποτελέσματα δείχνουν ότι οι ένθετοι ταξινομητές έχουν καλύτερη απόδοση στους τίτλους ειδήσεων, ενώ η απόδοσή τους μειώνεται σε κειμενικά δεδομένα χρηστών στο Twitter όπως συμβαίνει και με την απόδοση των βασικών ταξινομητών. Όσον αφορά τις μεθόδους δημιουργίας ομαδοποιημένων ταξινομητών, η μέθοδος *boosting* παρουσιάζει καλύτερα αποτελέσματα από την μέθοδο *bagging* και την αρχή της πλειοψηφίας. Πιο συγκεκριμένα, ο ομαδοποιημένος ταξινομητής που δημιουργήθηκε με την μέθοδο *boosting* παρουσιάζει καθολικά καλύτερα αποτελέσματα, τόσο σε σχέση με τους βασικούς ταξινομητές όσο και με τους ομαδοποιημένους ταξινομητές που ακολουθούν άλλες μεθόδους ομαδοποίησης, σε όλα τα είδη κειμενικών δεδομένων. Το γεγονός αυτό φαίνεται να είναι αποτέλεσμα της έμφασης που δίνει η μέθοδος *boosting* σε δεδομένα που είναι πιο δύσκολο να ταξινομηθούν σωστά. Η ορθότητα των ομαδοποιημένων ταξινομητών στα διάφορα είδη κειμενικών δεδομένων παρουσιάζονται στην Εικόνα 5.14.



Εικόνα 5.14 Η ορθότητα των ομαδοποιημένων ταξινομητών σε κάθε είδος δεδομένων κειμένου.

Στα πλαίσια της ανάλυσης της απόδοσης των ομαδοποιημένων ταξινομητών έγινε σύγκριση της απόδοσης τους σε σχέση με την απόδοση του ταξινομητή SVM, ο οποίος παρουσίασε την καλύτερη απόδοση από τους βασικούς. Στόχος ήταν να καθοριστεί το ποσοστό της βελτίωσης της διαδικασίας αναγνώρισης από τα ομαδοποιημένα σχήματα ως προς την απόδοση του καλύτερου βασικού ταξινομητή. Από τα αποτελέσματα προκύπτει ότι ο ομαδοποιημένος ταξινομητής που ακολουθεί την αρχή της πλειοψηφίας παρουσιάζει ορθότητα η οποία είναι υψηλότερη κατά 7.3% σε σχέση με την χρήση του ταξινομητή SVM. Επίσης, η ακρίβεια της μεθόδου bagging είναι υψηλότερη του SVM κατά 8.5% ενώ η ακρίβεια της μεθόδου boosting είναι υψηλότερη κατά 10.9% αντιστοίχως.

Στην συνέχεια, έγινε υπολογισμός του στατιστικού δείκτη Cohen's Kappa (Cohen, 1960), για την αποτίμηση της απόδοσης των ταξινομητών και για τον προσδιορισμό της αξιοπιστίας του βαθμού συμφωνίας μεταξύ του εμπειρογνώμονα και των ταξινομητών. Ο δείκτης Cohen's Kappa ορίζεται ως ένα μέτρο για τον προσδιορισμό της αξιοπιστίας της συμφωνίας ανάμεσα σε δύο αξιολογητές. Θεωρούμε ότι δυο αξιολογητές ταξινομήσαν N στοιχεία σε k αλληλο-αποκλειόμενες κατηγορίες. Τα αποτελέσματα παρουσιάζονται σε ένα πίνακα γειτνίασης $k \times k$. Το p_{ij} , όπου $i, j = 1, 2, \dots, k$ ορίζεται ως εξής:

$$p_{ij} = \frac{N_{ij}}{N}$$

όπου N_{ij} ο αριθμός των περιπτώσεων που, ενώ από τον αξιολογητή A (Rater A) κατηγοριοποιήθηκαν στην κατηγορία i , από τον αξιολογητή B (Rater B)

κατηγοριοποιήθηκαν στην κατηγορία j . Για τον υπολογισμό του k συμπληρώνεται ο Πίνακας 5.5, όπου:

$$p_{.j} = \sum_{i=1}^k p_{ij} \text{ και } p_{i.} = \sum_{j=1}^k p_{ij}$$

Πίνακας 5.5 Υπολογισμός Cohen's Kappa

RATER A	RATER B				
	1	2	...	k	Total
1	p_{11}	p_{12}	...	p_{1k}	$p_{1.}$
2	p_{21}	p_{22}	...	p_{2k}	$p_{2.}$
⋮	⋮	⋮	...	⋮	⋮
n	p_{n1}	p_{n2}	...	p_{nk}	$p_{n.}$
Total	$p_{.1}$	$p_{.2}$	...	$p_{.k}$	1

Οι διαγώνιοι p_{ii} αντιπροσωπεύουν το ποσοστό της παρατηρηθείσας συμφωνίας μεταξύ των αξιολογητών σε κάθε κατηγορία. Το συνολικό ποσοστό της παρατηρούμενης συμφωνίας είναι:

$$p_o = \sum_{i=1}^k p_{ii}$$

και το συνολικό ποσοστό της αναμενόμενης συμφωνίας είναι:

$$p_e = \sum_{i=1}^k p_{i.} \cdot p_{.i}$$

Η μετρική Kappa, που αντιπροσωπεύει τον βαθμό συμφωνίας των αξιολογητών, υπολογίζεται ως εξής:

$$k = \frac{p_o - p_e}{1.0 - p_e}$$

Στην περίπτωση που υπάρχει πλήρη συμφωνία, η μετρική k θα είναι 1. Γενικά, ο βαθμός συμφωνίας με βάση την μετρική k μπορεί να χαρακτηριστεί ως :

Μη συμφωνία	για $k < 0$
Ισχνή Συμφωνία	για k μεταξύ 0.01-0.20
Καλή συμφωνία	για k μεταξύ 0.21-0.40
Μέτρια συμφωνία	για k μεταξύ 0.41-0.60
Ισχυρή συμφωνία	για k μεταξύ 0.61-0.80
Σχεδόν τέλεια συμφωνία	για k μεταξύ 0.81-0.99

Στα πλαίσια της μελέτης, χρησιμοποιήθηκε η μετρική Cohen's Kappa για να υπολογίσει τον βαθμό αξιοπιστίας της συμφωνίας μεταξύ των εκτιμήσεων του εμπειρογνώμονα και των ταξινομητών στο σύνολο των δεδομένων του πειράματος. Τα αποτελέσματα παρουσιάζονται στον Πίνακα 5.6.

Πίνακας 5.6. Αποτελέσματα Cohen's kappa μετρικής

	Knowledge-based tool	SVM	Naïve Bayes	Majority	Bagging	Boosting
Expert Human Annotator	0.53	0.61	0.58	0.72	0.74	0.78

Τα αποτελέσματα της μελέτης έδειξαν ότι υπήρξε μια μέτρια συμφωνία μεταξύ του μηχανισμού βασισμένου στην γνώση και του εμπειρογνώμονα δεδομένου ότι η μετρική Kappa ήταν $k = 0.53$ (confidence interval 95%, $p < 0.0005$). Επίσης, το ίδιο προέκυψε και για τον ταξινομητή naïve bayes, όπου υπήρξε μια μέτρια συμφωνία και η μετρική k υπολογίστηκε ότι ήταν $k = 0.58$. Όσον αφορά τον ταξινομητή SVM, τα αποτελέσματα έδειξαν ότι υπήρξε ισχυρή συμφωνία με τον εμπειρογνώμονα και η μετρική $k = 0.61$. Από τα αποτελέσματα φαίνεται ότι υπήρξε μια ισχυρή συμφωνία μεταξύ του εμπειρογνώμονα και των ομαδοποιημένων ταξινομητών, όπου η μετρική Kappa υπολογίστηκε ότι ήταν $k = 0.72$ για τον ομαδοποιημένο ταξινομητή που ακολουθεί την αρχή της πλειοψηφίας,

$k = 0.74$ για τον ομαδοποιημένο ταξινομητή που σχηματίστηκε με την μέθοδο bagging και $k = 0.78$ για τον ομαδοποιημένο ταξινομητή που ακολουθεί την μέθοδο boosting.

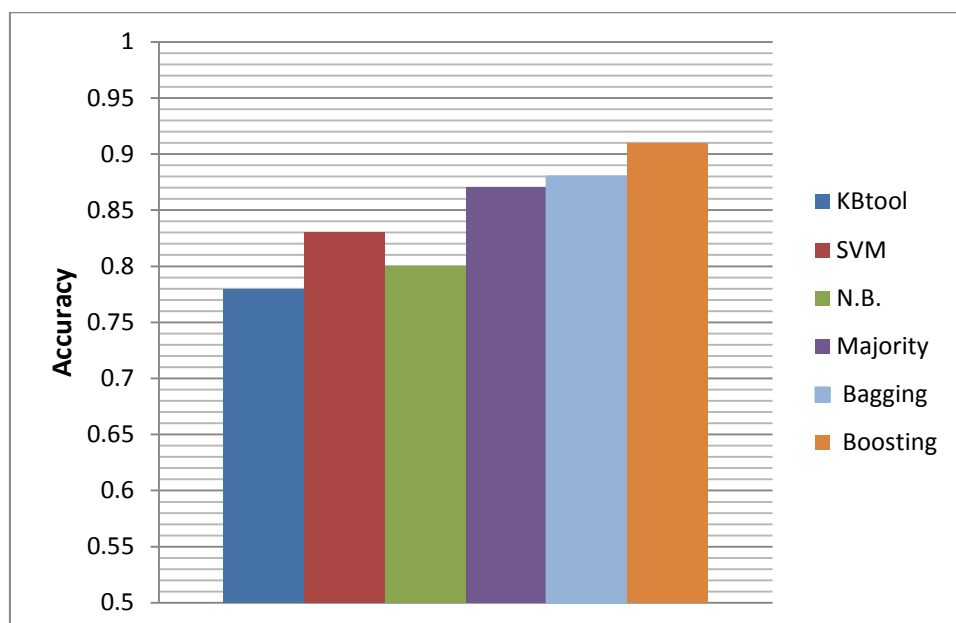
5.6.3 Αξιολόγηση Προσδιορισμού Συναισθηματικής Πολικότητας

Στο δεύτερο μέρος της πειραματικής μελέτης έγινε αξιολόγηση των επιδόσεων των ταξινομητών και των ομαδοποιημένων σχημάτων ταξινόμησης ως προς τον καθορισμό της συναισθηματικής πολικότητας των προτάσεων και τον χαρακτηρισμό τους ως συναισθηματικά θετικές ή αρνητικές. Το πρόβλημα αντιμετωπίστηκε ξανά ως πρόβλημα ταξινόμησης των προτάσεων στην κατάλληλη κατηγορία και είχε ως στόχο να αξιολογήσει την απόδοση τόσο των βασικών ταξινομητών όσο και των μεθόδων ομαδοποίησης ταξινομητών. Στα πλαίσια της πειραματικής μελέτης, επιλέχθηκαν οι προτάσεις του συνόλου δεδομένων που χαρακτηρίστηκαν ως συναισθηματικές από τον ειδικό. Για την αξιολόγηση της απόδοσης υπολογίστηκαν οι μετρικές της ορθότητας (accuracy), ακρίβειας (precision), ευαισθησίας (sensitivity) και εξειδίκευσης (specificity), δεδομένου ότι η κλάση εξόδου είναι μια δυαδική μεταβλητή. Τα αποτελέσματα που προέκυψαν παρουσιάζονται στον Πίνακα 5.7.

Πίνακας 5.7 Η απόδοση των ταξινομητών ως προς τη συναισθηματική πολικότητα

	Knowledge Based Tool	SVM	Naïve Bayes	Majority	Bagging	Boosting
Accuracy	0.78	0.83	0.80	0.87	0.88	0.90
Precision	0.83	0.89	0.86	0.90	0.91	0.91
Sensitivity	0.79	0.78	0.83	0.89	0.88	0.91
Specificity	0.76	0.87	0.82	0.84	0.85	0.85

Τα αποτελέσματα δείχνουν πολύ καλή απόδοση για τους βασικούς και τους ομαδοποιημένους ταξινομητές. Οι ομαδοποιημένοι ταξινομητές παρουσιάζουν ξανά καθολικά καλύτερη απόδοση σε όλα τα πειράματα σε σχέση με τους βασικούς ταξινομητές. Ο προσδιορισμός της συναισθηματικής πολικότητας των προτάσεων πραγματοποιείται με πολύ καλή επίδοση και από τους τρεις βασικούς ταξινομητές, γεγονός που συμβάλει στην καλύτερη συνολική απόδοση των ομαδοποιημένων ταξινομητών. Η ακρίβεια των ταξινομητών στο σύνολο δεδομένων της πειραματικής μελέτης παρουσιάζεται στην Εικόνα 5.15.



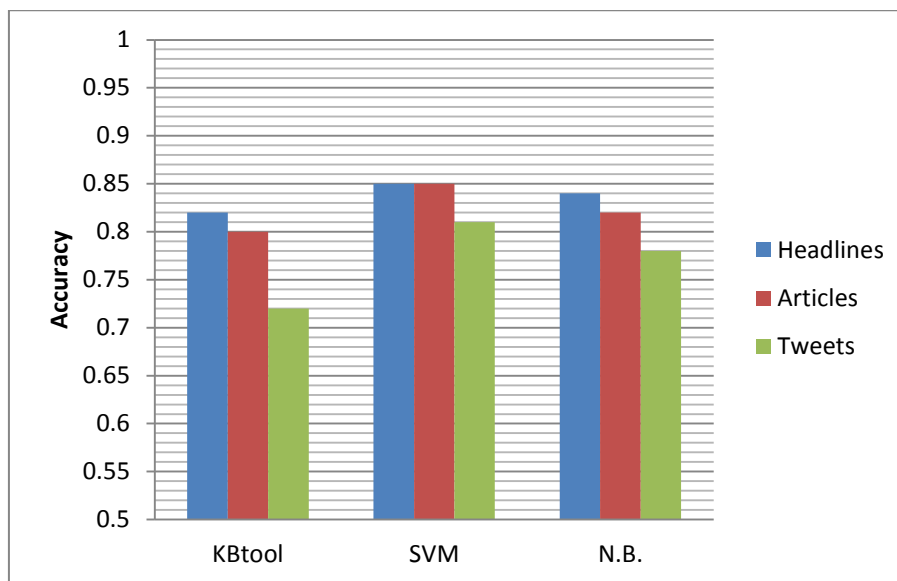
Εικόνα 5.15 Η ορθότητα των ταξινομητών σε όλο το σύνολο δεδομένων

Τα αναλυτικά αποτελέσματα των επιδόσεων των βασικών ταξινομητών και των ομαδοποιημένων σχημάτων για κάθε τύπο κειμενικών δεδομένων παρουσιάζονται στον Πίνακα 5.8 και 5.9 αντίστοιχα.

Πίνακας 5.8 Ανάλυση της απόδοσης των βασικών ταξινομητών σε κάθε είδος κειμενικών δεδομένων

	Headlines			Articles			Tweets		
	KBtool	SVM	N.B.	KBtool	SVM	N.B.	KBtool	SVM	N.B.
Accuracy	0.82	0.85	0.84	0.80	0.85	0.82	0.72	0.81	0.78
Precision	0.87	0.93	0.89	0.77	0.89	0.85	0.84	0.88	0.86
Sensitivity	0.84	0.85	0.83	0.82	0.76	0.80	0.71	0.77	0.87
Specificity	0.79	0.90	0.87	0.74	0.86	0.85	0.70	0.85	0.77

Τα αποτελέσματα δείχνουν ότι οι τρεις βασικοί ταξινομητές έχουν πολύ καλή απόδοση στην αναγνώριση της συναισθηματικής πολικότητας των προτάσεων. Η καλύτερη απόδοση επιτυγχάνεται από τον SVM, όπου η ορθότητα και η ακρίβειά του είναι καλύτερες από εκείνες του μηχανισμού βασισμένου στην γνώση (knowledge based), καθώς και του ταξινομητή Naïve Bayes. Η ορθότητα των βασικών ταξινομητών στα διάφορα είδη κειμενικών δεδομένων παρουσιάζονται στην Εικόνα 5.16.



Εικόνα 5.16 Η ορθότητα των βασικών ταξινομητών σε κάθε είδος δεδομένων κειμένου

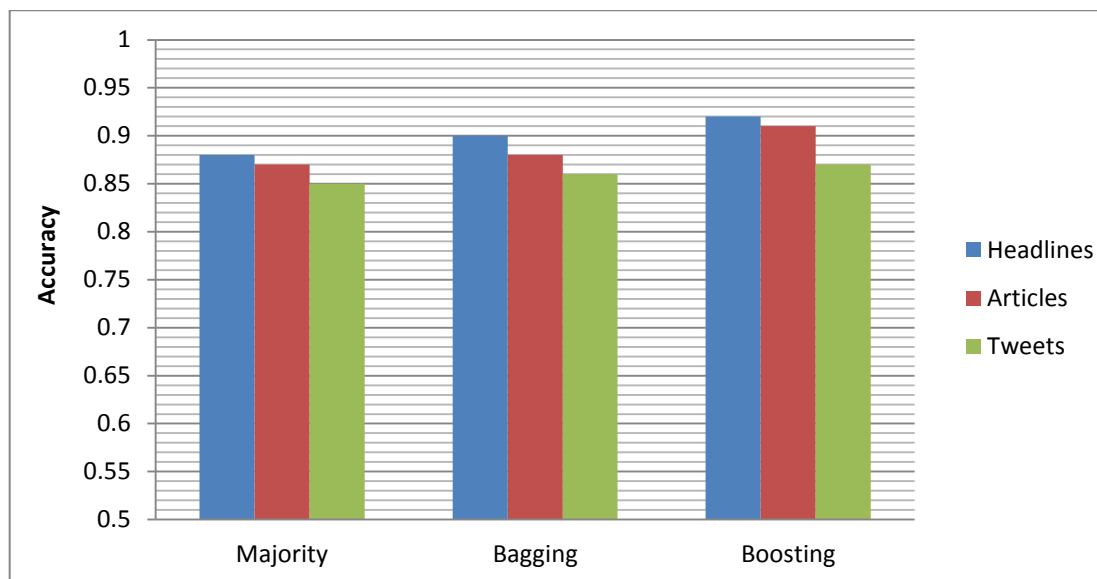
Ο μηχανισμός που βασίζεται στην γνώση (knowledge-based) παρουσιάζει καλύτερη απόδοση σε τίτλους και σε άρθρα, ενώ η απόδοσή του μειώνεται σε posts χρηστών στο Twitter. Κύριος λόγος για αυτήν την απόδοση αποτελεί ο τρόπος που το εργαλείο αναλύει και διαχειρίζεται συντακτικά τις προτάσεις. Πιο συγκεκριμένα, στην περίπτωση των άρθρων και των τίτλων, οι προτάσεις είναι ορθά συνταγμένες και η ακριβής ανάλυση των συντακτικών εξαρτήσεων συμβάλει στον ακριβή καθορισμό του συναισθηματικού περιεχομένου καθώς και στην ποσοτικοποίησή του. Αντιθέτως, στην περίπτωση των posts χρηστών στο Twitter, οι προτάσεις περιέχουν αναφορικές εκφράσεις και αμφισβητούμενες γεγονόσεις που δυσκολεύει την ακριβή μορφολογική ανάλυσή τους και συνεπώς την δημιουργία συναισθηματικών δομών. Επίσης, το εργαλείο σε σχέση με τις στατιστικές προσεγγίσεις διαχειρίζεται τις αρνήσεις πιο αναλυτικά και συστηματικά και είναι σε θέση να αναγνωρίζει αρνήσεις που αντιστρέφουν την σημασιολογία των λέξεων και την συναισθηματική πολικότητα των προτάσεων. Αξιοσημείωτο σημείο της απόδοσης του μηχανισμού βασισμένου στη γνώση αποτελεί η ισορροπία της απόδοσής του στον καθορισμό της θετικής και αρνητικής συναισθηματικής πολικότητας, κάτι που φαίνεται από το ίδιο επίπεδο των μετρικών της ευαισθησίας και εξειδίκευσης. Μεταξύ των βασικών ταξινομητών, ο ταξινομητής SVM παρουσιάζει τα καλύτερα αποτελέσματα σε όλα τα είδη κειμενικών δεδομένων, έχοντας υψηλότερη ακρίβεια σε τίτλους ειδήσεων από ότι σε δεδομένα χρηστών στο twitter. Οι ομαδοποιημένοι ταξινομητές παρουσιάζουν και πάλι καθολικά καλύτερα αποτελέσματα σε σχέση με τους βασικούς ταξινομητές σε όλα τα είδη

κειμενικών δεδομένων. Η ανάλυση της απόδοσης των ομαδοποιημένων ταξινομητών σε κάθε είδος κειμενικών δεδομένων παρουσιάζεται στον Πίνακα 5.9.

Πίνακας 5.9 Ανάλυση της απόδοσης των ομαδοποιημένων ταξινομητών σε κάθε είδος κειμενικών δεδομένων

	Headlines			Articles			Tweets		
	Majority	Bagging	Boosting	Majority	Bagging	Boosting	Majority	Bagging	Boosting
Accuracy	0.88	0.90	0.92	0.87	0.88	0.91	0.85	0.86	0.87
Precision	0.95	0.92	0.93	0.90	0.92	0.91	0.85	0.88	0.88
Sensitivity	0.90	0.91	0.91	0.90	0.88	0.93	0.86	0.86	0.89
Specificity	0.88	0.87	0.90	0.83	0.88	0.85	0.82	0.80	0.80

Όσον αφορά τις μεθόδους δημιουργίας των ομαδοποιημένων ταξινομητών, η μέθοδος boosting παρουσιάζει καλύτερα αποτελέσματα από την μέθοδο bagging και την αρχή της πλειοψηφίας και ο ομαδοποιημένος ταξινομητής που δημιουργήθηκε με την μέθοδο boosting παρουσιάζει καθολικά καλύτερα αποτελέσματα τόσο σε σχέση με τους βασικούς ταξινομητές όσο και με τους ομαδοποιημένους ταξινομητές που ακολουθούν άλλες μεθόδους ομαδοποίησης σε όλα τα είδη κειμενικών δεδομένων. Το γεγονός είναι αποτέλεσμα της έμφασης που δίνει η μέθοδος boosting σε δεδομένα που είναι πιο δύσκολο να ταξινομηθούν σωστά. Το ομαδοποιημένο σχήμα που συνδυάζει τους βασικούς με την αρχή της πλειοψηφίας βελτιώνει την απόδοση της αναγνώρισης της πολικότητας του συναισθηματικού περιεχομένου και παρουσιάζει καλύτερα αποτελέσματα από τον καλύτερο βασικό ταξινομητή.



Εικόνα 5.17 Η ορθότητα των ομαδοποιημένων ταξινομητών σε κάθε είδος δεδομένων κειμένου.

Επίσης, έγινε σύγκριση της απόδοσης των ομαδοποιημένων ταξινομητών με τον καλύτερο βασικό ταξινομητή με στόχο να καθοριστεί το ποσοστό της βελτίωσης της απόδοσης της ταξινόμησης από τα ένθετα σχήματα. Από τα αποτελέσματα προκύπτει ότι ο ομαδοποιημένος ταξινομητής που ακολουθεί την αρχή της πλειοψηφίας παρουσιάζει ορθότητα η οποία είναι υψηλότερη κατά 4.8% σε σχέση με την χρήση του ταξινομητή SVM. Επίσης, η ακρίβεια της μεθόδου bagging είναι υψηλότερη του SVM κατά 6.0%, ενώ η ακρίβεια της μεθόδου boosting είναι υψηλότερη κατά 8.4% αντιστοίχως.

Επιπροσθέτως, έγινε υπολογισμός της συμφωνίας μεταξύ των προσδιορισμών των ταξινομητών και του εμπειρογνώμονα ειδικού με βάση την μετρική Cohen's Kappa. Η μετρική Cohen's Kappa υπολογίστηκε στο σύνολο των δεδομένων για να προσδιορίσει αν υπήρχε και σε ποιο βαθμό συμφωνία μεταξύ των ταξινομητών και του εμπειρογνώμονα. Τα αποτελέσματα της μετρικής Cohen's Kappa για τους ομαδοποιημένους ταξινομητές παρουσιάζονται στον Πίνακα 5.10.

Πίνακας 5.10. Cohen's kappa μετρική

	Knowledge-based tool	SVM	Naïve Bayes	Majority	Bagging	Boosting
Expert Human Annotator	0.55	0.60	0.57	0.70	0.72	0.75

Από την ανάλυση των αποτελεσμάτων προέκυψε ότι μεταξύ του μηχανισμού βασισμένου στην γνώση και των προσδιορισμών του εμπειρογνώμονα υπήρξε μια μέτρια συμφωνία δεδομένου ότι η μετρική Kappa ήταν $k = 0.55$ (confidence interval 95%, $p < 0.0005$). Επίσης, το ίδιο προέκυψε και για τον ταξινομητή naïve bayes, όπου υπήρξε μια μέτρια συμφωνία και η μετρική Kappa υπολογίστηκε ότι ήταν $k = 0.57$. Όσον αφορά τον ταξινομητή SVM, τα αποτελέσματα έδειξαν ότι υπήρξε μέτρια συμφωνία με τον εμπειρογνώμονα και η μετρική Kappa ήταν $k = 0.60$. Από τα αποτελέσματα φαίνεται ότι υπήρξε μια ισχυρή συμφωνία μεταξύ του εμπειρογνώμονα και των ομαδοποιημένων ταξινομητών, όπου η μετρική Kappa υπολογίστηκε ότι ήταν $k = 0.70$ για τον ομαδοποιημένο ταξινομητή που ακολουθεί την αρχή της πλειοψηφίας, $k = 0.72$ για τον ομαδοποιημένο ταξινομητή που σχηματίστηκε με την μέθοδο bagging και $k = 0.76$ για τον ομαδοποιημένο ταξινομητή που ακολουθεί την μέθοδο boosting.

5.6.4 Αξιολόγηση Αναγνώρισης Συναισθηματικών Καταστάσεων

Στο τρίτο μέρος της πειραματικής μελέτης έγινε αξιολόγηση των επιδόσεων των βασικών ταξινομητών στην αναγνώριση του συναισθηματικού περιεχομένου των προτάσεων στην κλίμακα του μοντέλου Ekman. Το πρόβλημα αντιμετωπίστηκε ξανά ως πρόβλημα ταξινόμησης των προτάσεων στην κατάλληλη κατηγορία, σύμφωνα με το μοντέλο, και είχε ως στόχο να αξιολογηθεί η απόδοση των βασικών ταξινομητών. Στα πλαίσια της πειραματικής μελέτης, επιλέχθηκαν οι προτάσεις του συνόλου δεδομένων που χαρακτηρίστηκαν ως συναισθηματικές από τον εμπειρογνώμονα-ειδικό. Για την αξιολόγηση της απόδοσης υπολογίστηκαν οι μετρικές της ακρίβειας (precision), της ανάκλησης (recall) καθώς και η μετρική F-measure, δεδομένου ότι η έξοδος αποτελείται από περισσότερες από δύο κλάσεις. Τα γενικά αποτελέσματα της απόδοσης των ταξινομητών που προέκυψαν από την πειραματική μελέτη παρουσιάζονται στον Πίνακα 5.11.

Πίνακας 5.11 Συνολική Απόδοση των ταξινομητών

	Knowledge Based Tool	SVM	Naïve Bayes
Precision	0.54	0.64	0.58
Recall	0.54	0.63	0.57
F-measure	0.54	0.63	0.57

Από τα αποτελέσματα φαίνεται ότι ο SVM έχει την καλύτερη απόδοση μεταξύ των ταξινομητών στο σύνολο των κειμενικών δεδομένων. Στην συνέχεια, στον Πίνακα 5.12 παρουσιάζονται τα αναλυτικά αποτελέσματα της απόδοσης των ταξινομητών ανά είδος κειμενικών δεδομένων.

Πίνακας 5.12 Ανάλυση της απόδοσης των βασικών ταξινομητών σε κάθε είδος κειμενικών δεδομένων.

	Headlines			Articles			Tweets		
	KBtool	SVM	N.B.	KBtool	SVM	N.B.	KBtool	SVM	N.B.
Precision	0.60	0.65	0.63	0.55	0.61	0.58	0.50	0.60	0.54
Recall	0.60	0.66	0.62	0.54	0.60	0.56	0.49	0.60	0.55
F-measure	0.60	0.66	0.63	0.54	0.61	0.57	0.50	0.60	0.54

Τα αποτελέσματα δείχνουν ότι οι τρεις βασικοί ταξινομητές έχουν πολύ καλή απόδοση στην αναγνώριση της συναισθηματικής πολικότητας των προτάσεων. Η καλύτερη απόδοση επιτυγχάνεται από τον ταξινομητή SVM, όπου η F-measure είναι καλύτερη από εκείνες του μηχανισμού βασισμένου στην γνώση (knowledge based) και του ταξινομητή Naive Bayes σε κάθε είδος κειμενικών δεδομένων. Στην συνέχεια, έγινε προσδιορισμός των επιδόσεων των ταξινομητών στην αναγνώριση κάθε ενός βασικού συναισθήματος, σύμφωνα με το μοντέλο Ekman, και τα αποτελέσματα παρουσιάζονται στον Πίνακα 5.13.

Πίνακας 5.13 Ανάλυση της απόδοσης των ταξινομητών σε κάθε μια συναισθηματική κατηγορία.

	Knowledge-based tool			SVM			Naïve Bayes		
	Precision	Recall	F-measure	Precision	Recall	F-measure	Precision	Recall	F-measure
Anger	0.49	0.45	0.47	0.57	0.57	0.57	0.55	0.55	0.55
Disgust	0.49	0.53	0.51	0.60	0.59	0.59	0.55	0.59	0.57
Fear	0.47	0.43	0.45	0.56	0.50	0.53	0.58	0.52	0.55
Joy	0.74	0.78	0.76	0.74	0.77	0.75	0.66	0.68	0.67
Sadness	0.71	0.62	0.66	0.73	0.72	0.73	0.66	0.59	0.62
Surprise	0.44	0.50	0.47	0.53	0.57	0.55	0.51	0.60	0.56

Από την πλευρά των συναισθημάτων, οι ταξινομητές παρουσιάζουν την καλύτερη απόδοση τους στα συναισθήματα της χαράς (Joy) και της λύπης (Sadness). Ένα αξιοσημείωτο στοιχείο αφορά την απόδοση του εργαλείου βασισμένου στην γνώση, του οποίου η απόδοσή στις δυο αυτές κατηγορίες είναι καλύτερη από τον ταξινομητή Naive Bayes και περίπου στο ίδιο επίπεδο με τον ταξινομητή SVM στο συναίσθημα της χαράς. Η απόδοση οφείλεται και στο γεγονός ότι τα συναισθήματα της χαράς και της λύπης εκφράζονται στις περισσότερες περιπτώσεις με έντονα συναισθηματικά φορτισμένες λέξεις. Επίσης, μεταξύ των ταξινομητών, την καλύτερη απόδοση στα συναισθήματα αυτά παρουσιάζει το εργαλείο που βασίζεται στην γνώση, ιδίως στους τίτλους ειδήσεων, κάτι που οφείλεται στην ακριβή ποσοτικοποίηση του περιεχομένου από το εργαλείο και την αναγνώριση τόσο της ύπαρξης του συγκεκριμένου συναισθήματος όσο και της απουσίας του μέσω της αντιστροφής του συναισθηματικού περιεχομένου του. Από την άλλη πλευρά, η απόδοση του εργαλείου που βασίζεται στην γνώση μειώνεται στην αναγνώριση των συναισθημάτων της έκπληξης και του φόβου, κυρίως λόγω του ότι στις περισσότερες περιπτώσεις μπορεί να εκφράζονται στις προτάσεις περιφραστικά και χωρίς την ύπαρξη ισχυρά συναισθηματικών λέξεων. Η ίδια απόδοση ισχύει και για τον ταξινομητή SVM, που παρουσιάζει χαμηλή ακρίβεια στα δύο αυτά συναισθήματα.

Επίσης, έγινε υπολογισμός της συμφωνίας μεταξύ των ταξινομητών με βάση την μετρική Cohen's Kappa statistic. Η μετρική kappa υπολογίστηκε στο σύνολο των δεδομένων για να προσδιορίσει αν υπήρχε και σε ποιο βαθμό συμφωνία μεταξύ των ταξινομητών και του εμπειρογνώμονα. Τα αποτελέσματα παρουσιάζονται στον Πίνακα 5.14.

Πίνακας 5.14. Cohen's kappa μετρική

	Knowledge-based tool	SVM	Naïve Bayes
Expert Human Annotator	0.50	0.57	0.52

Από την ανάλυση των αποτελεσμάτων προέκυψε ότι μεταξύ του μηχανισμού βασισμένου στην γνώση και των προσδιορισμών του εμπειρογνώμονα υπήρξε μια μέτρια συμφωνία δεδομένου ότι η μετρική Kappa ήταν $k = 0.50$ (confidence interval 95%, $p < 0.0005$).

Επίσης, το ίδιο προέκυψε και για τον ταξινομητή Naïve Bayes, όπου υπήρξε μια μέτρια συμφωνία και η μετρική k υπολογίστηκε ότι ήταν $k = 0.52$. Όσον αφορά τον ταξινομητή SVM, τα αποτελέσματα έδειξαν ότι και δω υπήρξε μέτρια συμφωνία με τον εμπειρογνώμονα, καθώς η τιμή της μετρικής ήταν $k = 0.57$.

Τα αποτελέσματα δείχνουν μια πολύ καλή απόδοση των ταξινομητών και των ομαδοποιημένων σχημάτων. Τα βασικά πλεονεκτήματα της προσέγγισής μας σε σχέση με την διεθνή βιβλιογραφία εντοπίζονται στις εξής συνιστώσες. Αρχικά, το εργαλείο που βασίζεται στην γνώση, αποτελεί ένα γενικό μηχανισμό αναγνώρισης του συναισθηματικού περιεχομένου του κειμένου ο οποίος δεν απαιτεί εκπαίδευση σε σύνολο δεδομένων και μπορεί να χρησιμοποιηθεί γενικά, άμεσα και αποτελεσματικά σε πλήθος διαφόρων ειδών κειμενικών δεδομένων. Επίσης, ένα δεύτερο πλεονέκτημα αφορά την χρήση ομαδοποιημένων σχημάτων ταξινόμησης, γεγονός που προσφέρει μεγαλύτερη κλιμάκωση σε νέα ετερογενή δεδομένα τα οποία έχουν προέλθει από διαφορετικές πηγές. Αξίζει να σημειωθεί ότι, η χρήση των ομαδοποιημένων σχημάτων ταξινόμησης, εκτός από το γεγονός ότι επιτυγχάνει καλύτερη απόδοση από την χρήση ενός μόνο ταξινομητή, παρουσιάζει καλύτερη γενίκευση και, λόγω της φύσης τους, καλύτερη σταθερότητα σε νέα δεδομένα. Επίσης, αξίζει να αναφέρουμε ότι στην διεθνή βιβλιογραφία, οι περισσότερες εργασίες αφορούν την εκπαίδευση και την χρήση μιας μεθόδου αναγνώρισης του συναισθηματικού περιεχομένου του κειμένου και συνήθως σε ένα πεδίο. Η προσέγγισή μας, λόγω της γενικής φύσης του μηχανισμού του βασισμένου σε γνώση καθώς και της λειτουργίας των ομαδοποιημένων σχημάτων, μπορεί να χρησιμοποιηθεί άμεσα σε γενικής φύσης κειμενικά δεδομένα παρουσιάζοντας καλή απόδοση. Επιπροσθέτως, οι περισσότερες προσεγγίσεις της διεθνούς βιβλιογραφίας βασίζονται στην χρήση ενός ταξινομητή ο οποίος εκπαιδεύεται ταυτόχρονα σε συναισθηματικά και ουδέτερα κειμενικά δεδομένα με τελικό στόχο την αναγνώριση του συναισθηματικού περιεχομένου του κειμένου. Μια λειτουργική διαφοροποίηση της προσέγγισής μας σε σχέση με προσεγγίσεις της διεθνούς βιβλιογραφίας αποτελεί, η χρήση μηχανισμών για τον προσδιορισμό αρχικά της ύπαρξης ή όχι συναισθηματικού περιεχομένου στα κειμενικά δεδομένα και στην συνέχεια για την αναγνώριση των συγκεκριμένων συναισθημάτων που περιέχονται στο κείμενο, γεγονός που σε συνδυασμό με τα προαναφερόμενα χαρακτηριστικά των μηχανισμών συμβάλει στην καλή απόδοση και κλιμάκωση σε νέα κειμενικά δεδομένα.

5.7 Πλαίσιο Συναισθηματικής Μοντελοποίησης Κειμενικής Συζήτησης

5.7.1 Ορισμός του Πλαισίου

Στην ενότητα αυτή παρουσιάζουμε μια προσέγγιση για την μοντελοποίηση του συναισθηματικού περιεχομένου μιας κειμενικής συζήτησης με στόχο να αναπαρασταθεί η συζήτηση σε συναισθηματικό επίπεδο. Το πλαίσιο που σχεδιάστηκε παρέχει ένα δομημένο τρόπο αναπαράστασης των συναισθηματικών καταστάσεων μια συζήτησης, με βάση τις επιμέρους αναλύσεις του συναισθηματικού περιεχομένου κάθε μέρους της συζήτησης. Η αναπαράσταση ακολουθεί μια γραφοειδή προσέγγιση, όπου οι κόμβοι του συναισθηματικού γράφου που δημιουργείται αναπαριστούν τις συναισθηματικές καταστάσεις και οι ακμές αναπαριστούν τις μεταβάσεις από μια συναισθηματική κατάσταση σε μια άλλη κατά την διάρκεια της συζήτησης (βλ. Εικόνα 5.18). Πιο συγκεκριμένα, η συναισθηματική μοντελοποίηση μια συζήτησης αναπαρίσταται με ένα κατευθυνόμενο γράφημα $G=(V,E)$, όπου οι κορυφές $V_1, V_2, \dots, V_n$ αναπαριστούν τις δυνατές συναισθηματικές καταστάσεις, με βάση το συναισθηματικό μοντέλο που ακολουθείται και τους μηχανισμούς αναγνώρισης των συναισθημάτων αυτών. Τα μεγέθη των κόμβων αναπαριστούν την κατανομή και το ποσοστό των αντίστοιχων συναισθημάτων στα κειμενικά δεδομένα της συζήτησης. Μια κατευθυνόμενη ακμή από ένα κόμβο του γραφήματος V_1 προς έναν άλλο V_2 , αναπαριστά μια συναισθηματική μετάβαση κατά την συζήτηση: από την συναισθηματική κατάσταση που αναπαριστά ο κόμβος V_1 στην συναισθηματική κατάσταση V_2 . Το βάρος της ακμής αναπαριστά το βαθμό της αλληλεπίδρασης μεταξύ των δύο αυτών συναισθημάτων και υπολογίζεται ως:

$$W_{V_1 \rightarrow V_2} = \frac{\sum_{i=1}^m C_{part_i} * val_i}{\sum_{j=1}^k C_{part_j} * val_j}$$

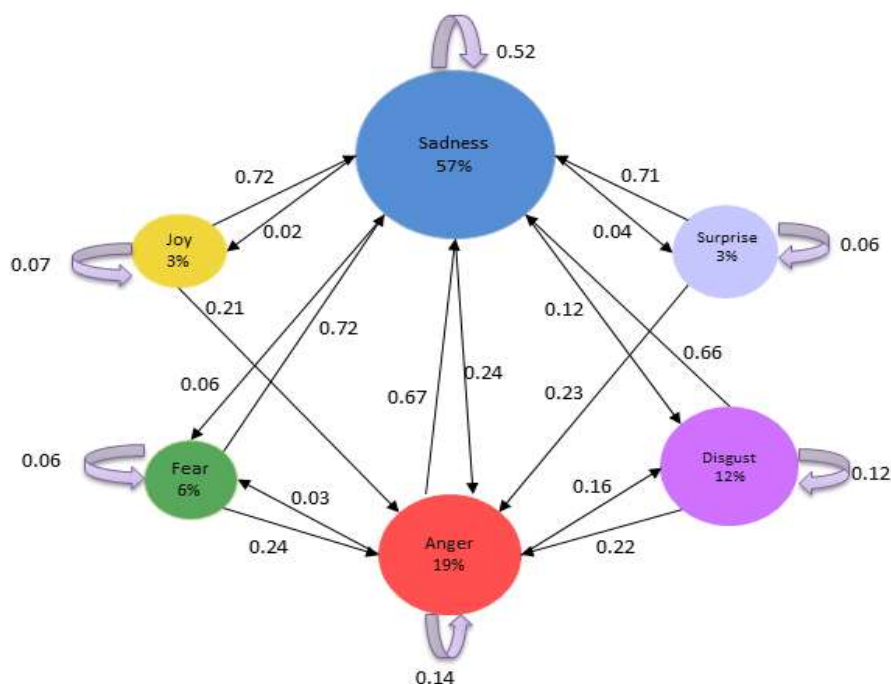
όπου C_{part} αναπαριστά ένα μέρος συζήτησης που αναπαριστά κείμενο χρήστη σε μια συζήτηση. Η παράμετρος k αναπαριστά το πλήθος των μερών της συζήτησης, που αναγνωρίστηκαν ότι έχουν το συναισθηματικό περιεχόμενο που αναπαριστά ο κόμβος V_1 . Η παράμετρος val αναπαριστά την ένταση του συναισθήματος στο μέρος της συζήτησης C_{part} . Η παράμετρος m αναπαριστά το πλήθος των μερών της συζήτησης που αναγνωρίστηκαν ότι φέρουν το συναισθηματικό περιεχόμενο που αναπαριστά ο κόμβος V_2 και τα οποία έγιναν σε απάντηση μέρους της συζήτησης που έφερε το συναισθήμα V_1 .

Στα πλαίσια της χρησιμοποίησης του μοντέλου Ekman και την αναγνώριση του συναισθηματικού περιεχομένου σε προτάσεις φυσικής γλώσσας ως προς αυτό, ο συναισθηματικός γράφος που κατασκευάζεται για κάθε θεματική συζήτηση αποτελείται από 6 κόμβους. Κάθε κόμβος αναπαριστά μια συναισθηματική κατάσταση με βάση το μοντέλο Ekman και επομένως οι κόμβοι είναι οι εξής: “anger”, “fear”, “disgust”, “joy”, “sadness” και “surprise”. Στην δημιουργία του συναισθηματικού γράφου, ο χρωματισμός των κορυφών ακολουθεί τον χαρακτηρισμό των συναισθηματικών καταστάσεων του τροχού συναισθημάτων του Plutchik.

5.7.2 Εφαρμογή του Πλαισίου

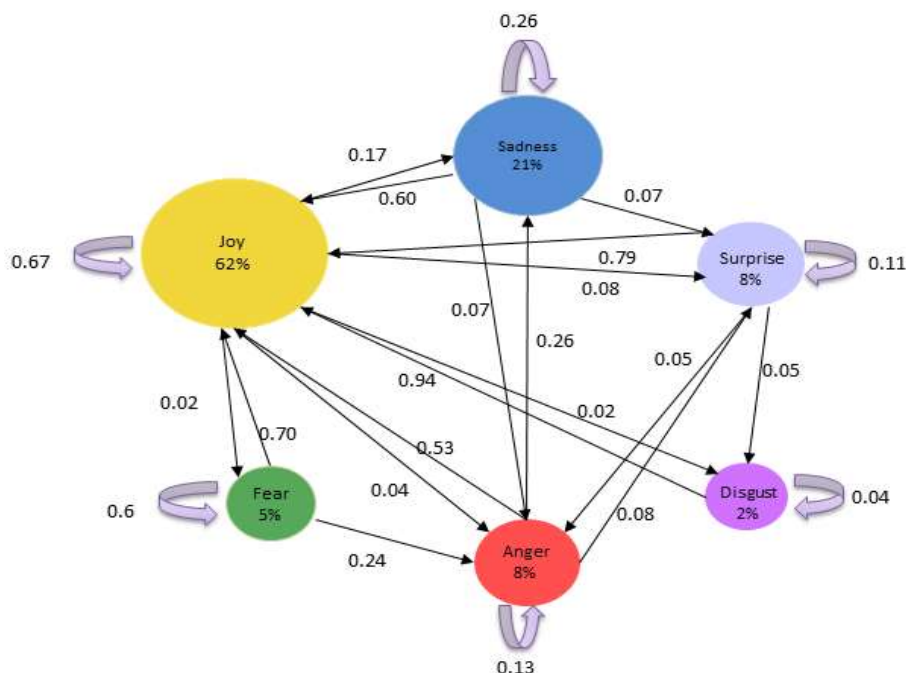
Το πλαίσιο της μοντελοποίησης της συναισθηματικής πληροφορίας μελετήθηκε και εφαρμόστηκε σε συζητήσεις χρηστών σε κοινωνικά δίκτυα σε διάφορες θεματικές περιοχές. Αρχικά, έγινε ανάκτηση των συζητήσεων των χρηστών για θεματικές περιοχές και στην συνέχεια έγινε ανάλυση των μερών της συζήτησης και αναγνώριση του συναισθηματικού περιεχομένου που φέρουν. Μετά την διαδικασία αναγνώρισης του περιεχομένου των μερών της συζήτησης, τα δεδομένα αναλύονται περαιτέρω και δημιουργείται αρχικά το γράφημα κατανομών συναισθημάτων της συζήτησης και στην συνέχεια ο συναισθηματικός γράφος. Στα πλαίσια της εφαρμογής του πλαισίου, δεδομένου ότι γίνεται προσδιορισμός της ύπαρξης των συναισθημάτων, η τιμή του βάρους κάθε ακμής του γραφήματος μπορεί να παίρνει τιμές στο πεδίο $[0,1]$, όπου μια τιμή κοντά στο 1 δείχνει ότι μέρη της συζήτησης που έφεραν το συναισθηματικό περιεχόμενο του κόμβου V_1 προκάλεσαν μέρη συζήτησης που έφεραν το συναισθηματικό περιεχόμενο του κόμβου V_2 .

Στα πλαίσια της μελέτης μας, έγινε ανάλυση διαφόρων συζητήσεων στο κοινωνικό δίκτυο Twitter και εφαρμογή του πλαισίου στην μοντελοποίηση του συναισθηματικού περιεχομένου του. Πιο συγκεκριμένα, ανακτήθηκαν τα σχόλια χρηστών στην θεματική περιοχή “Syrian civil war” (#syriancivilwar) και στην περιοχή “Valentine’s day” (#valentinesday) προκειμένου να γίνει η συναισθηματική μοντελοποίηση τους. Τα κείμενα αναλύθηκαν αυτόματα και αρχικά προσδιορίστηκε η ύπαρξη συναισθηματικού περιεχομένου από το σχήμα ταξινόμησης boosting και στην συνέχεια τα συναισθηματικά κείμενα αναλύθηκαν και καθορίστηκε το ακριβές συναισθηματικό τους περιεχόμενο από τον ταξινομητή SVM. Τα αποτελέσματα της εφαρμογής του πλαισίου στις δυο θεματικές περιοχές απεικονίζεται στις Εικόνες 5.18 και 5.19 αντίστοιχως.



Εικόνα 5.18 Συναισθηματικό γράφημα της θεματικής περιοχής “Syrian civil war”

Η αναπαράσταση της συζήτησης δείχνει ότι το συναίσθημα της λύπης είναι το επικρατέστερο στην συζήτηση και περιέχεται στο 57% των μερών της συζήτησης, σε αντίθεση με τα συναισθήματα της χαράς και της έκπληξης που περιέχονται στο 3% των μερών της συζήτησης το κάθε ένα. Στις περισσότερες περιπτώσεις που εκφράζεται το συναίσθημα της στενοχώριας στη συζήτηση, ακολουθείται από μέρος που εκφράζει στενοχώρια στο 52% των περιπτώσεων της συζήτησης, από μέρος που εκφράζει θυμό στο 24%, αηδία στο 12%, φόβο στο 6%, έκπληξη στο 4% και χαρά στο 2% των περιπτώσεων.



Εικόνα 5.19 Συναισθηματικό γράφημα της θεματικής περιοχής “Valentine’s day”

Στην αναπαράσταση της συζήτησης της περιοχής “Valentine’s day”, το συναίσθημα της χαράς περιέχεται στο 62% των μερών της συζήτησης και είναι το επικρατέστερο, ενώ σε αντίθεση το συναίσθημα της αηδίας περιέχεται στο 2% των μερών της συζήτησης. Ο συναισθηματικός γράφος μπορεί δυναμικά να μεταβάλλεται χρονικά κατά την εξέλιξη μιας συζήτησης. Η μοντελοποίηση του συναισθηματικού περιεχομένου των συζητήσεων και η αναλυτική παρουσίασή τους με έναν εύληπτο τρόπο μπορεί να παρέχει άμεσα πληροφορίες για την συναισθηματικό επίπεδο της συζήτησης. Σε συνδυασμό με την αυτόματη ανάλυση των προτάσεων φυσικής γλώσσας και τον ακριβή προσδιορισμό του συναισθηματικού τους περιεχομένου σε πραγματικό χρόνο μπορεί να παρέχει πληροφορίες και για την εξέλιξη του συναισθηματικού περιεχομένου μιας συζήτησης καθώς αυτή διαδραματίζεται και να αλλάζει κάθε χρονική στιγμή.

6

Αναγνώριση Συναισθηματικής Κατάστασης Εκφράσεων Προσώπου

6.1 Εισαγωγή

Ένα από τα βασικότερα χαρακτηριστικά των υπολογιστικών συστημάτων είναι η δυνατότητα αλληλεπίδρασης που προσφέρουν στον χρήστη. Για να γίνει η επικοινωνία ανθρώπου - υπολογιστή ποιοτικότερη και πιο φυσική, είναι απαραίτητο τα υπολογιστικά συστήματα να έχουν την ικανότητα να αναγνωρίζουν και να προσαρμόζονται στις καταστάσεις του χρήστη και να διαμορφώνουν κατάλληλα τη συμπεριφορά τους. Μια βασική παράμετρος αφορά την προσαρμογή των υπολογιστικών συστημάτων σε συναισθηματικό επίπεδο, πράγμα που απαιτεί την αναγνώριση της συναισθηματικής κατάστασης του χρήστη. Η αναγνώριση συναισθημάτων από ένα υπολογιστικό σύστημα είναι μια σύνθετη διαδικασία και αφορά κυρίως την προσπάθεια να αναλύσει και να προσδιορίσει την συναισθηματική κατάσταση ενός χρήστη που έρχεται σε επαφή με αυτό.

Η σημασία της αναγνώρισης συναισθημάτων μελετήθηκε από την Rosalind Ricard, που εισήγαγε τον όρο της συναισθηματικής υπολογιστικής (Affective Computing) (Picard, 1997), περιγράφοντας και δημιουργώντας ένα διεπιστημονικό πεδίο που εκτείνεται στην επιστήμη των υπολογιστών, την ψυχολογία και τη γνωστική επιστήμη. Το πεδίο της συναισθηματικής υπολογιστικής σχετίζεται με την μελέτη και ανάπτυξη μεθόδων και συστημάτων τα οποία μπορούν να αναγνωρίσουν, να επεξεργαστούν και να εκφράσουν τα ανθρώπινα συναισθήματα. Δεδομένου ότι οι άνθρωποι επικοινωνούν όχι μόνο με τον λόγο αλλά και έμμεσα μέσω εκφράσεων του προσώπου, μέθοδοι και συστήματα που θα μπορούσαν να εντοπίσουν, να επεξεργαστούν και να αναγνωρίσουν το συναισθηματικό

περιεχόμενο των εκφράσεων του προσώπου θα ήταν ιδιαίτερα χρήσιμες και θα μπορούσαν να συμβάλουν αποτελεσματικά σε διάφορα πεδία και εφαρμογές. Οι εκφράσεις του προσώπου αποτελούν τον πιο εκφραστικό τρόπο με τον οποίο οι άνθρωποι μπορούν να εκδηλώσουν τα συναισθήματά τους και μια παγκόσμια γλώσσα έκφρασης συναισθημάτων, όπου μπορούν να εκφράσουν άμεσα και έντονα ένα ευρύ φάσμα ανθρώπινων συναισθηματικών καταστάσεων. Επιστημονικές έρευνες έχουν δείξει ότι κατά τη διάρκεια της ανθρώπινης επικοινωνίας, μόνο το 7% της πληροφορίας μεταφέρεται με το γλωσσικό μέρος της επικοινωνίας, όπως οι λέξεις που επιλέγονται (verbal part), το 38% της πληροφορίας μεταφέρεται με την ομιλία (paralanguage vocal part) και το 55% μεταφέρεται μέσω των εκφράσεων του προσώπου (Mehrabian, 1968). Συνεπώς, οι εκφράσεις του προσώπου αποτελούν το σημαντικότερο μέσο επικοινωνίας κατά την πρόσωπο-με-πρόσωπο επικοινωνία.

Ο σχεδιασμός μεθόδων και συστημάτων για την επεξεργασία εκφράσεων προσώπου και την αυτόματη αναγνώριση των συναισθημάτων τους αποτελεί μια πολύ σημαντική και δύσκολη διεργασία που βρίσκει εφαρμογή σε ένα ευρύ φάσμα πεδίων. Για παράδειγμα, οι εικονικοί πράκτορες και τα ρομπότ, έχοντας μηχανισμούς που θα τους δίνουν την δυνατότητα να αναγνωρίζουν και να εκφράζουν συναισθήματά, θα συνέβαλλαν ώστε να γίνει πιο φυσική η επικοινωνία μεταξύ του ανθρώπου – υπολογιστή (Alepis, & Virvou, 2011; Clavel, 2015). Στα ευφυή συστήματα διδασκαλίας, τα συναισθήματα και η μάθηση είναι άρρηκτα συνδεδεμένα μεταξύ τους (Shen et al., 2009) και συνεπώς, η αναγνώριση της συναισθηματικής κατάστασης του εκπαιδευόμενου θα μπορούσε να βελτιώσει σημαντικά την αποτελεσματικότητα των μαθησιακών διαδικασιών (Akputu et al., 2013; Finch et al., 2015). Επίσης, στο πεδίο της επίβλεψης (surveillance) και σε εφαρμογές, όπως για παράδειγμα η παρακολούθηση της κατάστασης οδηγού οχήματος και η παρακολούθηση της κατάστασης ηλικιωμένων ατόμων, συστήματα ανάλυσης και προσδιορισμού συναισθηματικής κατάστασης θα ήταν χρήσιμα προσφέροντας την δυνατότητα για καλύτερη κατανόηση της κατάστασης του χρήστη. Ωστόσο, η ανάλυση των εκφράσεων του προσώπου και η ακριβής αναγνώριση του συναισθηματικού περιεχομένου τους αποτελεί μια πολυδιάστατη και δύσκολη διαδικασία (Koutlas & Fotiadis, 2008).

Στα πλαίσια αυτού του κεφαλαίου παρουσιάζουμε μια προσέγγιση για την αυτόματη ανάλυση εκφράσεων προσώπου και την αναγνώριση του συναισθηματικού τους περιεχομένου. Στα πλαίσια της προσέγγισης αυτής, οι εκφράσεις του προσώπου αναλύονται και εξάγονται κατάλληλα χαρακτηριστικά ακολουθώντας μια αναλυτική, *τοπικής-εμβέλειας*

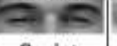
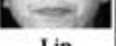
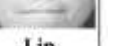
(*local-based*) προσέγγιση. Δοθείσας μιας νέα εικόνας, αρχικά χρησιμοποιείται ο αλγόριθμος Viola-Jones (Viola & Jones, 2004) για να προσδιορίσει την ύπαρξη προσώπων στην εικόνα και μετά την ανίχνευση προσώπου, συγκεκριμένες περιοχές όπως τα μάτια, τα φρύδια και το στόμα αναλύονται εις βάθος και εξάγονται γεωμετρικά χαρακτηριστικά, τα οποία μοντελοποιούν την κατάσταση κάθε μέρους του προσώπου. Μετά την ανάλυση, εξάγονται συνολικά 25 χαρακτηριστικά τα οποία αναπαριστούν την έκφραση του προσώπου. Με βάση τα χαρακτηριστικά που εξάγονται, μέθοδοι ταξινομητών χρησιμοποιούνται για να προσδιορίσουν την ύπαρξη ή όχι συναισθημάτων σε μια έκφραση προσώπου και να αναγνωρίσουν την ακριβή συναισθηματική του κατάσταση.

6.2 Βασικές Έννοιες και Σχετικές Εργασίες

6.2.1 Βασικές έννοιες

6.2.1.1 Facial Action Coding System (FACS)

Το Facial Action Coding System (FACS) (Ekman & Friesen, 1978) είναι ένα μοντέλο περιγραφής των καταστάσεων των εκφράσεων του προσώπου. Το μοντέλο FACS αναπτύχθηκε από τους Paul Ekman και W.V. Friesen το 1978, οι οποίοι καθόρισαν πώς οι συσπάσεις κάθε μυ του προσώπου, μεμονωμένα ή σε συνδυασμό με άλλους μύες, αλλάζουν την εμφάνιση του προσώπου. Το μοντέλο FACS είναι ανατομικά προσανατολισμένο στις εκφράσεις προσώπου και βασίζεται στον ορισμό των Action Units (AU), οι οποίες προκαλούν τις κινήσεις του προσώπου. Πιο συγκεκριμένα, κάθε AU αντιστοιχεί στην ταυτόχρονη δράση μιας ομάδας μυών, οι οποίοι διαμορφώνουν μια συγκεκριμένη δράση στο πρόσωπο. Το μοντέλο επικεντρώνεται στους μύες του προσώπου οι οποίοι, ανεξάρτητα ή κατά ομάδες, προκαλούν αλλαγές στην εμφάνιση του προσώπου. Οι αλλαγές αυτές καθώς και οι μύες που προκάλεσαν τις αλλαγές καλούνται Action Units. Μερικά παραδείγματα βασικών AUs παρουσιάζονται στην Εικόνα 6.1.

Upper Face Action Units					
AU 1	AU 2	AU 4	AU 5	AU 6	AU 7
					
Inner Brow Raiser	Outer Brow Raiser	Brow Lowerer	Upper Lid Raiser	Cheek Raiser	Lid Tightener
*AU 41	*AU 42	*AU 43	AU 44	AU 45	AU 46
					
Lid Droop	Slit	Eyes Closed	Squint	Blink	Wink
Lower Face Action Units					
AU 9	AU 10	AU 11	AU 12	AU 13	AU 14
					
Nose Wrinkler	Upper Lip Raiser	Nasolabial Deepener	Lip Corner Puller	Cheek Puffer	Dimpler
AU 15	AU 16	AU 17	AU 18	AU 20	AU 22
					
Lip Corner Depressor	Lower Lip Depressor	Chin Raiser	Lip Puckerer	Lip Stretcher	Lip Funneler
AU 23	AU 24	*AU 25	*AU 26	*AU 27	AU 28
					
Lip Tightener	Lip Pressor	Lips Part	Jaw Drop	Mouth Stretch	Lip Suck

Εικόνα 6.1. Μερικά AUs του προσώπου και συνδυασμοί τους

Για παράδειγμα, το AU 1 αφορά το σήκωμα του εσωτερικού τμήματος του φρυδιού, το AU2 αφορά το σήκωμα του εξωτερικού τμήματος του φρυδιού και το AU 26 αφορά το κατέβασμα του σαγονιού. Επίσης, στην διεθνή βιβλιογραφία έχουν αναπτυχθεί και άλλα μοντέλα εκφράσεων προσώπου, πολλά από τα οποία έχουν βασιστεί ή αποτελούν παραλλαγές του συστήματος FACS, όπως για παράδειγμα τα FAST (Facial Action Scoring Technique), EMFACS (Emotional Facial Action Coding System), FAST (Facial Action Scoring Technique), MAX (Maximally Discriminative Facial Movement Coding System) και το EMG (Facial Electromyography) (Ekman & Rosenberg, 1997).

6.2.1.2 Παράμετροι απόδοσης κίνησης προσώπου (Facial Animation Parameters -FAPs)

Το μοντέλο FACS ενέπνευσε και τη δημιουργία των παραμέτρων απόδοσης κίνησης προσώπου στο πλαίσιο του προτύπου ISO MPEG-4. Πιο συγκεκριμένα, το σύνολο παραμέτρων απόδοσης κίνησης προσώπου (Facial Animation Parameter, FAPs) έχει σχεδιαστεί για να επιτρέπει τον ορισμό του σχήματος και της κίνησης του προσώπου, την αναπαράσταση εκφράσεων, των συναισθημάτων καθώς και της ομιλίας. Οι παράμετροι κίνησης προσώπου βασίζονται στη μελέτη των κινήσεων του προσώπου και συνδέονται στενά με τις κινήσεις των μυών. Οι παράμετροι κίνησης αποτελούν ένα σύνολο παραμέτρων που αντιπροσωπεύουν τις ενέργειες κίνησης του προσώπου που

περιλαμβάνουν τις κινήσεις του κεφαλιού, του ελέγχου της γλώσσας, των ματιών, του στόματος κ.α. Πιο συγκεκριμένα, κάθε FAP είναι μια ενέργεια του προσώπου που παραμορφώνει ένα μοντέλο προσώπου από την ουδέτερη του κατάσταση. Η τιμή του FAP δείχνει το μέγεθος του FAP, το οποίο με τη σειρά του δείχνει το μέγεθος της παραμόρφωσης που προκαλείται με βάση το ουδέτερο πρόσωπο. Το σύνολο παραμέτρων απόδοσης κίνησης προσώπου (FAPs) είναι βασισμένο στην μελέτη των ελάχιστων δράσεων του προσώπου και είναι στενά συνδεδεμένο με τις δράσεις των μυών. Οι δράσεις, όπως για παράδειγμα η σύμπτυξη των φρυδιών και το άνοιγμα του στόματος, επιτρέπουν την αναπαράσταση των πιο φυσικών μορφοποιήσεων του προσώπου. Για να χρησιμοποιηθούν οι τιμές των FAP σε κάθε μοντέλο προσώπου, το πρότυπο MPEG-4 καθόρισε την κανονικοποίηση των τιμών των FAP. Η κανονικοποίηση έγινε χρησιμοποιώντας τις παραμέτρους Facial Animation Parameter Units (FAPUs), όπου ορίζονται ως το κλάσμα της απόστασης μεταξύ των βασικών χαρακτηριστικών του προσώπου. Οι μονάδες αυτές έχουν σχεδιαστεί έτσι ώστε να επιτρέπουν την απεικόνιση των παραμέτρων FAPs σε κάθε μοντέλο προσώπου με έναν συνεχή τρόπο.

6.2.2 Σχετικές Εργασίες

Στο πεδίο της αναγνώρισης συναισθημάτων από εκφράσεις προσώπου, δύο κύριοι τύποι προσεγγίσεων για την μοντελοποίηση των εκφράσεων είναι οι *ολιστικές μέθοδοι* (holistic methods) και οι *αναλυτικές ή τοπικής-εμβέλειας (local-based) μέθοδοι* (Pantic & Rothkrantz, 2000). Οι ολιστικές μέθοδοι προσπαθούν να μοντελοποιήσουν τις παραμορφώσεις του ανθρώπινου προσώπου αναπαριστώντας ολόκληρο το πρόσωπο σαν μια οντότητα. Από την άλλη πλευρά, οι αναλυτικές μέθοδοι μοντελοποιούν και αναλύουν τοπικές, διακριτές παραμορφώσεις του προσώπου όπως, τα μάτια, τα φρύδια, τη μύτη, το στόμα κ.α., καθώς και τις γεωμετρικές σχέσεις τους, προκειμένου να δημιουργήσουν περιγραφικά και εκφραστικά μοντέλα της έκφρασης (Arca et al., 2004). Επίσης, για την διαδικασία εξαγωγής χαρακτηριστικών και την αναπαράσταση της έκφρασης υπάρχουν δύο κύρια είδη προσεγγίσεων, οι μέθοδοι που βασίζονται σε γεωμετρικά χαρακτηριστικά (geometric-based methods) και οι μέθοδοι που βασίζονται στην εμφάνιση (appearance-based methods). Οι μέθοδοι που βασίζονται στα γεωμετρικά χαρακτηριστικά, εξάγουν χαρακτηριστικά που περιγράφουν και μοντελοποιούν τα γεωμετρικά χαρακτηριστικά του προσώπου και των βασικών περιοχών του, όπως τη θέση και το σχήμα. Οι μέθοδοι που βασίζονται στην εμφάνιση (appearance-based methods) χρησιμοποιούν φίλτρα εικόνας,

όπως φίλτρα Gabor σε όλο το πρόσωπο ή σε συγκεκριμένα μέρη για να εξαχθούν διανύσματα χαρακτηριστικών γνωρισμάτων (Ghimire & Lee, 2013).

Οι συγγραφείς, στην εργασία (Aung & Aye, 2012), παρουσιάζουν μια προσέγγιση που βασίζεται στην μορφολογική ανάλυση της περιοχής του στόματος για την αναγνώριση πέντε συναισθηματικών καταστάσεων εκφράσεων. Πιο συγκεκριμένα, οι συγγραφείς αναγνωρίζουν τις καταστάσεις της χαράς (joy), του θυμού (anger), της στενοχώριας (sadness), της έκπληξης (surprise) καθώς και την ουδέτερη κατάσταση (neutral) βασιζόμενοι στο ιστόγραμμα της περιοχής του στόματος. Στην πειραματική μελέτη της προσέγγισης σε ένα μέρος εκφράσεων προσώπου της βάσης δεδομένων JAFFE (Lyons et al., 1998), η ακρίβεια της μεθόδου ήταν 81.6%. Οι συγγραφείς, στην εργασία (Kittusamy, & Chakrapani, 2012) παρουσιάζουν μια προσέγγιση για την αναγνώριση των 6 βασικών συναισθημάτων του μοντέλου Ekman που βασίζεται στους χώρους ιδιο-διανυσμάτων (eigen spaces). Στην πειραματική μελέτη, οι συγγραφείς αναφέρουν ακρίβεια 79.5% στην βάση δεδομένων Cohn Kanade (Kanade et al., 2000) και 76.6% στην βάση δεδομένων JAFFE. Οι συγγραφείς στην εργασία (Thai et al., 2011) παρουσιάζουν μια προσέγγιση που βασίζεται στην τεχνική αναγνώρισης ακμών Canny (Canny, 1986) και την μέθοδο PCA (principal component analysis) για την ανάλυση των εκφράσεων του προσώπου και την εξαγωγή χαρακτηριστικών. Η αναγνώριση πραγματοποιείται από ένα νευρωνικό δίκτυο με βάση τα χαρακτηριστικά κάθε έκφρασης και στην πειραματική μελέτη, οι συγγραφείς αναφέρουν ακρίβεια 85.7% στην βάση δεδομένων JAFFE. Στην εργασία (Prinosil et al., 2008), οι συγγραφείς παρουσιάζουν μια προσέγγιση που βασίζεται σε τεχνικές ανάλυσης και εξαγωγής χαρακτηριστικών όπως, τα φίλτρα Gabor, την ανάλυση LDA (Linear Discrimination analysis) και την PCA (principal component analysis). Οι συγγραφείς αναφέρουν ακρίβεια 93,8% στην βάση δεδομένων JAFFE και επίσης ακρίβεια 79% στην αναγνώριση χαμογελαστών-μη χαμογελαστών εκφράσεων στην βάση δεδομένων ORL. Στην εργασία (Michel & Kaliouby, 2005), οι συγγραφείς παρουσιάζουν μια προσέγγιση για την αναγνώριση των συναισθημάτων του μοντέλου Ekman σε εκφράσεις προσώπου. Στα πλαίσια της προσέγγισης, γίνεται εξαγωγή 22 χαρακτηριστικών από τα μέρη του προσώπου και στην συνέχεια ένας ταξινομητής SVM χρησιμοποιείται για να αναγνωρίσει την συναισθηματική κατάσταση της νέας έκφρασης. Στην πειραματική μελέτη που έγινε στην βάση δεδομένων Cohn-Kanade, οι συγγραφείς αναφέρουν ακρίβεια 87.9%. Στην εργασία (Kotsia & Pitas, 2007), οι συγγραφείς παρουσιάζουν μια προσέγγιση που βασίζεται στην μοντελοποίηση του προσώπου με πλέγμα που αποτελείται από 104 κόμβους και στην συνέχεια στην εκπαίδευση και χρήση ταξινομητή SVM, για την αναγνώριση του

συναισθηματικού περιεχομένου της έκφρασης. Στην πειραματική μελέτη οι συγγραφείς αναφέρουν ακρίβεια έως 99,7% στην βάση δεδομένων Cohn-Kanade, μια απόδοση που είναι μια από τις καλύτερες στην διεθνή βιβλιογραφία. Επίσης, στην εργασία (Shih et al., 2008), οι συγγραφείς παρουσιάζουν μια προσέγγιση που βασίζεται στην χρήση της τεχνικής linear discriminant analysis (LDA) για την εξαγωγή χαρακτηριστικών για χρήση σε ταξινομητή SVM για την αναγνώριση των συναισθημάτων. Η απόδοση της προσέγγισης στην βάση δεδομένων JAFFE ήταν 94.1% και είναι μια από τις καλύτερες στην διεθνή βιβλιογραφία για την συγκεκριμένη βάση.

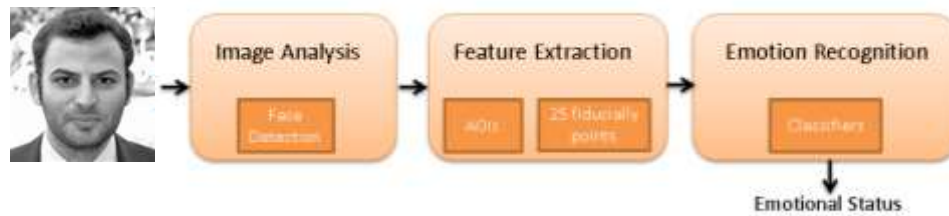
Αξίζει να σημειωθεί ότι τα τελευταία χρόνια προσεγγίσεις ομαδοποίησης ταξινομητών (ensemble classifiers) έχουν προσελκύσει το ενδιαφέρον της ερευνητικής κοινότητας και έχει μελετηθεί η απόδοσή τους στην αναγνώριση συναισθημάτων από εκφράσεις προσώπου. Στην εργασία (Zavaschi et al., 2011), οι συγγραφείς μελετούν την απόδοση ομαδοποιημένων ταξινομητών στην αναγνώριση του συναισθηματικού περιεχομένου στις εκφράσεις της βάσης δεδομένων JAFFE και της Cohn-Kanade και βασικό συμπέρασμα είναι ότι οι ομαδοποιημένοι ταξινομητές βελτιώνουν κατά 5% με 10% την απόδοση των βασικών ταξινομητών.

6.3 Σύστημα Συναισθηματικής Αναγνώρισης Εκφράσεων Προσώπου

Σε αυτή την ενότητα, παρουσιάζουμε την προσέγγιση για την αυτόματη ανάλυση εκφράσεων προσώπου και την αναγνώριση του συναισθηματικού περιεχομένου που φέρουν. Η προσέγγιση μας ακολουθεί μια αναλυτική μέθοδο για την μοντελοποίηση των εκφράσεων και την αναγνώριση των συναισθημάτων του μοντέλου Ekman. Επίσης, ακολουθεί μια γεωμετρική μέθοδο για την εξαγωγή των χαρακτηριστικών της έκφρασης, καθώς οι γεωμετρικές μέθοδοι, αν και είναι υπολογιστικά πιο απαιτητικές, είναι πιο εύρωστες σε διακυμάνσεις της θέσης, του μεγέθους και του προσανατολισμού του προσώπου (Zhang et al., 1998)

Στα πλαίσια της προσέγγισης, αρχικά ανιχνεύονται τα ανθρώπινα πρόσωπα χρησιμοποιώντας τη μέθοδο Viola-Jones και έπειτα γίνεται ανάλυση της έκφρασης του προσώπου, όπου αναλύονται εις βάθος περιοχές όπως τα μάτια, τα φρύδια και το στόμα, και εξάγονται χαρακτηριστικά που μοντελοποιούν την έκφραση του προσώπου. Στην συνέχεια, με βάση τα χαρακτηριστικά που εξάγονται, σχήματα ομαδοποιημένων ταξινομητών

προσδιορίζουν τη συναισθηματική κατάσταση της έκφρασης. Τα βασικά στάδια της διαδικασίας παρουσιάζονται στην Εικόνα 6.2.



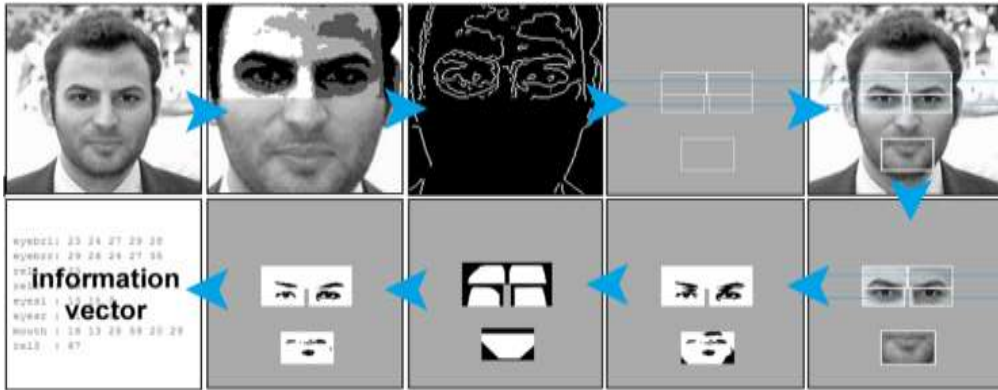
Εικόνα 6.2 Διάγραμμα Ροής του Συστήματος

Πιο συγκεκριμένα, όταν ένα πρόσωπο ανιχνεύεται στην εικόνα, η περιοχή του προσώπου εξάγεται και γίνεται περαιτέρω ανάλυση του. Αρχικά, η αντίθεση (contrast) ενισχύεται και στην συνέχεια, γίνεται προσδιορισμός συγκεκριμένων περιοχών του προσώπου, οι οποίες συμβάλλουν στην αναγνώριση του συναισθηματικού περιεχόμενου της έκφρασης του. Οι συγκεκριμένες περιοχές ονομάζονται *περιοχές ενδιαφέροντος* (Areas of Interests-AOIs) και είναι, η περιοχή των ματιών, η περιοχή του στόματος και η περιοχή των φρυδιών, και οι οποίες αναλύονται περαιτέρω και εξάγονται κατάλληλα χαρακτηριστικά. Η διαδικασία εξαγωγής χαρακτηριστικών αναλύει τις AOIs και εξαγάγει κατάλληλα χαρακτηριστικά, τα οποία περιγράφουν την κατάσταση του κάθε μέρους του προσώπου. Το σύστημα ακολουθεί μια αναλυτική τοπικής-εμβέλειας προσέγγιση και έτσι η εξαγωγή των χαρακτηριστικών γίνεται στις περιοχές ενδιαφέροντος του προσώπου. Έπειτα, η αναγνώριση των συναισθηματικών καταστάσεων πραγματοποιείται από ταξινομητές, που προσδιορίζουν το κατάλληλο συναισθηματικό περιεχόμενο της έκφρασης του προσώπου. Οι βασικοί ταξινομητές που χρησιμοποιούνται είναι: ένας support vector machine (SVM), ένας neuro-fuzzy και ένας random forest. Μελετάται η απόδοσή κάθε ενός καθώς και της ομαδοποίησής τους σε ένα σχήμα ταξινόμησης που χρησιμοποιεί πλειοψηφική ψηφοφορία (majority voting).

6.3.1 Ανάλυση Προσώπου

Όπως προαναφέραμε, στα πλαίσια της ανάλυσης της εικόνας και της αναγνώρισης προσώπου, χρησιμοποιείται ο αλγόριθμος Viola-Jones για να ανιχνεύσει τα πρόσωπα που ενδεχομένως να απεικονίζονται στην εικόνα. Σε περίπτωση που περισσότερα από ένα πρόσωπα ανιχνευθούν, κάθε ένα πρόσωπο αναλύεται ξεχωριστά. Αρχικά, με στόχο την ενίσχυση της αντίθεσης, γίνεται ισοστάθμιση ιστογράμματος (histogram equalization), με βάση το ιστόγραμμα της εικόνας του προσώπου, και στη συνέχεια εφαρμόζονται φίλτρα

sober (Sober, 1990) για να τονιστούν οι άκρες του προσώπου. Έπειτα, χωρίζεται η περιοχή του προσώπου σε οριζόντιες ζώνες και γίνεται υπολογισμός της πυκνότητας πληροφορίας σε κάθε ζώνη. Η περιοχή με την μεγαλύτερη πυκνότητα περιέχει τα μάτια και στην συνέχεια χρησιμοποιούνται φίλτρα για να προσδιοριστούν τα φρύδια, επάνω από τα μάτια, καθώς και η περιοχή του στόματος κάτω από αυτά. Στην Εικόνα 6.3 παρουσιάζονται τα στάδια της ανάλυσης ενός προσώπου.

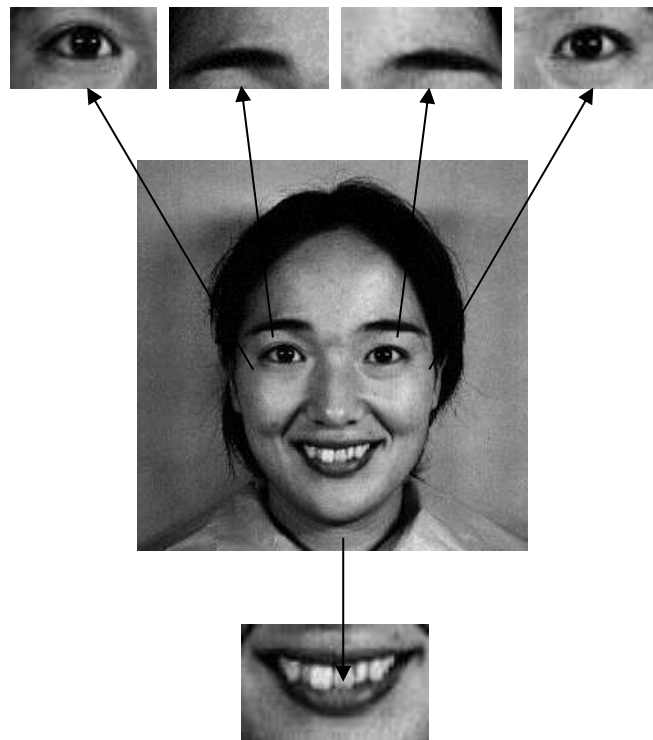


Εικόνα 6.3 Τα στάδια της ανάλυσης των εικόνων προσώπων

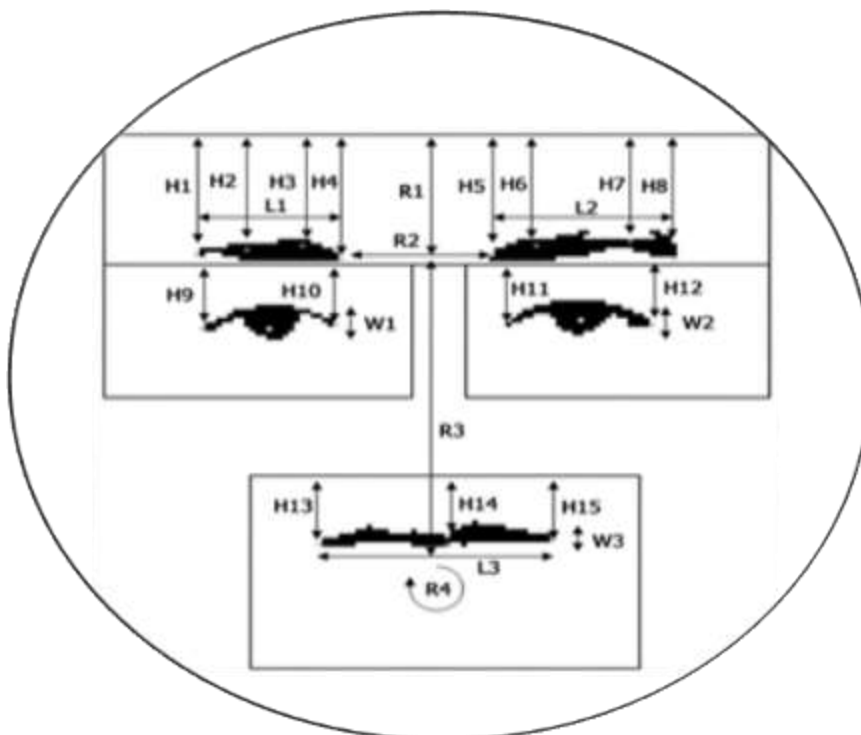
6.3.2 Εξαγωγή Χαρακτηριστικών

Μετά τον προσδιορισμό των περιοχών ενδιαφέροντος (AOIs), κάθε μια περιοχή αναλύεται και εξάγονται χαρακτηριστικά που περιγράφουν και μοντελοποιούν την κατάσταση και τον σχηματισμό της περιοχής. Τα χαρακτηριστικά έχουν ως στόχο να περιγράψουν γεωμετρικά χαρακτηριστικά της όπως είναι, το ύψος, το πλάτος και το μήκος, καθώς και το σχήμα της. Αρχικά, γίνεται εφαρμογή φίλτρων φωτεινότητας για την απλοποίηση μιας AOI, απομακρύνοντας περιττά στοιχεία που δεν φέρουν πληροφορία. Οι περιοχές ενδιαφέροντος, μετά την εφαρμογή των φίλτρων, αναπαρίστανται πιο ομαλά και τα χαρακτηριστικά τους μπορούν να εξαχθούν ευκολότερα. Επίσης, για να μετρηθεί η μέση φωτεινότητα του προσώπου, μια μάσκα προσώπου εφαρμόζεται για την απομόνωση της περιοχής προσώπου από την εικόνα φόντου και επίσης από τα μαλλιά του μοντέλου, αφού αυτές οι περιοχές δεν επηρεάζουν τη φωτεινότητα του.

Στην Εικόνα 6.4 παρουσιάζονται οι περιοχές ενδιαφέροντος και στην Εικόνα 6.5 παρουσιάζονται τα χαρακτηριστικά που εξάγονται από κάθε μια περιοχή για την αναπαράσταση της έκφρασης ενός προσώπου.



Εικόνα 6.4. Απομόνωση περιοχών ενδιαφέροντος



Εικόνα 6.5 Εξαγωγή Χαρακτηριστικών από κάθε περιοχή ενδιαφέροντος και η μοντελοποίηση της έκφρασης

Τα χαρακτηριστικά αυτά αναλυτικά έχουν ως εξής:

Αριστερό Φρύδι

- H1: Το ύψος του σημείου τέρμα αριστερά.
- H2: Το ύψος του σημείου στο $1/3$ της απόστασης μεταξύ του αριστερού και του δεξιού άκρου.
- H3: Το ύψος του σημείου στα $2/3$ της απόστασης μεταξύ του αριστερού και του δεξιού άκρου.
- H4: Το ύψος του σημείου τέρμα δεξιά.
- L1: Το μήκος του φρυδιού.

Δεξί Φρύδι

- H5: Το ύψος του σημείου τέρμα αριστερά.
- H6: Το ύψος του σημείου στο $1/3$ της απόστασης μεταξύ του αριστερού και του δεξιού άκρου.
- H7: Το ύψος του σημείου στα $2/3$ της απόστασης μεταξύ του αριστερού και του δεξιού άκρου.
- H8: Το ύψος του σημείου τέρμα δεξιά.
- L2: Το μήκος του φρυδιού.

Αριστερό μάτι

- H9: Το ύψος του σημείου τέρμα αριστερά.
- H10: Το ύψος του σημείου τέρμα δεξιά
- W1: Το άνοιγμα του ματιού.

Δεξί μάτι

- H11: Το ύψος του σημείου τέρμα αριστερά.
- H12: Το ύψος του σημείου τέρμα δεξιά.
- W2: Το άνοιγμα του ματιού.

Στόμα

- H13: Το ύψος του σημείου τέρμα αριστερά.
- H14: Το ύψος του σημείου στο μέσο του στόματος.
- H15: Το ύψος του σημείου τέρμα δεξιά.
- W3: Το άνοιγμα του στόματος.
- L3: Το μήκος του στόματος.

Επιπλέον χαρακτηριστικά

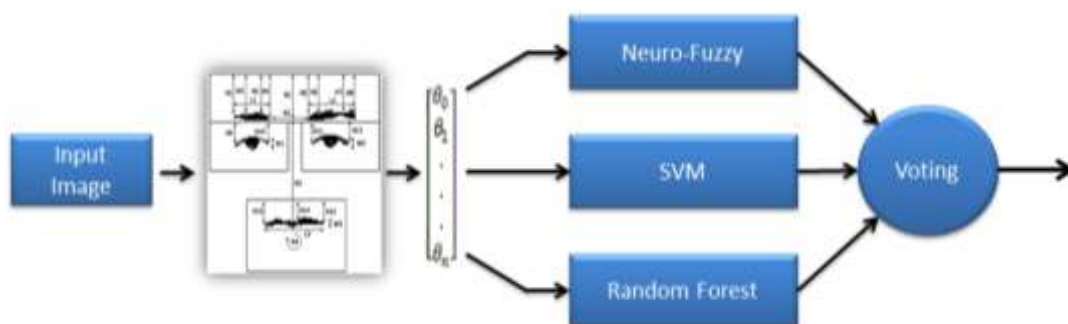
- R1: Το μέσο ύψος των φρυδιών.
- R2: Η απόσταση μεταξύ των φρυδιών.
- R3: Η απόσταση μεταξύ των φρυδιών και του στόματος.

- R4: Ο λόγος του ανοίγματος του στόματος προς το μήκος του.

Επομένως, οι εκφράσεις του προσώπου αναπαρίστανται από τα 25 χαρακτηριστικά, τα οποία μοντελοποιούν την έκφραση και μεταφέρουν την απαραίτητη σημασιολογική πληροφορία ώστε να γίνει η αναγνώριση του συναισθηματικού περιεχομένου της έκφρασης από τους ταξινομητές. Στα πλαίσια της προσέγγισης χρησιμοποιήσαμε πέντε περιοχές ενδιαφέροντος και για κάθε περιοχή εξαγάγαμε τα αντίστοιχα χαρακτηριστικά. Καθορίσαμε το πλαίσιο της μοντελοποίησης να αποτελείται από τα παραπάνω 25 χαρακτηριστικά, καθώς μπορούν να αποτυπώσουν σε μεγάλο βαθμό την έκφραση και ταυτόχρονα είναι εφικτό να προσδιοριστεί κάθε ένα με μεγάλη ακρίβεια.

6.3.3 Ταξινομητές Αναγνώρισης Συναισθηματικών Καταστάσεων

Μετά την ανάλυση μιας έκφρασης και την εξαγωγή των κατάλληλων χαρακτηριστικών της, πραγματοποιείται η αναγνώριση του συναισθηματικού περιεχομένου που φέρει από ταξινομητές μηχανικής μάθησης. Στα πλαίσια της μελέτης, μελετήθηκε η απόδοση τριών αντιπροσωπευτικών ταξινομητών, ενός νευροασαφούς (βασισμένου στο μοντέλο ANFIS), ενός στατιστικού ταξινομητή (βασισμένου στη μέθοδο SVM) και ενός ταξινομητή τυχαίου δάσους (Random Forest) που βασίζεται σε δέντρα αποφάσεων. Με βάση τις επιδόσεις των ταξινομητών, σχεδιάστηκε ομαδοποιημένο σχήμα που συνδυάζει τους παραπάνω ταξινομητές σύμφωνα με την αρχή της πλειοψηφίας. Η διάρθρωσή του παρουσιάζεται στην Εικόνα 6.6.



Εικόνα 6.6 Βασική αρχιτεκτονική του ομαδοποιημένου ταξινομητή.

6.3.3.1 Νευρο-ασαφής (Neuro-Fuzzy) Ταξινομητής

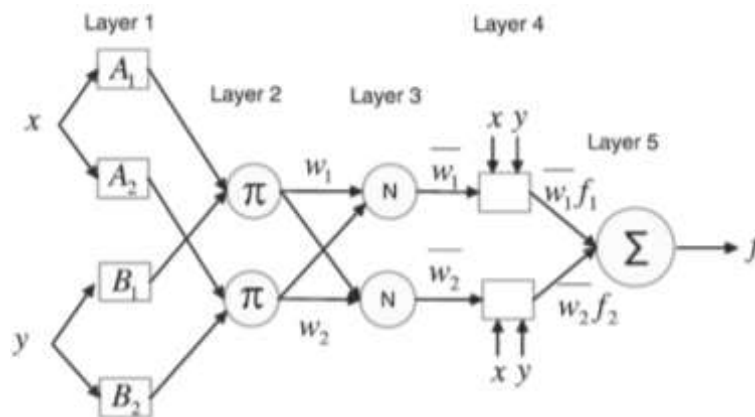
Τα νευρο-ασαφή (Neuro-fuzzy) συστήματα συνδυάζουν τις τεχνολογίες των νευρωνικών δικτύων και της ασαφούς λογικής. Σκοπός ενός τέτοιου συνδυασμού είναι η αξιοποίηση

των πλεονεκτημάτων και των δύο μεθόδων και κυρίως της υπολογιστική δύναμης των νευρωνικών δικτύων και της επικοινωνίας υψηλού επιπέδου με τον χρήστη μέσω των λεκτικών μεταβλητών που χρησιμοποιεί η ασαφής λογική. Τα νευρο-ασαφή συστήματα συνδυάζουν τις τεχνικές των νευρωνικών δικτύων και των συστημάτων ασαφούς συμπερασμού. Το πιο γνωστό ίσως μοντέλο τέτοιου συνδυασμού είναι το ANFIS (Adaptive Neuro-Fuzzy Inference System). Το ANFIS, συνδυάζοντας πτυχές των νευρωνικών δικτύων με πτυχές των ασαφών συστημάτων συμπερασμού (Fuzzy Inference Systems-FIS), συνδυάζει το πλεονέκτημα της εύκολης εφαρμογής με την ικανότητα εκπαίδευσης (Jang, 1993). Ένα FIS προσομοιώνει τη συμπεριφορά κανόνων *if-then* μέσω της γνώσης των ειδικών-εμπειρογνομόνων ή με τη βοήθεια μίας διαθέσιμης βάσης δεδομένων. Για λόγους απλότητας, υποθέτουμε ότι ένα ασαφές σύστημα συμπερασμού έχει δύο εισόδους, x και y , και μία έξοδο f . Η βάση των κανόνων περιέχει δύο ασαφείς κανόνες *if-then*, τύπου ασαφούς μοντέλου Takagi και Sugeno, οι οποίοι εκφράζονται ως εξής:

Rule1: *if x is A_1 and y is B_1 then $f_1 = p_1x + q_1y + r_1$*

Rule2: *if x is A_2 and y is B_2 then $f_2 = p_2x + q_2y + r_2$*

όπου A_i και B_i είναι ασαφή σύνολα (ασαφείς τιμές), f_i είναι οι έξοδοι που ορίζονται εντός της ασαφούς περιοχής από τον ασαφή κανόνα, και p_i , q_i , r_i είναι οι παράμετροι που καθορίζονται κατά τη διάρκεια της εκπαίδευσης.



Εικόνα 6.7 Αρχιτεκτονική ANFIS

Στην Εικόνα 6.7, αποτυπώνεται η αρχιτεκτονική του νευρο-ασαφούς μοντέλου ANFIS, όπου οι κυκλικοί είναι οι σταθεροί κόμβοι, ενώ οι τετράγωνοι κόμβοι με παραμέτρους είναι

οι προσαρμοστικοί κόμβοι. Η αρχιτεκτονική του νευρο-ασαφούς συστήματος συμπερασμού χωρίζεται σε 5 επίπεδα (Jang, 1993).

Το μοντέλο ANFIS που δημιουργήσαμε αποτελείται από πέντε επίπεδα και από 25 μεταβλητές εισόδου $x_1, x_2, x_3, \dots, x_{25}$, που αναπαριστούν τα χαρακτηριστικά που εξάγονται από την έκφραση. Στην περίπτωση μας, το επίπεδο 1, που είναι το επίπεδο εισόδου του νευρο-ασαφούς μοντέλου, έχει 25 νευρώνες που αντιπροσωπεύουν τις είκοσι-πέντε μεταβλητές εισόδου, οι οποίες είναι τα χαρακτηριστικά που εξάγονται από τις περιοχές ενδιαφέροντος. Το επίπεδο εισόδου στέλνει τις εξωτερικές ξερές τιμές (crisp values) απ' ευθείας στο επόμενο στρώμα (επίπεδο 2), το οποίο είναι το επίπεδο ασαφοποίησης, και οι ξερές τιμές από το πρώτο επίπεδο μετασχηματίζονται στις κατάλληλες λεκτικές (ασαφείς) τιμές (low, average ή high). Κάθε κόμβος του επιπέδου 2 αναπαριστά το ασαφές σύνολο μίας από τις λεκτικές μεταβλητές εισόδου. Κάθε κόμβος στο επίπεδο είναι ένας σταθερός (fixed) κόμβος και έχει ως ρόλο του την εκτέλεση του πολλαπλασιασμού των σημάτων εξόδου του πρώτου επιπέδου. Η κάθε έξοδος αυτού του επιπέδου αντιστοιχεί στην ισχύ ενεργοποίησης (firing strength) ενός κανόνα. Οι νευρώνες του επιπέδου 3 (επίπεδο ασαφών κανόνων) αναπαριστούν τους ασαφείς κανόνες (fuzzy rules) ($R_1, \dots, R_{154}$) και κάθε κόμβος υπολογίζει το λόγο του βαθμού ενεργοποίησης (firing strength). Οι έξοδοι αυτού του επιπέδου ονομάζονται κανονικοποιημένοι βαθμοί ενεργοποίησης (normalized firing strengths). Στο επίπεδο 4 (έξοδος επιπέδου συναρτήσεων συμμετοχής) κάθε κόμβος είναι ένας προσαρμοστικός κόμβος και δημιουργεί έναν απλό γραμμικό συνδυασμό των εισόδων του συστήματος και ενός συνόλου παραμέτρων με τις εξόδους του επιπέδου 3, και μετά υπολογίζει τη συνεισφορά του κάθε κανόνα προς τη συνολική έξοδο. Οι παράμετροι σε αυτό το επίπεδο ορίζονται ως συνεπαγόμενοι παράμετροι (consequent parameters). Το επίπεδο 5 (επίπεδο ασαφοποίησης) αποτελείται από ένα μόνο σταθερό κόμβο και υπολογίζει τη συνολική έξοδο σαν το συνολικό άθροισμα όλων των εισερχόμενων σημάτων – συναρτήσεων.

Στο μοντέλο ANFIS, οι έξοδοι σε κάθε επίπεδο αποτελούν μέρος των εισόδων του επόμενου επιπέδου. Η εκπαίδευση του ANFIS βασίζεται σε έναν υβριδικό αλγόριθμο εκμάθησης (Jang, 1993). Ο αλγόριθμος περιλαμβάνει ένα προς τα εμπρός έλεγχο (forward pass) και ένα προς τα πίσω πέρασμα (backward pass). Στον προς τα εμπρός έλεγχο, ο αλγόριθμος χρησιμοποιεί την μέθοδο των ελαχίστων τετραγώνων για να αναγνωρίσει τις παραμέτρους συμπερασμού του επιπέδου 4. Ο υβριδικός αλγόριθμος έχει αποδειχθεί ότι είναι πολύ αποτελεσματικός σε μοντέλα ANFIS. Στο πλαίσιο της μελέτης μας, το μοντέλο

ANFIS εκπαιδεύτηκε με βάση το σύνολο δεδομένων των εκφράσεων των βάσεων δεδομένων. Μετά την εκπαίδευση του, το νευρο-ασαφές μοντέλο είναι πλήρες και μπορεί να χρησιμοποιηθεί για την ταξινόμηση νέων ασκήσεων.

6.3.3.2 Ταξινομητής Μηχανής Διανυσμάτων Υποστήριξης (Support Vector Machine)

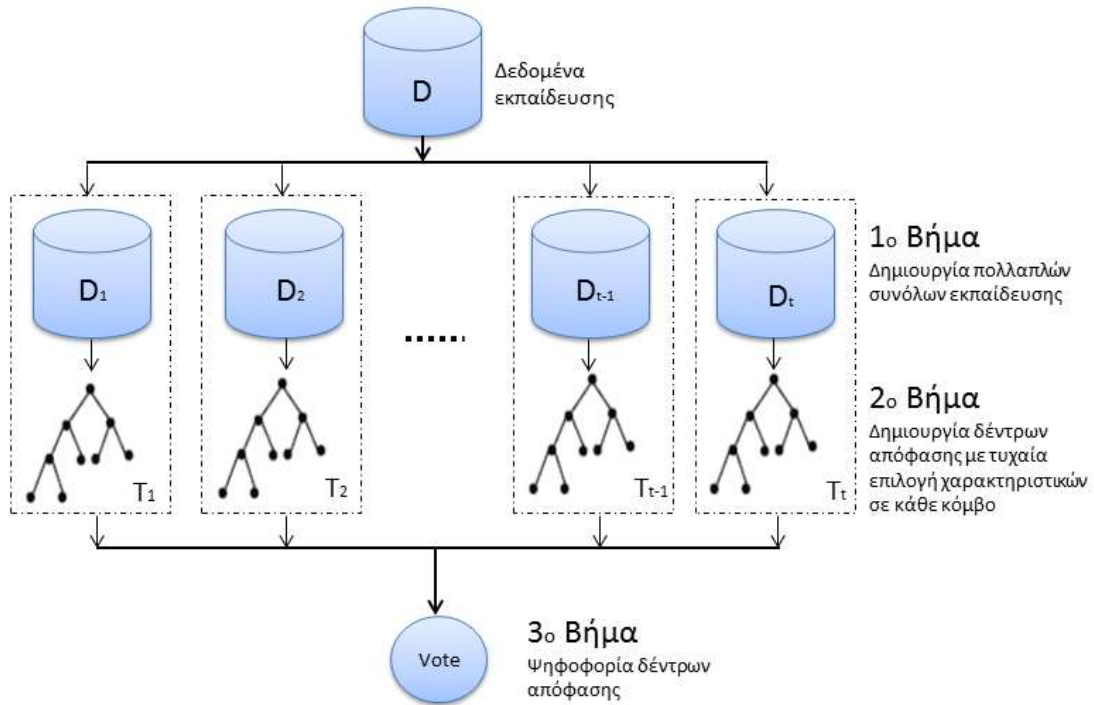
Οι μηχανές διανυσμάτων υποστήριξης (Support Vector Machines- SVMs) (Cortes & Vapnik, 1990) στηρίζονται στη θεωρία της στατιστικής μάθησης και στα νευρωνικά δίκτυα και αποτελούν μια από τις πιο διαδεδομένες μεθόδους (γραμμικής και μη) ταξινόμησης. Οι μηχανές διανυσμάτων υποστήριξης για την δημιουργία ενός μοντέλου ταξινόμησης προσπαθούν να βρουν μια υπερεπιφάνεια (hypersurface), που να διαχωρίζει στο χώρο τα στιγμιότυπα των δύο κλάσεων. Η υπερεπιφάνεια αυτή επιλέγεται ώστε να απέχει όσο το δυνατόν περισσότερο από τα κοντινότερα στιγμιότυπα των δύο κλάσεων (maximum margin hypersurface). Από γεωμετρική σκοπιά, η προσέγγιση των μηχανών διανυσμάτων υποστήριξης προσπαθεί να βρει στον $|T|$ -διάστατο χώρο, ανάμεσα από τα διαχωριστικά υπερεπίπεδα $h_1, h_2, h_3 \dots h_n$, εκείνο το υπερεπίπεδο που διαχωρίζει τα στιγμιότυπα των κλάσεων και έχει το μεγαλύτερο εύρος ανάμεσά τους.

6.3.3.3 Ταξινομητής Τυχαίου Δάσους (Random Forest)

Ο αλγόριθμος τυχαίου δάσους (Random Forest) (Breiman, 2001) βασίζεται στα δέντρα απόφασης (Decision Trees) και σχηματίζει ένα δάσος, όπου κάθε δέντρο εκτελεί την διεργασία της ταξινόμησης ανεξάρτητα. Τα δέντρα απόφασης είναι ιεραρχικές δομές όπου κάθε εσωτερικός κόμβος περιλαμβάνει έναν έλεγχο πάνω σε ένα γνώρισμα, κάθε κλαδί αντιστοιχεί στο αποτέλεσμα αυτού του ελέγχου και κάθε φύλλο δίνει μια πρόγνωση για την τιμή της μεταβλητής της κατηγορίας.

Βασικό χαρακτηριστικό που εισάγεται στα τυχαία δάση είναι το επίπεδο τυχαιότητας, που αφορά τον τρόπο δόμησης των δέντρων απόφασης. Πιο συγκεκριμένα, εισάγεται τυχαιότητα στην κατασκευή του δέντρου, η οποία περιλαμβάνει την τυχαία επιλογή διαμέλισης, καθώς και την τυχαία επιλογή δεδομένων εισόδου. Η τυχαία επιλογή διαμέλισης καθορίζει ότι κάθε κόμβος διαμελίζεται επιλέγοντας τυχαία μία από τις καλύτερες διαμελίσεις στο κόμβο αυτό, ενώ η την τυχαία επιλογή εισόδου καθορίζει ότι η διαμέλιση του κόμβου αποφασίζεται από ένα τυχαία επιλεγμένο υποσύνολο των δεδομένων εισόδου. Κατά την φάση της εκπαίδευσης, κάθε δέντρο εκπαιδεύεται σε ένα υποσύνολο του συνόλου δεδομένων και η επιλογή του επόμενου χωρίσματος στο κλαδί ενός δέντρου υπολογίζεται τυχαία. Έτσι, με αυτό τον τρόπο έχουμε τυχαιότητα όχι μόνο στο επίπεδο των

δεδομένων εκπαίδευσης, αλλά και στο επίπεδο του μοντέλου. Η ανάλυση του τρόπου λειτουργίας των παρουσιάζεται στην εικόνα 6.9.



Εικόνα 6.8 Λειτουργία τυχαιών δασών.

Για να ταξινομήσουμε ένα νέο άγνωστο διάνυσμα εισόδου με το συγκεκριμένο ταξινομητή, τροφοδοτούμε το διάνυσμα σε κάθε δένδρο του δάσους και το συγκεκριμένο δένδρο κάνει τη δική του πρόβλεψη. Γενικά, ο αριθμός των δέντρων προσδιορίζεται με βάση τα χαρακτηριστικά και τις παραμέτρους του προβλήματος. Έρευνες έχουν δείξει ότι ο ενδεδειγμένος αριθμός κυμαίνεται συνήθως από 64-128 (Oshiro et al., 2012). Στην περίπτωση μας, ο ταξινομητής τυχαίου δάσους που εκπαιδεύτηκε, αποτελείται από 80 επιμέρους δέντρα, αριθμός που καθορίστηκε εμπειρικά με βάση τις δοκιμές και τα πειραματικά αποτελέσματα. Στο τέλος, οι ψήφοι όλων των δένδρων συνεκτιμώνται σε ένα σχήμα ψηφοφορίας και η τελική πρόβλεψη του διανύσματος από τον ταξινομητή είναι η κατηγορία εκείνη που συγκέντρωσε τις περισσότερες ψήφους από τα δένδρα ταξινόμησης. Τα τυχαία δάση αποτελούν έναν αποτελεσματικό τρόπο να περιορίζεται η επιρροή που έχουν μικρές τυχαίες παραλλαγές στο σύνολο δεδομένων σε ένα δέντρο απόφασης που εκπαιδεύτηκε στο σύνολο δεδομένων.

6.4 Πειραματική Αξιολόγηση

Για την αξιολόγηση της απόδοσης των ταξινομητών σχεδιάστηκε και πραγματοποιήθηκε μια εκτεταμένη πειραματική μελέτη. Η πειραματική μελέτη, είχε δύο κύρια μέρη, όπου κάθε ένα αξιολόγησε διαφορετικές συνιστώσες των συστημάτων. Πιο συγκεκριμένα, στα πλαίσια του πρώτου μέρους της πειραματικής μελέτης, έγινε μελέτη και αξιολόγηση της απόδοσης των ταξινομητών στην αναγνώριση της ύπαρξης συναισθηματικού περιεχομένου σε εκφράσεις προσώπου. Στην συνέχεια, στο δεύτερο μέρος, έγινε αξιολόγηση της απόδοσης στην αναγνώριση του είδους του συναισθηματικού περιεχομένου των εκφράσεων, σύμφωνα με το μοντέλο συναισθημάτων του Ekman. Επίσης, στα πλαίσια της μελέτης, χρησιμοποιήθηκαν δεδομένα από διάφορες βάσεις εκφράσεων καθώς επίσης και εικόνες φυσικών προσώπων εκτός των βάσεων δεδομένων και δημιουργήθηκε ένα σύνολο δεδομένων με στόχο να γίνει πολύπλευρη αξιολόγηση της απόδοσης των συστημάτων σε διαφορετικούς τύπους εκφράσεων.

6.4.1 Δημιουργία Συνόλου Δεδομένων

Στα πλαίσια της πειραματικής μελέτης, έγινε αξιολόγηση της απόδοσης των ταξινομητών σε δεδομένα των βάσεων δεδομένων Jaffe και Cohn-Kanade οι οποίες περιέχουν εκφράσεις προσώπου. Οι δύο βάσεις έχουν χρησιμοποιηθεί ευρέως στην αποτίμηση της απόδοσης συστημάτων και προσεγγίσεων για την ανάλυση. Η βάση προσώπων Jaffe (Japanese Female Facial Expression) (Lyons et al., 1998) περιέχει συνολικά 213 εικόνες προσώπων από 10 ανθρώπους, οι οποίοι εκφράζουν τις 6 βασικές συναισθηματικές καταστάσεις του μοντέλου Ekman καθώς και την ουδέτερη κατάσταση. Η βάση Cohn-Kanade (Kanade et al., 2000) περιέχει 486 ακολουθίες εικόνων από 97 άτομα. Οι εκφράσεις προσώπου ταξινομούνται στις εξής συναισθηματικές καταστάσεις: θυμός (anger), περιφρόνηση (contempt), απέχθεια (disgust), φόβος (fear), χαρά (joy), ουδέτερη (neutral), λύπη (sadness) και έκπληξη (surprise). Η βάση αποτελείται από ενήλικες, 69% γυναίκες και 31% άντρες, ηλικίας 18 ως 50 χρονών, διαφόρων εθνικοτήτων. Στα πλαίσια της μελέτης, χρησιμοποιήθηκαν όλες οι συναισθηματικές εκφράσεις του μοντέλου Ekman, πλην της περιφρόνησης.

6.4.2 Αξιολόγηση προσδιορισμού ύπαρξης Συναισθηματικού Περιεχομένου

Στο πρώτο μέρος της πειραματικής μελέτης έγινε αξιολόγηση των επιδόσεων των ταξινομητών ως προς τον καθορισμό της ύπαρξης συναισθηματικού περιεχομένου σε

εκφράσεις προσώπου του συνόλου δεδομένων. Για την εκπαίδευση των ταξινομητών έγινε χωρισμός του συνόλου δεδομένων σε τρία ίσα μέρη, όπου τα δύο χρησιμοποιήθηκαν για την εκπαίδευση και το ένα για την αποτίμηση της απόδοσης των ταξινομητών. Για την αξιολόγηση της απόδοσης των ταξινομητών, χρησιμοποιήθηκαν οι μετρικές: ορθότητα (accuracy), ακρίβεια (precision), ευαισθησία (sensitivity) και εξειδίκευση (specificity), δεδομένου ότι η κλάση εξόδου είναι μια δυαδική μεταβλητή (binary class output variable) (Sokolova & Lapalme, 2009). Τα αποτελέσματα που προέκυψαν παρουσιάζονται στον Πίνακα 6.1.

Πίνακας 6.1. Συνολικά αποτελέσματα απόδοσης ταξινομητών

	Neuro Fuzzy	SVM	Random Forest	Majority
Accuracy	0.91	0.93	0.92	0.96
Precision	0.93	0.95	0.94	0.96
Sensitivity	0.93	0.94	0.96	0.96
Specificity	0.82	0.88	0.77	0.91

Η καλύτερη απόδοση παρουσιάζεται από τον ταξινομητή SVM και είναι ελαφρώς καλύτερη από την απόδοση του ταξινομητή τυχαίου δάσους. Η ομαδοποίηση των τριών ταξινομητών σε σχήμα που ακολουθεί την αρχή της πλειοψηφίας (majority voting) παρουσιάζει καλύτερη απόδοση από όλους τους επιμέρους βασικούς ταξινομητές. Στην συνέχεια, στον Πίνακα 6.2 παρουσιάζονται τα αναλυτικά αποτελέσματα των επιδόσεων των ταξινομητών σε κάθε βάση δεδομένων.

Πίνακας 6.2. Ανάλυση της απόδοσης των ταξινομητών σε κάθε βάση δεδομένων

	Jaffe				Cohn-kanade			
	Neuro Fuzzy	SVM	Random Forest	Majority	Neuro Fuzzy	SVM	Random Forest	Majority
Accuracy	0.91	0.92	0.91	0.95	0.90	0.94	0.92	0.96
Precision	0.94	0.95	0.91	0.95	0.92	0.96	0.93	0.98
Sensitivity	0.92	0.93	0.97	0.95	0.95	0.95	0.97	0.96

Specificity	0.86	0.87	0.73	0.90	0.77	0.90	0.82	0.95
-------------	------	------	------	------	------	------	------	------

Από τα αποτελέσματα φαίνεται ότι η απόδοση των ταξινομητών είναι πολύ ικανοποιητική. Ο ταξινομητής SVM παρουσιάζει την καλύτερη απόδοση τόσο στα δεδομένα της βάσης δεδομένων Jaffe όσο και σ' αυτά της βάσης Cohn-Kanade. Η απόδοση του ταξινομητή random forest είναι στο ίδιο επίπεδο με την απόδοση του SVM στην βάση Jaffe, ενώ είναι ελαφρώς χαμηλότερη στην βάση Cohn-Kanade. Ο νεύρο-ασαφής ταξινομητής παρουσιάζει σχεδόν την ίδια απόδοση και στις δύο βάσεις δεδομένων.

Από την πλευρά των βάσεων δεδομένων, ο νεύρο-ασαφής ταξινομητής παρουσιάζει καλύτερη απόδοση στην αναγνώριση της ύπαρξης συναισθήματος σε δεδομένα της βάσης Jaffe από ότι σε δεδομένα της βάσης Cohn-Kanade, ενώ οι υπόλοιποι ταξινομητές παρουσιάζουν καλύτερη απόδοση στην βάση Cohn-Kanade. Επίσης, αξίζει να σημειωθεί ότι οι ταξινομητές παρουσιάζουν πολύ υψηλή ευαισθησία (sensitivity), η οποία είναι υψηλότερη από την εξειδίκευση (specificity), κάτι που δείχνει ότι παρουσίασαν καλή απόδοση στην αναγνώριση εκφράσεων που ήταν συναισθηματικά ουδέτερες, οι οποίες αναγνωρίστηκαν αποτελεσματικά από τον μηχανισμό και ταξινομήθηκαν σωστά στην ουδέτερη κατηγορία. Επιπλέον, η χαμηλή εξειδίκευση (specificity) είναι αποτέλεσμα της ταξινόμησης συναισθηματικών εκφράσεων στην ουδέτερη κατηγορία. Αυτό οφείλεται στο γεγονός ότι ορισμένες εικόνες των βάσεων εκφράζουν το συναίσθημα χωρίς ιδιαίτερη ένταση και ισχυρές παραμορφώσεις στα χαρακτηριστικά του προσώπου.

Η ομαδοποίηση των ταξινομητών παρουσιάζει καθολικά καλύτερα αποτελέσματα σε σχέση με τους βασικούς ταξινομητές. Για την αποτίμηση απόδοσής της και της συνεισφοράς της στην βελτίωσης της απόδοσης των βασικών ταξινομητών, έγινε σύγκριση της με τον καλύτερο βασικό ταξινομητή. Στόχος είναι να καθοριστεί το ποσοστό της βελτίωσης της απόδοσης της ταξινόμησης από το ομαδοποιημένο σχήμα ταξινόμησης. Από τα αποτελέσματα προκύπτει ότι ο ομαδοποιημένος ταξινομητής που ακολουθεί την αρχή της πλειοψηφίας παρουσιάζει ορθότητα η οποία είναι υψηλότερη κατά 3.2% σε σχέση με την χρήση μόνο του ταξινομητή SVM, ο οποίος παρουσιάζει την καλύτερη απόδοση.

Επιπροσθέτως, χρησιμοποιήθηκε η μετρική Cohen's Kappa για να υπολογίσει τον βαθμό αξιοπιστίας της συμφωνίας μεταξύ των εκτιμήσεων του εμπειρογνώμονα και των ταξινομητών στο σύνολο των δεδομένων του πειράματος. Τα αποτελέσματα παρουσιάζονται στον Πίνακα 6.3.

Πίνακας 6.3. Αποτελέσματα μετρικής Cohen's kappa

	Neuro Fuzzy	SVM	Random Forest	Majority
Annotated Data	0.77	0.83	0.81	0.87

Τα αποτελέσματα της μελέτης έδειξαν ότι υπήρξε μια σχεδόν τέλεια συμφωνία μεταξύ του ταξινομητή SVM και των δεδομένων των δύο βάσεων, δεδομένου ότι η μετρική Kappa ήταν $k = 0.83$ (confidence interval 95%, $p < 0.0005$). Επίσης, το ίδιο προέκυψε και για τον ταξινομητή Random Forest, όπου υπήρξε μια σχεδόν τέλεια συμφωνία και η μετρική Kappa υπολογίστηκε ότι ήταν $k = 0.81$. Όσον αφορά τον νευρο-ασαφή ταξινομητή, τα αποτελέσματα έδειξαν ότι υπήρξε ισχυρή συμφωνία και η μετρική Kappa υπολογίστηκε ότι ήταν $k = 0.77$. Επίσης, όσον αφορά το ομαδοποιημένο σχήμα ταξινόμησης, από τα αποτελέσματα φαίνεται ότι υπήρξε μια ισχυρή συμφωνία όπου η μετρική Kappa υπολογίστηκε ότι ήταν $k = 0.87$.

Ένα βασικό πλεονέκτημα της προσέγγισης μας και των ομαδοποιημένων ταξινομητών σε σχέση με τις καλύτερες μεθόδους της διεθνούς βιβλιογραφίας αποτελεί η κλιμάκωση σε ετερογενή δεδομένα όπως από διαφορετικές βάσεις εικόνων, καθώς και η ενσωμάτωση σε εφαρμογές με δεδομένα του πραγματικού κόσμου.

6.4.3 Αξιολόγηση Αναγνώρισης Συναισθηματικών Καταστάσεων

Στο δεύτερο μέρος της πειραματικής μελέτης έγινε αξιολόγηση των επιδόσεων των ταξινομητών στην αναγνώριση του συναισθηματικού περιεχομένου των εκφράσεων στην κλίμακα του μοντέλου Ekman. Το πρόβλημα αντιμετωπίστηκε ξανά ως πρόβλημα ταξινόμησης των εκφράσεων στην κατάλληλη κατηγορία. Στα πλαίσια της πειραματικής μελέτης, επιλέχθηκαν οι εκφράσεις του συνόλου δεδομένων που χαρακτηρίστηκαν ως συναισθηματικές και για τις οποίες είχε προσδιοριστεί σωστά η ύπαρξη συναισθηματικού περιεχομένου τους στα προηγούμενα μέρη της διαδικασίας αξιολόγησης. Έγινε αξιολόγηση των επιδόσεων των βασικών ταξινομητών καθώς και του ομαδοποιημένου σχήματος σύνθεσης τους με την αρχή της πλειοψηφίας. Για την αξιολόγηση της απόδοσης υπολογίστηκαν οι μετρικές της ακρίβειας (precision), της ανάκλησης (recall) καθώς και η μετρική F-measure δεδομένου ότι η κλάση εξόδου είναι μια μεταβλητή πολλών τιμών. Τα

γενικά αποτελέσματα της απόδοσης των ταξινομητών που προέκυψαν από την πειραματική μελέτη παρουσιάζονται στον Πίνακα 6.4.

Πίνακας 6.4. Συνολικά αποτελέσματα απόδοσης ταξινομητών

	Neuro Fuzzy	SVM	Random Forest
Precision	0.82	0.87	0.85
Recall	0.80	0.88	0.84
F-measure	0.81	0.87	0.85

Πίνακας 6.5. Ανάλυση της απόδοσης των ταξινομητών σε κάθε βάση δεδομένων

	Jaffe			Cohn-Kanade		
	Neuro Fuzzy	SVM	Random Forest	Neuro Fuzzy	SVM	Random Forest
Precision	0.80	0.86	0.84	0.85	0.89	0.87
Recall	0.78	0.86	0.83	0.85	0.89	0.86
F-measure	0.79	0.86	0.83	0.85	0.89	0.87

Τα αποτελέσματα δείχνουν ότι οι ταξινομητές έχουν πολύ καλή απόδοση στην αναγνώριση της συναισθηματικής κατάστασης των εκφράσεων. Η καλύτερη απόδοση επιτυγχάνεται από τον ταξινομητή SVM, όπου η F-measure είναι καλύτερη από εκείνες του νευρο-ασαφούς και του ταξινομητή τυχαίου δάσους. Η απόδοσή των ταξινομητών είναι καλύτερη στα δεδομένα της βάσης Cohn-Kanade σε σχέση με τα δεδομένα της βάσης Jaffe. Στην συνέχεια έγινε προσδιορισμός των επιδόσεων των ταξινομητών στην αναγνώριση κάθε ενός βασικού συναισθήματος σύμφωνα με το μοντέλο Ekman και τα αποτελέσματα παρουσιάζονται στον Πίνακα 6.6.

Πίνακας 6.6 Ανάλυση της απόδοσης των ταξινομητών σε κάθε μια συναισθηματική κατηγορία.

Neuro-fuzzy	SVM	Random Forest
-------------	-----	---------------

	Precision	Recall	F-measure	Precision	Recall	F-measure	Precision	Recall	F-measure
Anger	0.89	0.76	0.82	0.86	0.86	0.86	0.94	0.76	0.84
Disgust	0.82	0.74	0.78	0.90	0.95	0.92	0.78	0.95	0.86
Fear	0.73	0.76	0.74	0.85	0.76	0.80	0.80	0.76	0.78
Joy	0.82	1.00	0.90	0.90	0.96	0.93	0.81	0.96	0.88
Sadness	0.80	0.80	0.80	0.89	0.85	0.87	0.94	0.75	0.83
Surprise	0.89	0.84	0.86	0.85	0.89	0.87	0.85	0.89	0.87

Τα αποτελέσματα της ανάλυσης της απόδοσης των ταξινομητών στην αναγνώριση κάθε μιας συναισθηματικής κατάστασης δείχνουν ότι και οι τρεις ταξινομητές έχουν σταθερή απόδοση και παρουσιάζουν πολύ καλή συμπεριφορά στην αναγνώριση του συναισθήματος της χαράς. Η συμπεριφορά αυτή οφείλεται στο γεγονός ότι το συναίσθημα της χαράς αποτελεί ένα δυνατό συναίσθημα και εκφράζεται στις περισσότερες περιπτώσεις με έντονο τρόπο και σημαντικές τροποποιήσεις των χαρακτηριστικών του προσώπου. Από την άλλη πλευρά, η απόδοσή τους μειώνεται στην αναγνώριση του συναισθήματος του φόβου. Ο ταξινομητής SVM παρουσιάζει καλή απόδοση στην αναγνώριση όλων των συναισθημάτων και η ακρίβειά του είναι σταθερή και περίπου στο ίδιο επίπεδο (0.85-0.90) για όλες τις συναισθηματικές κατηγορίες. Σε ορισμένες συναισθηματικές κατηγορίες όπως του φόβου εμφανίζει χαμηλή ανάκληση, αποτέλεσμα της ταξινόμησης εκφράσεων φόβου λανθασμένα σε κάποια άλλη κατηγορία. Η απόδοσή του είναι υψηλότερη για τα συναισθήματα της χαράς και της απέχθειας, ενώ μειώνεται στην αναγνώριση του συναισθήματος του φόβου. Ο ταξινομητής τυχαίου δάσους παρουσιάζει αντιστοίχως καλή συμπεριφορά στην αναγνώριση της χαράς ($F=0.88$), της έκπληξης ($F=0.88$) και της απέχθειας ($F=0.86$) και χαμηλότερη στην αναγνώριση του θυμού. Αξιοσημείωτο σημείο είναι η ανάκληση που παρουσιάζει ο νεύρο-ασαφής ταξινομητής στο συναίσθημα της χαράς, κάτι που δείχνει ότι όλες οι εκφράσεις που φέρουν το συναίσθημα της χαράς αναγνωρίστηκαν σωστά από τον νεύρο-ασαφή ταξινομητή.

Επίσης, χρησιμοποιήθηκε η μετρική Cohen's Kappa για να υπολογίσει τον βαθμό αξιοπιστίας της συμφωνίας μεταξύ των δεδομένων των βάσεων και των ταξινομητών στο σύνολο των δεδομένων του πειράματος. Τα αποτελέσματα παρουσιάζονται στον Πίνακα 6.7.

Πίνακας 6.7. Αποτελέσματα Cohen's kappa μετρικής

	Neuro Fuzzy	SVM	Random Forest
Annotated Data	0.80	0.85	0.83

Τα αποτελέσματα της μελέτης έδειξαν ότι υπήρξε μια σχεδόν τέλεια συμφωνία μεταξύ των ταξινομητών και των δεδομένων των βάσεων. Η μετρική Kappa υπολογίστηκε ότι ήταν για τον ταξινομητή SVM $k = 0.85$ (confidence interval 95%, $p < 0.0005$), για τον Random Forest $k = 0.83$ και για τον νευρο-ασαφή ταξινομητή 0.80. Γενικά, τα αποτελέσματα δείχνουν μια πολύ καλή απόδοση των ταξινομητών στην αναγνώριση των συναισθηματικών καταστάσεων, κάτι που δείχνει και ότι τα χαρακτηριστικά που εξάγονται διατηρούν και αναπαριστούν σημασιολογικά την συναισθηματική πληροφορία των εκφράσεων.

Τα αποτελέσματα δείχνουν μια πολύ καλή απόδοση η οποία είναι κοντά στις αποδόσεις των καλύτερων προσεγγίσεων στην διεθνή βιβλιογραφία και οι οποίες αναφέρουν ακρίβεια στο εύρος 0.85-0.99 στις βάσεις δεδομένων που χρησιμοποιήθηκαν για την πειραματική μας αξιολόγησης. Ένα βασικό πλεονέκτημα της προσέγγισής μας αφορά την χρήση ομαδοποιημένων σχημάτων ταξινόμησης, κάτι που προσφέρει μεγαλύτερη κλιμάκωση σε νέα ετερογενή δεδομένα, τα οποία έχουν προέλθει από διαφορετικές πηγές, όπως βάσεις δεδομένων προσώπου ή ακόμη και από τον πραγματικό κόσμο. Ακόμη, αξίζει να σημειωθεί ότι, οι περισσότερες προσεγγίσεις στην διεθνή βιβλιογραφία αφορούν την εκπαίδευση και την χρήση μιας μεθόδου αναγνώρισης του συναισθηματικού περιεχομένου εκφράσεων και πραγματοποιούν στις περισσότερες περιπτώσεις αποτίμηση της απόδοσής της σε μια βάση δεδομένων εκφράσεων προσώπων. Επίσης, οι περισσότερες προσεγγίσεις βασίζονται στην χρήση ενός ταξινομητή, ο οποίος εκπαιδεύεται ταυτόχρονα σε συναισθηματικές και ουδέτερες εκφράσεις με τελικό στόχο την αναγνώριση του συναισθηματικού περιεχομένου των εκφράσεων. Στην προσέγγισή μας, μια διαφοροποίηση είναι η χρήση μηχανισμών για τον προσδιορισμό αρχικά της ύπαρξης ή όχι συναισθηματικού περιεχομένου στις εκφράσεις και στην συνέχεια η αναγνώριση των συγκεκριμένων συναισθημάτων που φέρουν, γεγονός που σε συνδυασμό με τα

χαρακτηριστικά των μηχανισμών ταξινόμησης που χρησιμοποιούνται, συμβάλει στην καλή απόδοση και κλιμάκωση σε νέα δεδομένα εκφράσεων προσώπου.

7

Συμπεράσματα και Μελλοντική Έρευνα

7.1 Συμπεράσματα

Στο Κεφάλαιο αυτό παρουσιάζεται μια επισκόπηση της συνεισφοράς της παρούσας διατριβής και παρατίθενται τα κύρια συμπεράσματα που προέκυψαν. Επίσης, παρουσιάζονται κύριες κατευθύνσεις στις οποίες θα μπορούσε να εστιάσει η μελλοντική έρευνα.

Μια πρώτη συνιστώσα της έρευνας ήταν η ανάλυση της αλληλεπίδρασης του χρήστη και η παροχή ανάδρασης σε αυτόν, στα πλαίσια ευφών εκπαιδευτικών συστημάτων. Αξίζει να σημειωθεί ότι στη διεθνή βιβλιογραφία δεν υπάρχουν αναφορές σε γενικές μεθοδολογίες που να μοντελοποιούν τη διαδικασία της δημιουργίας ανάδρασης σε ευφυή συστήματα. Στα πλαίσια της διατριβής, σχεδιάστηκε μια μεθοδολογία που περιλαμβάνει μια γενική, συστηματική και αναλυτική μοντελοποίηση των τύπων των ενεργειών και των απαντήσεων των εκπαιδευόμενων στα πλαίσια της ενασχόλησης τους με διάφορων ειδών εκπαιδευτικές δραστηριότητες. Επίσης, σχεδιάστηκε ένα ιεραρχικό μοντέλο οργάνωσης της ανάδρασης σε επίπεδα. Στα πλαίσια της παροχής ανάδρασης, χρησιμοποιούνται κανόνες για να αναλύσουν την συμπεριφορά του εκπαιδευόμενου και να καθορίσουν την παροχή της κατάλληλης ανάδρασης από το σύστημα. Τα παραπάνω αναπτύχθηκαν σ' ένα ευφές σύστημα διδασκαλίας σε θέματα λογικής ως γλώσσας αναπαράστασης γνώσης. Τα αποτελέσματα της πειραματικής μελέτης έδειξαν ότι το πλαίσιο της ανάδρασης είναι αποδοτικό, συμβάλλοντας ώστε να γίνει η μάθηση των εκπαιδευόμενων πιο αποτελεσματική. Το

πλαίσιο ανάδρασης είναι γενικό και μπορεί να εφαρμοστεί και σε άλλα ευφυή συστήματα διδασκαλίας.

Μια δεύτερη ερευνητική συνιστώσα της διατριβής ήταν να σχεδιάσουμε και να αναπτύξουμε μηχανισμούς και μεθοδολογίες για την αυτόματη ανάλυση κειμένου και τον προσδιορισμό του συναισθηματικού περιεχομένου που φέρει. Αρχικά, για τον σκοπό αυτό έγινε ανάπτυξη ενός μηχανισμού βασισμένου στην γνώση, που πραγματοποιεί εις βάθος ανάλυση της φυσικής γλώσσας και προσδιορίζει το συναισθηματικό περιεχόμενο χρησιμοποιώντας λεξικολογικούς πόρους. Ένα βασικό χαρακτηριστικό του εργαλείου είναι η γενική χρήση του, καθώς δεν απαιτεί εκπαίδευση σε δεδομένα. Επίσης, σχεδιάστηκαν σχήματα ομαδοποίησης ταξινομητών, που συνδυάζουν το εργαλείο με ταξινομητές από το πεδίο της μηχανικής μάθησης, όπως Naïve Bayes και Support Vector Machines. Η ανάπτυξη σχημάτων ομαδοποίησης ταξινομητών είχε ως στόχο τόσο την αύξηση της απόδοσης της διεργασίας της αναγνώρισης συναισθηματικού περιεχομένου όσο και την αποτελεσματική γενίκευση σε μεγάλο πλήθος διαφορετικών ειδών κειμενικών δεδομένων. Τα πειραματικά αποτελέσματα της αξιολόγησης έδειξαν ότι τα ομαδοποιημένα σχήματα παρουσίασαν μια αύξηση της απόδοσης κατά 4.8 έως 8.4% σε σχέση με τον καλύτερο ταξινομητή. Επίσης, όσον αφορά τους επιμέρους ταξινομητές, το εργαλείο που βασίζεται στην γνώση είχε πολύ καλή απόδοση στην ανάλυση ορθά δομημένων προτάσεων που ακολουθούν συντακτικούς κανόνες και η γενική φύση λειτουργίας του σε συνδυασμό με την φύση λειτουργίας των ομαδοποιημένων σχημάτων ταξινόμησης μπορούν να εξασφαλίσουν μια καλή απόδοση και κλιμάκωση σε νέα δεδομένα.

Μια ακόμη συνεισφορά της διατριβής αποτελεί η εισαγωγή ενός τρόπου αναπαράστασης του συναισθηματικού περιεχομένου της συζήτησης σε θεματική περιοχή. Για τον σκοπό αυτό, σχεδιάστηκε μια προσέγγιση που δημιουργεί το συναισθηματικό γράφημα μιας συζήτησης, το οποίο παρέχει ένα δομημένο τρόπο αναπαράστασης των συναισθηματικών καταστάσεων με βάση τις επιμέρους αναλύσεις του συναισθηματικού περιεχομένου κάθε μέρους της συζήτησης. Η μοντελοποίηση του συναισθηματικού περιεχομένου των συζητήσεων και η αναλυτική παρουσίασή τους με έναν εύληπτο τρόπο μέσω του συναισθηματικού γράφου που κατασκευάζεται, μπορεί να προσφέρει άμεσα πληροφορίες για την συναισθηματικό επίπεδο της συζήτησης οι οποίες δεν είναι εμφανείς διαφορετικά. Επίσης, η προσέγγιση μπορεί παρέχει πληροφορίες και για την χρονική εξέλιξη του συναισθηματικού περιεχομένου μιας συζήτησης καθώς αυτή διαδραματίζεται και αλλάζει κάθε χρονική στιγμή.

Μια τρίτη συνιστώσα της έρευνας ήταν η κατανόηση κειμένου και η αυτόματη μετατροπή προτάσεων φυσικής γλώσσας (ΦΓ) σε κατηγορηματική λογική πρώτης τάξης (ΚΛΠΤ). Για τον σκοπό αυτό, αρχικά δημιουργήσαμε μια προσέγγιση για τον προσδιορισμό του επιπέδου δυσκολίας της διαδικασίας φορμαλισμού των προτάσεων η οποία βασίζεται στην ανάλυση της πρότασης ΦΓ, στην ανάλυση της αντίστοιχης έκφρασης ΚΛΠΤ, καθώς και σε χαρακτηριστικά που εξάγονται από την αντιστοίχισή τους. Στην συνέχεια, δημιουργήσαμε μια νέα, καινοτόμα μεθοδολογία για την αυτόματη μετατροπή προτάσεων ΦΓ σε ΚΛΠΤ, η οποία βασίζεται στην συλλογιστική βασισμένη στις περιπτώσεις (case based reasoning). Σχεδιάστηκε με βάση την αρχή ότι προτάσεις ΦΓ που έχουν ίδια ή παρόμοια συντακτική δομή και δέντρα εξαρτήσεων (dependency trees) θα έχουν ίδια ή παρόμοια δομή έκφρασης σε ΚΛΠΤ. Για τον προσδιορισμό της ομοιότητας μεταξύ προτάσεων φυσικής γλώσσας χρησιμοποιήθηκε η μετρική tree edit distance, η οποία υπολογίζεται στα δέντρα εξαρτήσεων. Τα πειραματικά αποτελέσματα της αξιολόγησης έδειξαν ότι η προσέγγιση είναι αποτελεσματική στις περιπτώσεις που η ομοιότητα μεταξύ της νέας προς μετατροπή πρότασης και της πιο όμοιας πρότασης από την βάση γνώσης είναι μεγάλη (0.80-1.00). Όσον αφορά την πολυπλοκότητα του φορμαλισμού προτάσεων, η μεθοδολογία είναι ακριβής στον προσδιορισμό του κατάλληλου επιπέδου δυσκολίας (μέση ακρίβεια 0.94) κάτι που επαληθεύεται και από τις επιδόσεις του μηχανισμού αυτόματης μετατροπής.

Μια άλλη συνιστώσα της έρευνας ήταν ο σχεδιασμός και η μελέτη μιας προσέγγισης για την αυτόματη ανάλυση προτάσεων και τη δημιουργία παραφράσεών τους (λογικά ισοδύναμων προτάσεων). Για το σκοπό αυτό, αναπτύχθηκαν τεχνικές παράφρασης, οι οποίες σχεδιάστηκαν για να μετατρέπουν δομικά μέρη του δέντρου εξάρτησης των προτάσεων σε ισοδύναμες μορφές, καθώς και να τροποποιούν κατάλληλα συγκεκριμένες λέξεις των προτάσεων. Τα πειραματικά αποτελέσματα της αξιολόγησης έδειξαν μια καλή συμπεριφορά των τεχνικών και του εργαλείου που σχεδιάστηκε. Οι λογικά ισοδύναμες παραφράσεις που δημιουργούνται είναι κατάλληλες και διατηρούν την σημασιολογία στο μεγαλύτερο μέρος (76.4%) των περιπτώσεων.

Μια επιπλέον συνιστώσα της έρευνας αφορούσε την ανάλυση εκφράσεων προσώπου και την αυτόματη αναγνώριση του συναισθηματικού τους περιεχομένου. Για το σκοπό αυτό, σχεδιάστηκε ένα πλαίσιο μοντελοποίησης των εκφράσεων του προσώπου, το οποίο ακολουθεί μια αναλυτική, τοπικής εμβέλειας προσέγγιση και αναπαριστά μια έκφραση με 25 χαρακτηριστικά που εξάγει από σημεία όπως, τα μάτια το στόμα και τα φρύδια. Με

βάση τα χαρακτηριστικά που εξάγονται, μέθοδοι ταξινομητών χρησιμοποιούνται για να προσδιορίσουν την ύπαρξη ή όχι συναισθημάτων σε μια έκφραση προσώπου και να αναγνωρίσουν την ακριβή συναισθηματική του κατάσταση. Στην συνέχεια, μελετάται η απόδοση νεύρο-ασαφούς προσέγγισης και μεθόδων μηχανικής μάθησης, όπως random forest και support vector machines, στον προσδιορισμό της ύπαρξης ή όχι συναισθημάτων σε μια έκφραση καθώς και στην αναγνώριση της ακριβούς συναισθηματικής της κατάστασης. Έπειτα, με βάση την απόδοση τους, σχεδιάστηκε ομαδοποιημένο σχήμα ταξινόμησης, το οποίο συνδυάζει ταξινομητές βασισμένους στις τρεις παραπάνω μεθόδους σε ένα σχήμα πλειοψηφίας. Στόχος του ομαδοποιημένου σχήματος αποτελεί η αύξηση της απόδοσης και η αποτελεσματική γενίκευση σε νέα δεδομένα εκφράσεων προσώπου. Το ομαδοποιημένο σχήμα παρουσίασε μια αύξηση της απόδοσης της τάξης του 3.5% σε σχέση με τον καλύτερο ταξινομητή. Επίσης, η απόδοση των ταξινομητών σε γνωστές βάσεις δεδομένων ανέδειξε καλά αποτελέσματα όπου η ακρίβεια τους κυμαινόταν στο εύρος 0.82-0.87 και είναι κοντά στις καλύτερες επιδόσεις προσεγγίσεων της διεθνούς βιβλιογραφίας, οι οποίες αναφέρουν ακρίβεια στο εύρος 0.85-0.99 στις βάσεις δεδομένων που χρησιμοποιήθηκαν για την αξιολόγηση. Ένα βασικό πλεονέκτημα της προσέγγισης μας και των ομαδοποιημένων ταξινομητών σε σχέση με τις καλύτερες μεθόδους της διεθνούς βιβλιογραφίας αποτελεί η κλιμάκωση και η καλή απόδοση σε ετερογενή δεδομένα, όπως δεδομένα εκφράσεων από διαφορετικές βάσεις προσώπων, καθώς και η ενσωμάτωση σε εφαρμογές με δεδομένα του πραγματικού κόσμου.

7.2 Μελλοντική έρευνα

Τα αποτελέσματα της έρευνας της διατριβής θέτουν περαιτέρω κατευθύνσεις στις οποίες θα μπορούσε να κινηθεί η μελλοντική έρευνα. Στην συνέχεια, παρουσιάζονται μερικές πιθανές μελλοντικές κατευθύνσεις.

Μια μελλοντική κατεύθυνση, στο πεδίο της παροχής ανάδρασης στα εκπαιδευτικά συστήματα, αποτελεί η ενσωμάτωση μηχανισμών εξαγωγής γνώσης και μηχανικής μάθησης για την πρόβλεψη των επιδόσεων των εκπαιδευόμενων και την παροχή κατάλληλων υποδείξεων με βάση τις προβλέψεις. Μια ακόμη πιθανή κατεύθυνση που θα μπορούσε να εξετάσει η μελλοντική έρευνα αφορά την χρήση οντολογίας για την αναπαράσταση του πλαισίου κάτι, που σε συνδυασμό με τεχνικές αυτόματης ευθυγράμμισης οντολογιών, θα μπορούσε να συμβάλει στην επεκτασιμότητα του πλαισίου, στην άμεση ενσωμάτωση σε εκπαιδευτικά συστήματα που βασίζονται στις οντολογίες και επίσης στην διασύνδεση του

με λειτουργικές μονάδες όπως το πεδίο γνώσης (domain model) και το μοντέλο μαθητή (student model) των συστημάτων αυτών.

Στο πεδίο της αυτόματης αναγνώρισης του συναισθηματικού περιεχομένου προτάσεων φυσικής γλώσσας, μια μελλοντική ερευνητική κατεύθυνση αφορά την μελέτη επιπρόσθετων τεχνικών ομαδοποίησης ταξινομητών. Πιο συγκεκριμένα, τα ομαδοποιημένα σχήματα που μελετήθηκαν βασίζονται σε ομάδες ταξινομητών που έχουν ομαδοποιηθεί με χρήση τεχνικών bagging και boosting, οι οποίες είναι μέθοδοι διαμέλισης με βάση τα στιγμιότυπα (instance partitioning methods). Μια μελλοντική κατεύθυνση έρευνας αφορά την μελέτη της απόδοσης τεχνικών όπως η random subspace, η οποία είναι μέθοδος διαμέλισης με βάση τα χαρακτηριστικά (feature partitioning method). Επίσης, μια ακόμη κατεύθυνση θα μπορούσε να αποτελέσει ο συνδυασμός με τεχνικές αναπαράστασης κειμένων, όπως η feature hashing, και η μελέτη της απόδοσής τους στην διεργασία της ταξινόμησης.

Μια μελλοντική ερευνητική κατεύθυνση στο πεδίο της αυτόματης μετατροπής προτάσεων φυσικής γλώσσας σε κατηγορηματική λογική πρώτης τάξης αποτελεί η μελέτη και ανάπτυξη τεχνικών για την αυτόματη επέκταση της βάσης γνώσης της μεθοδολογίας με περιπτώσεις (cases). Πιο συγκεκριμένα, θα μπορούσε να μελετηθεί η συμπεριφορά τεχνικών αυτόματης μετατροπής ΚΛΠΤ εκφράσεων σε ισοδύναμη μορφή, οι οποίες σε συνδυασμό με τις τεχνικές αυτόματης παραγωγής παραφράσεων των προτάσεων φυσικής γλώσσας, θα μπορούσαν να προσθέσουν αυτόματα λογικά ισοδύναμες περιπτώσεις στην βάση γνώσης. Επίσης, μια ακόμη κατεύθυνση για μελλοντική έρευνα αφορά την εισαγωγή βαρών στον τρόπο που υπολογίζεται η ομοιότητα μεταξύ προτάσεων.

Στο πεδίο της αυτόματης παραγωγής λογικά ισοδύναμων παραφράσεων η μελλοντική έρευνα θα μπορούσε να εξετάσει αυτόματες μεθόδους αποτίμησης της συντακτικής και της σημασιολογικής ορθότητας των παραγόμενων παραφράσεων. Αρχικά, μια κατεύθυνση αφορά την μελέτη μεθόδων αυτόματου προσδιορισμού της συντακτικής ορθότητας των παραγόμενων παραφράσεων. Σε αυτή την κατεύθυνση, η ενσωμάτωση γραμματικών, όπως οι γραμματικές χωρίς συμφραζόμενα (context free grammars), θα μπορούσαν να εξεταστούν. Επίσης, μια άλλη κατεύθυνση έρευνας θα μπορούσε να εξετάσει τη χρησιμοποίηση τεχνικών για την σημασιολογική ορθότητα και την εννοιολογική διατήρηση του περιεχομένου μεταξύ των παραγόμενων παραφράσεων. Σε αυτή την κατεύθυνση, θα μπορούσαν να εξεταστούν τεχνικές, όπως η λανθάνουσα σημασιολογική δεικτοδότηση και

η εξαγωγή πληροφορίας, για την ομοιότητα εννοιών μεταξύ των παραγόμενων παραφράσεων.

Στο πεδίο της ανάλυσης εκφράσεων προσώπου και του προσδιορισμού του συναισθηματικού περιεχομένου που φέρουν, μια μελλοντική κατεύθυνση αφορά την μοντελοποίηση των χαρακτηριστικών ως γράφημα, όπου θα υπολογίζονται όχι μόνο η σχετική γεωμετρική θέση για κάθε σημείο ενδιαφέροντος, αλλά και η απόσταση του από όλα τα υπόλοιπα. Επίσης, μια ακόμη κατεύθυνση που θα μπορούσε να εξετάσει η μελλοντική έρευνα αφορά την μελέτη τεχνικών για την ακριβή εξαγωγή χαρακτηριστικών κάτω από τις δύσκολες συνθήκες του πραγματικού κόσμου, όπως σε διαφορετικές συνθήκες φωτισμού και κλήσης του προσώπου. Επιπρόσθετα, μια ακόμη ερευνητική κατεύθυνση αφορά την μελέτη *stacking* τεχνικών προσδιορισμού της τελικής απόφασης με βάση τις αποφάσεις των επιμέρους ταξινομητών.

Τέλος, μια ακόμη κατεύθυνση για μελλοντική έρευνα θα μπορούσε να είναι ο σχεδιασμός και η ανάπτυξη πολυμεσικής διεπαφής αλληλεπίδρασης μεταξύ χρήστη και εκπαιδευτικών συστημάτων, η οποία θα ενσωματώνει σε ένα ενιαίο περιβάλλον μηχανισμούς όπως η αυτόματη μετατροπή προτάσεων φυσικής γλώσσας σε ΚΛΠΤ για την κατανόηση της επικοινωνίας με τον χρήστη, η αναγνώριση της συναισθηματικής κατάστασης από την γραφή και την εικόνα με στόχο την καλύτερη προσαρμογή στις ανάγκες του χρήστη. Αυτή η κατεύθυνση αποτελεί μια βασική συνιστώσα εφαρμογής των μεθόδων που διερευνήθηκαν στην παρούσα διατριβή.

Λίστα Δημοσιεύσεων

Δημοσιεύσεις στα Πλαίσια της Διατριβής

Δημοσιεύσεις σε Διεθνή Περιοδικά:

- J4.** **Isidoros Perikos** and Ioannis Hatzilygeroudis, Recognizing Emotions in Text Using Ensemble of Classifiers, International Journal of Engineering Applications of Artificial Intelligence (EAAI), Elsevier, 2016
- J3.** **Isidoros Perikos**, Foteini Grivokostopoulou and Ioannis Hatzilygeroudis, Assistance and Feedback Mechanism in an Intelligent Tutoring System for Teaching Conversion of Natural Language into Logic, International Journal of Artificial Intelligence in Education (IJAIED), Springer, 2016
- J2.** **Isidoros Perikos**, Foteini Grivokostopoulou, Konstantinos Kovas, & Ioannis Hatzilygeroudis, (2016). Automatic estimation of exercises' difficulty levels in a tutoring system for teaching the conversion of natural language into first-order logic. Expert Systems, 33(6), 569-580. doi: 10.1111/exsy.12182
- J1.** Foteini Grivokostopoulou, Ioannis Hatzilygeroudis, **Isidoros Perikos**, Teaching assistance and automatic difficulty estimation in converting first order logic to clause form, Artificial Intelligence Review, Springer, Vol. 42, Issue 3, 2014

Δημοσιεύσεις σε Διεθνή Συνέδρια:

- C13.** **Isidoros Perikos** and Ioannis Hatzilygeroudis "A Case-Based Reasoning Approach to Convert Natural Language into First Order Logic" IEEE International Conference on Computational Science and Engineering (IEEE CSE 2016), Paris, France, 2016
- C12.** **Isidoros Perikos** and Ioannis Hatzilygeroudis "A Methodology for Generating Natural Language Paraphrases", International Conference on Information, Intelligence, Systems and Applications (IISA 2016), Chalkidiki, Greece, 2016
- C11.** **Isidoros Perikos** and Ioannis Hatzilygeroudis "A Classifier Ensemble Approach to Detect Emotions Polarity in Social Media", International Conference on Web Information Systems and Technologies (WEBIST 2016), Rome, Italy, 2016
- C10.** **Isidoros Perikos**, Epameinondas Ziakopoulos and Ioannis Hatzilygeroudis "Recognize Emotions from Facial Expressions Using a SVM and Neural Network Schema",

- International Conference Engineering Applications of Neural Networks (EANN 2015), Rodos, Greece, 2015
- C9. Isidoros Perikos**, Epameinondas Ziakopoulos and Ioannis Hatzilygeroudis "Recognizing Emotions from Facial Expressions using Neural Network", IFIP International Conference on Artificial Intelligence Applications and Innovations (AIAI 2014), September 2014, Rodos, Greece
- C8.** Foteini Grivokostopoulou, **Isidoros Perikos**, Ioannis Hatzilygeroudis, 'Using Semantic Web Technologies in a Web Based System for Personalized Learning AI Course', IEEE International Conference on Technology for Education (IEEE T4E 2014), 18-21 December, India, 2014
- C7.** Andreas Kanavos, **Isidoros Perikos**, Pantelis Vikatos, Ioannis Hatzilygeroudis, Christos Makris and Athanasios Tsakalidis, "Conversation Emotional Modelling in Social Networks", IEEE International Conference on Tools with Artificial Intelligence (IEEE ICTAI 2014), Limassol, Cyprus, November 2014
- C6.** Foteini Grivokostopoulou, **Isidoros Perikos**, Ioannis Hatzilygeroudis, "Utilizing Semantic Web Technologies and Data Mining Techniques to Analyze Students Learning and Predict Final Performance", IEEE International Conference on Teaching, Assessment and Learning for Engineering (IEEE TALE 2014), New Zealand, December 2014
- C5. Isidoros Perikos**, Ioannis Hatzilygeroudis "Recognizing Emotion Presence in Natural Language Sentences", International Conference Engineering Applications of Neural Networks (EANN 2013), Chalkidiki, Greece, September 2013
- C4. Isidoros Perikos**, Ioannis Hatzilygeroudis, Foteini Grivokostopoulou, "Help Generation in a System for Learning NL to FOL Conversion", IASTED International Conference on Computers and Advanced Technology in Education (CATE 2011), Cambridge, United Kingdom, 2011
- C3. Isidoros Perikos**, Ioannis Hatzilygeroudis, Foteini Grivokostopoulou, Konstantinos Kovas "Difficulty Estimator for Translating Natural Language into First Order Logic", 3rd International Conference on Intelligent Decision Technologies, (KES-IDT 2011), Piraeus, Greece, 2011
- C2. Isidoros Perikos**, Ioannis Hatzilygeroudis, Foteini Grivokostopoulou, "Teaching Assistant Tools for NL to FOL Conversion" IADIS Multi Conference on Computer Science and Information Systems, Rome, Italy, 2011

- C1. Isidoros Perikos** and Ioannis Hatzilygeroudis “A Web-Based Interactive System for Learning NL to FOL Conversion”, 2nd International Conference on Intelligent Interactive Multimedia Systems and Services (KES-IIMSS 2009), Venice, Italy, 2009

Δημοσιεύσεις σε Workshops:

- W1. Isidoros Perikos** and Ioannis Hatzilygeroudis “Hybrid Approaches to Determine Existence of Emotions in Facial Gestures” 6th International Workshop on Combinations of Intelligent Methods and Applications (CIMA 2016), The Hague, Holland, 2016

Επιπρόσθετες Δημοσιεύσεις

Δημοσιεύσεις σε Διεθνή Περιοδικά:

- J2.** Foteini Grivokostopoulou, **Isidoros Perikos**, Ioannis Hatzilygeroudis, "An Educational System for Learning Search Algorithms and Automatically Assessing Student Performance", International Journal of Artificial Intelligence in Education (IAIED), Springer, 2016
- J1.** Themistoklis Chronopoulos, Ioannis Hatzilygeroudis, **Isidoros Perikos**, Konstantinos Kovas "A web-based example-tracing tutor for building sentences in First Order Logic", Engineering Intelligent Systems, 2010, Vol 18, Nos 3/4 September/December 2010

Δημοσιεύσεις σε Βιβλία-Τόμους:

- B2.** Foteini Grivokostopoulou, **Isidoros Perikos**, Ioannis Hatzilygeroudis, "Difficulty estimation of exercises on tree-based search algorithms using neuro-fuzzy and neuro-symbolic approaches" , Advances in Combining Intelligent Methods, Springer 2016
- B1.** Foteini Grivokostopoulou, Ioannis Hatzilygeroudis, **Isidoros Perikos**, "An intelligent system for learning first order logic to clause form conversion", Multicultural Awareness and Technology in Higher Education: Global Perspectives, 2014.

Δημοσιεύσεις σε Διεθνή Συνέδρια:

- C18.** **Isidoros Perikos**, Foteini Grivokostopoulou, Konstantinos Kovas and Ioannis Hatzilygeroudis "Exploring the Educational Capabilities of Social Media in High Schools", International Conference on Computer Supported Education (CSEDU 2016), Rome, Italy, 2016
- C17.** Foteini Grivokostopoulou, **Isidoros Perikos** and Ioannis Hatzilygeroudis, "An Educational Game for Teaching Search Algorithms", International Conference on Computer Supported Education (CSEDU 2016), Rome, Italy, 2016

-
- C16.** Andreas Kanavos, **Isidoros Perikos**, Ioannis Hatzilygeroudis and Athanasios Tsakalidis "Integrating User's Emotional Behavior for Community Detection in Social Networks", International Conference on Web Information Systems and Technologies (WEBIST 2016), Rome, Italy, 2016
- C15.** **Isidoros Perikos**, Foteini Grivokostopoulou, Konstantinos Kovas and Ioannis Hatzilygeroudis "Utilizing Social Networks and Web 2.0 Technologies in Education", IEEE International Conference on Technology for Education (IEEE T4E 2015), 10-12 December, India, 2015
- C14.** **Isidoros Perikos**, Ioannis Hatzilygeroudis, "Embodied Conversational Agents: A methodology for Learning to Express Facial Emotions", International Conference on Conference on Information, Intelligence, Systems and Applications (IISA 2015), Corfu, Greece, 2015
- C13.** Foteini Grivokostopoulou, **Isidoros Perikos** and Ioannis Hatzilygeroudis "Estimating the Difficulty of Exercises on Search Algorithms using a Neuro-Fuzzy Approach", IEEE International Conference on Tools with Artificial Intelligence (IEEE ICTAI 2015), Vietri sul Mare, Italy, November 2015
- C12.** Andreas Kanavos, **Isidoros Perikos** "Towards Detecting Emotional Communities in Twitter", IEEE International Conference on Research Challenges in Information Science (IEEE RCIS 2015), Athens, Greece, May 2015
- C11.** **Isidoros Perikos**, Foteini Grivokostopoulou, Konstantinos Kovas and Ioannis Hatzilygeroudis "Assisting Tutors to Utilize WEB 2.0 Tools in Education", IADIS Multi Conference on Computer Science and Information Systems, Las Palmas, Spain, 2015
- C10.** Foteini Grivokostopoulou, **Isidoros Perikos**, Konstantinos Kovas, Ioannis Hatzilygeroudis, "Teaching Renewable Energy using 3D Virtual Worlds", IEEE International Conference on Advanced Learning Technologies (IEEE ICALT 2015), Taiwan, July 2015
- C9.** **Isidoros Perikos**, Panagiotis Angelopoulos, Michael Paraskevas, Thomas Zarouchas, Giannis Tzimas, "A Methodology for Evaluation and Knowledge Extraction from On-Going Learning Process", IEEE International Conferences on Computational Science and Engineering (IEEE CSE 2014), China, December 2014
- C8.** Andreas Kanavos, **Isidoros Perikos**, Pantelis Vikatos, Ioannis Hatzilygeroudis, Christos Makris and Athanasios Tsakalidis "Modeling ReTweet Diffusion using Emotional Content", IFIP International Conference on Artificial Intelligence Applications and Innovations (AIAI 2014), September 2014, Rodos, Greece
-

- C7.** **Isidoros Perikos**, Foteini Grivokostopoulou, Ioannis Hatzilygeroudis, "Automatic Marking of NL to FOL Conversions", IASTED International Conference on Computers and Advanced Technology in Education (CATE 2012), Napoli, Italy, 2012
- C6.** Foteini Grivokostopoulou, **Isidoros Perikos**, Ioannis Hatzilygeroudis, "A Web-Based Interactive System for Learning FOL to CF Conversions" IADIS Multi Conference on Computer Science and Information Systems, Lisbon, Portugal, 2012
- C5.** Ioannis Hatzilygeroudis, Foteini Grivokostopoulou, **Isidoros Perikos** "Teaching Aspects of Constraint Satisfaction Algorithms via a Game", Third AAAI Symposium on Educational Advances in Artificial Intelligence (EAAI 2012), Toronto, Canada, 2012
- C4.** Foteini Grivokostopoulou, **Isidoros Perikos**, Ioannis Hatzilygeroudis "An Automatic Marking System for FOL to CF Conversions", IEEE International Conference on Teaching, Assessment and Learning for Engineering 2012 (IEEE TALE 2012), Hong Kong, China, 2012
- C3.** Ioannis Hatzilygeroudis, Foteini Grivokostopoulou, **Isidoros Perikos** "Using Game-based Learning in Teaching CS Algorithms", IEEE International Conference on Teaching, Assessment and Learning for Engineering 2012, (IEEE TALE 2012), Hong Kong, China, 2012
- C2.** Foteini Grivokostopoulou, **Isidoros Perikos**, Ioannis Hatzilygeroudis "Assistant tools for teaching FOL to CF conversion", IFIP International Conference on Artificial Intelligence Applications and Innovations (AIAI 2012), Halkidiki, Greece, 2012
- C1.** Themistoklis Chronopoulos, Isidoros Perikos, Ioannis Hatzilygeroudis "An example-tracing tutor for teaching NL to FOL conversion", 6th IFIP Conference on Artificial Intelligence Applications and Innovations (AIAI 2010), Larnaca, Cyprus, 2010

Δημοσιεύσεις σε Workshops:

- W1.** Foteini Grivokostopoulou, Isidoros Perikos, Ioannis Hatzilygeroudis "An Intelligent Tutoring System for Teaching FOL Equivalence", First Workshop on AI-supported Education for Computer Science (AIEDCS 2013), Memphis, USA, July 2013, pp. 20-29

References

- Akputu, K. O., Seng, K. P., & Lee, Y. L. (2013). Facial emotion recognition for intelligent tutoring environment. In *2nd International Conference on Machine Learning and Computer Science (IMLCS 2013)* (pp. 9-13).
- Alepis, E., & Virvou, M. (2011). Automatic generation of emotions in tutoring agents for affective e-learning in medical education. *Expert Systems with Applications*, 38(8), 9840-9847.
- Aleven, V., & Koedinger, K. R. (2000). Limitations of student control: Do students know when they need help?. In *Intelligent tutoring systems* (pp. 292-303). Springer Berlin Heidelberg.
- Aleven, V., McLaren, B., Roll, I., & Koedinger, K. (2006). Toward meta-cognitive tutoring: A model of help seeking with a Cognitive Tutor. *International Journal of Artificial Intelligence and Education*, 16, 101 – 128.
- Aleven, V., Stahl, E., Schworm, S., Fischer, F., & Wallace, R. (2003). Help seeking and help design in interactive learning environments. *Review of Educational Research*, 73(3), 277-320.
- Alm, C. O., Roth, D., & Sproat, R. (2005). Emotions from text: machine learning for text-based emotion prediction. In *Proceedings of the conference on human language technology and empirical methods in natural language processing* (pp. 579-586). Association for Computational Linguistics.
- Androutsopoulos, I., & Malakasiotis, P. (2009). A survey of paraphrasing and textual entailment methods. arXiv preprint arXiv:0912.3747.
- Anusha, V., & Sandhya, B. (2015). A learning based emotion classifier with semantic text processing. In *Advances in Intelligent Informatics* (pp. 371-382). Springer International Publishing.
- Arca, S., Campadelli, P., & Lanzarotti, R. (2004). An automatic feature-based face recognition system. In *Proceedings of the 5th International Workshop on Image Analysis for Multimedia Interactive Services (WIAMIS'04)*.
- Aung, D. M., & Aye, N. (2012). A Facial Expression Classification using Histogram Based Method. *International Proceedings of Computer Science and Information Technology*, 58, 1.
- Baccianella, S., Esuli, A., & Sebastiani, F. (2010). SentiWordNet 3.0: An Enhanced Lexical Resource for Sentiment Analysis and Opinion Mining. In *LREC* (Vol. 10, pp. 2200-2204).
- Baral, C., Dzifcak, J., Gonzalez, M. A., & Zhou, J. (2011). Using inverse λ and generalization to translate english to formal languages. In *Proceedings of the Ninth International Conference on Computational Semantics* (pp. 35-44). Association for Computational Linguistics.
- Barker-Plummer, D., Cox, R., & Dale, R. (2009). Dimensions of difficulty in translating natural language into first order logic. *International Working Group on Educational Data Mining*.
- Barker-Plummer, D., Cox, R., & Dale, R. (2011). Student translations of natural language into logic: The Grade Grinder corpus release 1.0. In *Proceedings of the 4th international conference on educational data mining* (pp. 51-60).

- Barker-Plummer, D., Cox, R., Dale, R., & Etchemendy, J. (2008). An empirical study of errors in translating natural language into logic. In *Proceedings of the 30th Annual Meeting of the Cognitive Science Society/CogSci* (pp. 505-510).
- Barker-Plummer, D., Dale, R., & Cox, R. (2012). Using edit distance to analyse errors in a natural language to logic translation corpus. In *Proceedings of the International Conference on Educational Data Mining*, 134.
- Berant, J., & Liang, P. (2014). Semantic Parsing via Paraphrasing. In *ACL* (1) (pp. 1415-1425).
- Berlanga, A. J., Van Rosmalen, P., Boshuizen, H. P., & Sloep, P. B. (2012). Exploring formative feedback on textual assignments with the help of automatically created visual representations. *Journal of Computer Assisted Learning*, 28(2), 146-160
- Bettadapura, V. (2012). Face expression recognition and analysis: the state of the art. *arXiv preprint arXiv:1203.6722*.
- Bille, P. (2005). A survey on tree edit distance and related problems. *Theoretical computer science*, 337(1), 217-239.
- Blum, A. L., & Rivest, R. L. (1992). Training a 3-node neural network is NP-complete. *Neural Networks*, 5(1), 117-127.
- Bradley, M. M., & Lang, P. J. (1999). *Affective norms for English words (ANEW): Instruction manual and affective ratings* (pp. 1-45). Technical report C-1, the center for research in psychophysiology, University of Florida.
- Breiman, L. (1996). Bagging predictors. *Machine learning*, 24(2), 123-140.
- Brilis, S., Gkatzou, E., Koursoumis, A., Talvis, K., Kermanidis, K. L., & Karydis, I. (2012). Mood classification using lyrics and audio: A case-study in greek music. In *IFIP International Conference on Artificial Intelligence Applications and Innovations* (pp. 421-430). Springer Berlin Heidelberg.
- Calvo, R. A., & D'Mello, S. (2010). Affect detection: An interdisciplinary review of models, methods, and their applications. *IEEE Transactions on affective computing*, 1(1), 18-37.
- Canny, J. (1986). A computational approach to edge detection. *IEEE Transactions on pattern analysis and machine intelligence*, (6), 679-698.
- Chaffar, S., & Inkpen, D. (2011). Using a heterogeneous dataset for emotion analysis in text. In *Canadian Conference on Artificial Intelligence* (pp. 62-67). Springer Berlin Heidelberg.
- Chatterjee, C. (2016). Use of intelligent systems to teach computer engineering. In *Circuit, Power and Computing Technologies (ICCPCT), 2016 International Conference on* (pp. 1-4). IEEE.
- Chaumartin, F. R. (2007). UPAR7: A knowledge-based system for headline sentiment tagging. In *Proceedings of the 4th International Workshop on Semantic Evaluations* (pp. 422-425). Association for Computational Linguistics.
- Chronopoulos, T., Hatzilygeroudis, I., Perikos, I., & Kostas, K. (2010). A web-based example-tracing tutor for formalising sentences in first order logic. *International Journal of Engineering Intelligent Systems Electrical Engineering Communications*, 18(3), 159.

- Chronopoulos, T., Perikos, I., & Hatzilygeroudis, I. (2010). An example-tracing tutor for teaching NL to FOL conversion. In *IFIP International Conference on Artificial Intelligence Applications and Innovations* (pp. 170-178). Springer Berlin Heidelberg.
- Clavel, C. (2015). Surprise and human-agent interactions. *Review of Cognitive Linguistics*, 13(2), 461-477.
- Clavel, C., & Callejas, Z. (2016). Sentiment analysis: from opinion mining to human-agent interaction. *IEEE Transactions on Affective Computing*, 7(1), 74-93.
- Cohen, J. (1960). A coefficient of agreement for nominal scales. *Educational and psychological measurement*, 20(1), 37-46.
- Corbett, A. T., & Anderson, J. R. (2001). Locus of feedback control in computer-based tutoring: Impact on learning rate, achievement and attitudes. In *Proceedings of the SIGCHI conference on Human factors in computing systems* (pp. 245-252). ACM.
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine learning*, 20(3), 273-297.
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *psychometrika*, 16(3), 297-334.
- Costantini, S., Florio, N., & Paolucci, A. (2011). A framework for structured knowledge extraction and representation from natural language via deep sentence analysis. In *CILC* (pp. 297-310).
- Da Silva, N. F., Hruschka, E. R., & Hruschka, E. R. (2014). Tweet sentiment analysis with classifier ensembles. *Decision Support Systems*, 66, 170-179.
- Danisman, T., & Alpkocak, A. (2008). Feeler: Emotion classification of text using vector space model. In *AISB 2008 Convention Communication, Interaction and Social Intelligence* (Vol. 1, p. 53).
- De Choudhury, M., Gamon, M., & Counts, S. (2012). Happy, Nervous or Surprised? Classification of Human Affective States in Social Media. In *ICWSM*.
- De Mantaras, R. L., McSherry, D., Bridge, D., Leake, D., Smyth, B., Craw, S., ... & Keane, M. (2005). Retrieval, reuse, revision and retention in case-based reasoning. *The Knowledge Engineering Review*, 20(03), 215-240.
- De Marneffe, M. C., MacCartney, B., & Manning, C. D. (2006). Generating typed dependency parses from phrase structure parses. In *Proceedings of LREC* (Vol. 6, No. 2006, pp. 449-454).
- Dietterich, T. G. (2000). Ensemble methods in machine learning. In *International workshop on multiple classifier systems* (pp. 1-15). Springer Berlin Heidelberg.
- Dolan, B., Brockett, C., & Quirk, C. (2005). Microsoft research paraphrase corpus. Retrieved February, 2016.
- Dolan, B., Quirk, C., & Brockett, C. (2004). Unsupervised construction of large paraphrase corpora: Exploiting massively parallel news sources. In *Proceedings of the 20th international conference on Computational Linguistics* (p. 350). Association for Computational Linguistics.
- Ekman, P. (1999). Basic Emotions In T. Dalgleish and T. Power (Eds.) *The Handbook of Cognition and Emotion* Pp. 45-60.

- Ekman, P., & Friesen, W., (1978). Facial Action Coding System: A Technique for the Measurement of Facial Movement. Consulting Psychologists Press, Palo Alto
- Ekman, P., & Rosenberg, E. L. (1997). *What the face reveals: Basic and applied studies of spontaneous expression using the Facial Action Coding System (FACS)*. Oxford University Press, USA.
- Fiedler, A., & Tsovaltzi, D. (2003). Automating Hinting in an Intelligent Tutorial Dialog System for Mathematics. In Proceedings of the IJCAI 2003 Workshop on Knowledge Representation and Automated Reasoning for E-Learning Systems, Acapulco, Mexico
- Finch, D., Peacock, M., Lazdowski, D., & Hwang, M. (2015). Managing emotions: A case study exploring the relationship between experiential learning, emotions, and student performance. *The International Journal of Management Education*, 13(1), 23-36
- Finkel, J. R., Grenager, T., & Manning, C. (2005). Incorporating non-local information into information extraction systems by gibbs sampling. In Proceedings of the 43rd Annual Meeting on Association for Computational Linguistics ,363-370, Association for Computational Linguistics.
- Fujita, A., & Isabelle, P. (2015). Expanding paraphrase lexicons by exploiting lexical variants. In Proceedings of Human Language Technologies: The 2015 Annual Conference of the North American Chapter of the Association for Computational Linguistics (NAACL-HLT).
- Funkhouser, A. (2016). Annotation of Utterances for Conversational Nonverbal Behaviors.
- Ganitkevitch, J., & Callison-Burch, C. (2014). The Multilingual Paraphrase Database. In LREC (pp. 4276-4283).
- Gaur, S., Vo, N. H., Kashihara, K., & Baral, C. (2014). Translating simple legal text to formal representations. In JSAI International Symposium on Artificial Intelligence (pp. 259-273). Springer Berlin Heidelberg.
- Gerdes, A., Heeren, B., Jeuring, J., & van Binsbergen, L. T. (2016). Ask-Elle: an adaptable programming tutor for Haskell giving automated feedback. *International Journal of Artificial Intelligence in Education*, 1-36.
- Gheorghiu, R., Vanlehn, K. (2008). XTutor: an intelligent tutor system for science and math based on excel. In: Woolf BP et al. (eds) Intelligent tutoring systems: 9th international conference, ITS 2008. Springer, Germany
- Goldin, I. M., Koedinger, K. R., & Aleven, V. (2013). Hints: You can't have just one. In S. K. D'Mello, R. A. Calvo, & A. Olney (Eds.), *In Proceedings of the 6th International Conference on Educational Data Mining (EDM 2013)* (pp. 232–235).
- Gouli, E., Gogoulou, A., Papanikolaou, K. A., & Grigoriadou, M. (2006). An adaptive feedback framework to support reflection, guiding and tutoring. *Advances in web-based education: Personalized learning environments*, 178-202.
- Grefenstette, E., & Sadrzadeh, M. (2015). Concrete Models and Empirical Evaluations for the Categorical Compositional Distributional Model of Meaning. *Computational Linguistics*. Vol. 41, No. 1, pp. 71-118

- Grivokostopoulou, F., & Hatzilygeroudis, I. (2013). Teaching AI Search Algorithms in a Web-Based Educational System. *Proceedings of the IADIS International Conference e-Learning 2013*, pp. 83-90, 23 – 26 July, Prague, Czech Republic.
- Grivokostopoulou, F., Hatzilygeroudis, I., & Perikos, I. (2014a). Teaching assistance and automatic difficulty estimation in converting first order logic to clause form. *Artificial Intelligence Review*, 42(3), 347-367.
- Grivokostopoulou, F., Hatzilygeroudis, I., & Perikos, I. (2014b). An Intelligent System for Learning First Order Logic to Clause Form Conversion. *Multicultural Awareness and Technology in Higher Education: Global Perspectives: Global Perspectives*, 265.
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2016a). An Educational System for Learning Search Algorithms and Automatically Assessing Student Performance. *International Journal of Artificial Intelligence in Education*, 1-34.
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2016b). Difficulty estimation of exercises on tree-based search algorithms using neuro-fuzzy and neuro-symbolic approaches, *Advances in Combining Intelligent Methods*, Springer (to appear).
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2015). Estimating the Difficulty of Exercises on Search Algorithms Using a Neuro-Fuzzy Approach. In *Tools with Artificial Intelligence (ICTAI), 2015 IEEE 27th International Conference on* (pp. 866-872). IEEE.
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2014a). Utilizing semantic web technologies and data mining techniques to analyze students learning and predict final performance. In *Teaching, Assessment and Learning (TALE), 2014 International Conference on* (pp. 488-494). IEEE.
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2014b). Using Semantic Web Technologies in a Web Based System for Personalized Learning AI Course. In *Technology for Education (T4E), 2014 IEEE Sixth International Conference on* (pp. 257-260). IEEE.
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2013). An intelligent tutoring system for teaching fol equivalence. In *The First Workshop on AI-supported Education for Computer Science (AIEDCS 2013)* (p. 20-29).
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2012a). An automatic marking system for FOL to CF conversions. In *Teaching, Assessment and Learning for Engineering (TALE), 2012 IEEE International Conference on* (pp. H1A-7). IEEE.
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2012b). Assistant tools for teaching FOL to CF conversion. In *IFIP International Conference on Artificial Intelligence Applications and Innovations* (pp. 306-315). Springer Berlin Heidelberg.
- Grivokostopoulou, F., Perikos, I., & Hatzilygeroudis, I. (2012c). A web-based interactive system for learning FOL to CF conversion. In *Proceedings of the IADIS international conference e-learning* (pp. 287-294).
- Ghimire, D., & Lee, J. (2013). Geometric feature-based facial expression recognition in image sequences using multi-class adaboost and support vector machines. *Sensors*, 13(6), 7714-7734.

- Hatcher, L. (1994). A step-by-step approach to using the SAS(R) system for factor analysis and structural equation modeling. Cary, NC: SAS Institute.
- Hattie, J., & Timperley, H. (2007). The power of feedback. *Review of Educational Research*, 77(1), 81-112.
- Hatzilygeroudis, I. (2007). Teaching NL to FOL and FOL to CF Conversions. In *FLAIRS Conference* (pp. 309-314).
- Hatzilygeroudis, I., & Perikos, I. (2009). A web-based interactive system for learning NL to FOL conversion. In *New Directions in Intelligent Interactive Multimedia Systems and Services-2* (pp. 297-307). Springer Berlin Heidelberg.
- Havnes, A., Smith, K., Dysthe, O., & Ludvigsen, K. (2012). Formative assessment and feedback: Making learning visible. *Studies in Educational Evaluation*, 38(1), 21-27.
- Hirashima, T., Horiguchi, T., Kashiara, A., & Toyoda, J. (1998). Error-based simulation for error-visualization and its management. *International Journal of Artificial Intelligence in Education*, 9(1-2), 17-31.
- Ho, D. T., & Cao, T. H. (2012). A high-order hidden Markov model for emotion detection from textual data. In *Pacific Rim Knowledge Acquisition Workshop* (pp. 94-105). Springer Berlin Heidelberg.
- Hyafil, L., & Rivest, R. L. (1976). Constructing optimal binary decision trees is NP-complete. *Information Processing Letters*, 5(1), 15-17.
- Janssen, T. M. (1996). Compositionality. Institute for Logic, Language and Computation (ILLC), University of Amsterdam
- Kamp, H., & Reyle, U. (2013). From discourse to logic: Introduction to modeltheoretic semantics of natural language, formal logic and discourse representation theory (Vol. 42). Springer Science & Business Media.
- Kampeera, W., & Cardey-Greenfield, S. (2012). Building a Lexically and Semantically-Rich Resource for Paraphrase Processing. In *Advances in Natural Language Processing* (pp. 138-143). Springer Berlin Heidelberg.
- Kanade, T., Cohn, J. F., & Tian, Y. (2000). Comprehensive database for facial expression analysis. In *Automatic Face and Gesture Recognition, 2000. Proceedings. Fourth IEEE International Conference on* (pp. 46-53). IEEE.
- Kanavos A., Perikos I., Hatzilygeroudis I. and Tsakalidis A. (2016). Integrating User's Emotional Behavior for Community Detection in Social Networks. In *Proceedings of the 12th International Conference on Web Information Systems and Technologies - Volume 1: SRIS*, ISBN 978-989-758-186-1, (pp. 355-362.)
- Kanavos, A., Perikos, I., Vikatos, P., Hatzilygeroudis, I., Makris, C., & Tsakalidis, A. (2014). Conversation Emotional Modeling in Social Networks. In *2014 IEEE 26th International Conference on Tools with Artificial Intelligence* (pp. 478-484). IEEE.
- Kate, R. J., & Mooney, R. J. (2006). Using string-kernels for learning semantic parsers. In *Proceedings of the 21st International Conference on Computational Linguistics and the 44th*

- annual meeting of the Association for Computational Linguistics (pp. 913-920). Association for Computational Linguistics.
- Khan, A. M., & Lawo, M. (2016). From Physiological data to Emotional States: Conducting a User Study and Comparing Machine Learning Classifiers. *Sensors & Transducers*, 201(6), 78.
- Khan, A. M., & Lawo, M. (2016). Recognizing Emotional States Using Physiological Devices. In *International Conference on Human-Computer Interaction* (pp. 164-171). Springer International Publishing.
- Khan, A. M., & Lawo, M. (2016). Developing a System for Recognizing the Emotional States Using Physiological Devices. In *Intelligent Environments (IE), 2016 12th International Conference on* (pp. 48-53). IEEE.
- Kirk, N. H. (2013). Towards Structural Natural Language Formalization: Mapping Discourse to Controlled Natural Language. arXiv preprint arXiv:1312.2087.
- Kittusamy, S. R. V., & Chakrapani, V. (2012). Facial expressions recognition using eigenspaces. *Journal of Computer Science*, 8(10), 1674.
- Koedinger, K. R., & Aleven, V. (2007). Exploring the assistance dilemma in experiments with cognitive tutors. *Educational Psychology Review*, 19(3), 239-264.
- Kotsia, I., & Pitas, I. (2007). Facial expression recognition in image sequences using geometric deformation features and support vector machines. *IEEE Transactions on Image Processing*, 16(1), 172-187.
- Koutlas, A., & Fotiadis, D. I. (2008, October). An automatic region based methodology for facial expression recognition. In *Systems, Man and Cybernetics, 2008. SMC 2008. IEEE International Conference on* (pp. 662-666). IEEE.
- Kulhavy, R. W., & Stock, W. A. (1989). Feedback in written instruction: The place of response certitude. *Educational Psychology Review*, 1(4), 279-308.
- Kuncheva, L. I. (2004). *Combining pattern classifiers: methods and algorithms*. John Wiley & Sons.
- Levenshtein, V. I. (1966). Binary codes capable of correcting deletions, insertions and reversals. In *Soviet physics doklady* (Vol. 10, p. 707).
- Liebold, B., & Ohler, P. (2013). Multimodal Emotion Expressions of Virtual Agents, Mimic and Vocal Emotion Expressions and Their Effects on Emotion Recognition. In *Affective Computing and Intelligent Interaction (ACII), 2013 Humaine Association Conference on* (pp. 405-410). IEEE.
- Liebold, B., Richter, R., Teichmann, M., Hamker, F. H., & Ohler, P. (2015). Human Capacities for Emotion Recognition and their Implications for Computer Vision. *i-com*, 14(2), 126-137.
- Li, S., Zong, C., & Wang, X. (2007, August). Sentiment classification through combining classifiers with multiple feature sets. In *2007 International Conference on Natural Language Processing and Knowledge Engineering* (pp. 135-140). IEEE.
- Liu, B., & Zhang, L. (2012). A survey of opinion mining and sentiment analysis. In *Mining text data* (pp. 415-463). Springer US.
- Lopez, L. (2009). Effects of delayed and immediate feedback in the computer-based testing environment. ProQuest. Doctoral dissertation.

- Lozano-Monazor, E., López, M. T., Fernández-Caballero, A., & Vigo-Bustos, F. (2014). Facial expression recognition from webcam based on active shape models and support vector machines. In *International Workshop on Ambient Assisted Living* (pp. 147-154). Springer International Publishing.
- Lyons, M., Akamatsu, S., Kamachi, M., & Gyoba, J. (1998). Coding facial expressions with gabor wavelets. In *Automatic Face and Gesture Recognition, 1998. Proceedings. Third IEEE International Conference on* (pp. 200-205). IEEE.
- MacCartney, B., & Manning, C. D. (2007). Natural logic for textual inference. In *Proceedings of the ACL-PASCAL Workshop on Textual Entailment and Paraphrasing* (pp. 193-200). Association for Computational Linguistics.
- Madnani, N., & Dorr, B. J. (2010). Generating phrasal and sentential paraphrases: A survey of data-driven methods. *Computational Linguistics*, 36(3), 341-387.
- Maestro-Prieto, J. A., & Simon-Hurtado, A. (2015). Learner-Adaptive Pedagogical Model in SIAL, an Open-Ended Intelligent Tutoring System for First Order Logic. In *Proceedings of the International Conference on Artificial Intelligence in Education* (pp. 702-705). Springer International Publishing.
- Maestro-Prieto, J. A., Simón-Hurtado, M. A., de-Paz-Santana, J. F., & Villarrubia-González, G. (2013). A MAS for Teaching Computational Logic. In *Distributed Computing and Artificial Intelligence* (pp. 209-217). Springer International Publishing.
- Malakasiotis, P. (2009). Paraphrase recognition using machine learning to combine similarity measures. In *Proceedings of the ACL-IJCNLP 2009 Student Research Workshop* (pp. 27-35). Association for Computational Linguistics.
- Mandernach, B. J. (2005). Relative effectiveness of computer-based and human feedback for enhancing student learning. *The Journal of Educators Online*, 2(1).
- Marie-Catherine de Marneffe and Christopher D. Manning. 2008. Stanford typed dependencies manual. <http://nlp.stanford.edu/downloads/lex-parser.shtml>
- Mathan, S. A., & Koedinger, K. R. (2005). Fostering the intelligent novice: Learning from errors with metacognitive tutoring. *Educational Psychologist*, 40(4), 257-265.
- Mathan, S., & Koedinger, K. R. (2003). Recasting the feedback debate: Benefits of tutoring error detection and correction skills. In *Proceedings of the International Conference on Artificial Intelligence in Education* (pp. 13-20).
- McKendree, J. (1990). Effective feedback content for tutoring complex skills. *Human Computer Interaction*, 5, 381-413.
- Medhat, W., Hassan, A., & Korashy, H., 2014. Sentiment analysis algorithms and applications: A survey. *Ain Shams Engineering Journal*, 5(4), 1093-1113.
- Mehay, D. N., & White, M. (2012). Shallow and Deep Paraphrasing for Improved Machine Translation Parameter Optimization. In *The AMTA 2012 Workshop on Monolingual Machine Translation, MONOMT*.
- Mehrabian, A. (1968) Communication without words. *Psychology Today*, vol. 2, no. 4, pp. 53-56

- Michel, P., & El Kaliouby, R. (2005). Facial expression recognition using support vector machines. In *The 10th International Conference on Human-Computer Interaction, Crete, Greece*.
- Mills, M. T., & Bourbakis, N. G. (2014). Graph-based methods for natural language processing and understanding—a survey and analysis. *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, 44(1), 59-71.
- Mills, M., Psarologou, A., & Bourbakis, N. (2013). Modeling Natural Language Sentences into SPN Graphs. In *Tools with Artificial Intelligence (ICTAI)*, 2013 IEEE 25th International Conference on (pp. 889-896). IEEE.
- Mitrovic, A., Martin, B., & Suraweera, P. (2007). Intelligent tutors for all: The constraint-based approach. *IEEE Intelligent Systems*, (4), 38-45
- Mitrovic, A., Martin, B., Suraweera, P., Zakharov, K., Milik, N., Holland, J., & McGuigan, N. (2009). ASPIRE: an authoring system and deployment environment for constraint
- Moreno, R., & Mayer, R. (2007). Interactive multimodal learning environments. *Educational Psychology Review*, 19(3), 309-326.
- Mory, E. H. (2004). Feedback research revisited. *Handbook of research on educational communications and technology*, 2, 745-783.
- Narayan, S., Reddy, S., & Cohen, S. B. (2016). Paraphrase Generation from Latent-Variable PCFGs for Semantic Parsing. arXiv preprint arXiv:1601.06068.
- Narciss, S. (2008). Feedback strategies for interactive learning tasks. *Handbook of research on educational communications and technology*, 3, 125-144.
- Neviarouskaya, A., Prendinger, H., & Ishizuka, M. (2007). Textual affect sensing for sociable and expressive online communication. In *International Conference on Affective Computing and Intelligent Interaction* (pp. 218-229). Springer Berlin Heidelberg.
- Orrite, C., Rodríguez, M., Martínez, F., & Fairhurst, M. (2008, September). Classifier ensemble generation for the majority vote rule. In *Iberoamerican Congress on Pattern Recognition* (pp. 340-347). Springer Berlin Heidelberg.
- Ortony, A., Clore, G. L., & Collins, A. (1988). *The cognitive structure of emotions* Cambridge University Press Cambridge.
- Osherenko, A., & André, E. (2007). Lexical affect sensing: Are affect dictionaries necessary to analyze affect?. In *International Conference on Affective Computing and Intelligent Interaction* (pp. 230-241). Springer Berlin Heidelberg.
- Oshiro, T. M., Perez, P. S., & Baranauskas, J. A. (2012, July). How many trees in a random forest?. In *International Workshop on Machine Learning and Data Mining in Pattern Recognition* (pp. 154-168). Springer Berlin Heidelberg.
- Pantic, M., & Rothkrantz, L. J. M. (2000). Automatic analysis of facial expressions: The state of the art. *IEEE Transactions on pattern analysis and machine intelligence*, 22(12), 1424-1445.
- Parrott, W. G. (2001). *Emotions in social psychology: Essential readings*. Psychology Press.
- Perikos, I., Angelopoulos, P., Paraskevas, M., Zarouchas, T., & Tzimas, G. (2014, December). A Methodology for Evaluation and Knowledge Extraction from On-going Learning Process. In

- Computational Science and Engineering (CSE)*, 2014 IEEE 17th International Conference on (pp. 425-431). IEEE.
- Perikos, I., & Hatzilygeroudis, I. (2013). Recognizing emotion presence in natural language sentences. In *International Conference on Engineering Applications of Neural Networks* (pp. 30-39). Springer Berlin Heidelberg.
- Perikos, I., & Hatzilygeroudis, I. (2015). Embodied conversational agents: A methodology for learning to express facial emotions. In *Information, Intelligence, Systems and Applications (IISA)*, 2015 6th International Conference on (pp. 1-5). IEEE.
- Perikos, I., & Hatzilygeroudis, I. (2016, July). A methodology for generating natural language paraphrases. In *Information, Intelligence, Systems & Applications (IISA)*, 2016 7th International Conference on (pp. 1-5). IEEE
- Perikos, I., & Hatzilygeroudis, I. (2016). A Case-Based Reasoning Approach to Convert Natural Language into First Order Logic. *IEEE International Conference on Computational Science and Engineering (CSE 2016)*, France, 2016
- Perikos, I., & Hatzilygeroudis, I. (2016). A Classifier Ensemble Approach to Detect Emotions Polarity in Social Media. In *Proceedings of the 12th International Conference on Web Information Systems and Technologies* ISBN 978-989-758-186-1, pages 363-370. DOI: 10.5220/0005864503630370
- Perikos, I., & Hatzilygeroudis, I. (2016). Recognizing emotions in text using ensemble of classifiers. *Engineering Applications of Artificial Intelligence*, 51, 191-201.
- Perikos, I., Grivokostopoulou, F., & Hatzilygeroudis, I. (2012). Automatic marking of NL to FOL conversions. In *Proc. of 15th IASTED International Conference on Computers and Advanced Technology in Education (CATE)*, Napoli, Italy (pp. 227-233).
- Perikos, I., Grivokostopoulou, F., & Hatzilygeroudis, I. (2011a). Teaching assistant tools for NL to FOL Conversion. In *Proceedings of the IADIS International Conference e-Learning* (pp. 20-23).
- Perikos, I., Grivokostopoulou, F., Hatzilygeroudis, I., & Kivas, K. (2011b). Difficulty Estimator for Converting Natural Language into First Order Logic. In *Intelligent Decision Technologies* (pp. 135-144). Springer Berlin Heidelberg.
- Perikos, I., Grivokostopoulou, F., & Hatzilygeroudis, I. (2011c). Help Generation in a System for learning Natural Language to First Order Logic Conversion. In *Proceedings of the 14th IASTED International Conference on Computers and Advanced Technology in Education (CATE 2011)*.
- Perikos, I., Grivokostopoulou, F., Hatzilygeroudis, I. (2016) Assistance and Feedback Mechanism in an Intelligent Tutoring System for Teaching Conversion of Natural Language into Logic, *International Journal of Artificial Intelligence in Education*, 2016 (To Appear)
- Perikos, I., Grivokostopoulou, F., Kivas, K., & Hatzilygeroudis, I. (2015). Utilizing Social Networks and Web 2.0 Technologies in Education. In *2015 IEEE Seventh International Conference on Technology for Education (T4E)* (pp. 118-119). IEEE.
- Perikos, I., Grivokostopoulou, F., Kivas, K., & Hatzilygeroudis, I. (2015). Assisting Tutors to Utilize Web 2.0 Tools in Education. *International Association for Development of the Information Society*.

- Perikos, I., Grivokostopoulou, F., Kovas, K., & Hatzilygeroudis, I. (2016). Automatic estimation of exercises' difficulty levels in a tutoring system for teaching the conversion of natural language into first-order logic. *Expert Systems*, 33(6), 569-580. doi: 10.1111/exsy.12182
- Perikos, I., Ziakopoulos, E., & Hatzilygeroudis, I. (2014). Recognizing emotions from facial expressions using neural network. In *IFIP International Conference on Artificial Intelligence Applications and Innovations* (pp. 236-245). Springer Berlin Heidelberg.
- Perikos, I., Ziakopoulos, E., & Hatzilygeroudis, I. (2015). Recognize Emotions from Facial Expressions Using a SVM and Neural Network Schema. In *Engineering Applications of Neural Networks* (pp. 265-274). Springer International Publishing.
- Phucharasupa, K., & Netisopakul, P. (2012). Thai Sentence Paraphrasing from the Lexical Resource.
- Picard, R. (1997). *Affective Computing* The MIT Press. Cambridge, MA, USA.
- Plutchik, R. (2001). The Nature of Emotions Human emotions have deep evolutionary roots, a fact that may explain their complexity and provide tools for clinical practice. *American Scientist*, 89(4), 344-350.
- Prinosil, J., Smékal, Z., & Esposito, A. (2008). Combining features for recognizing emotional facial expressions in static images. In *Verbal and Nonverbal Features of Human-Human and Human-Machine Interaction* (pp. 56-69). Springer Berlin Heidelberg.
- Psarologou, A., Bourbakis, N., & Virvou, M. (2014). A mapping mechanism of NL sentences onto an SPN state machine for understanding purposes. In *Information, Intelligence, Systems and Applications, IISA 2014, The 5th International Conference on* (pp. 321-324). IEEE.
- Psarologou, A., Esposito, A., & Bourbakis, N. (2015). A Synthesis of Stochastic Petri Net (SPN) Graphs for Natural Language Understanding (NLU) Event/Action Association. In *Tools with Artificial Intelligence (ICTAI), 2015 IEEE 27th International Conference on* (pp. 218-225). IEEE.
- Ptaszynski, M., Dokoshi, H., Oyama, S., Rzepka, R., Kurihara, M., Araki, K., & Momouchi, Y. (2013). Affect analysis in context of characters in narratives. *Expert Systems with Applications*, 40(1), 168-176.
- Ptaszynski, M., Dokoshi, H., Oyama, S., Rzepka, R., Kurihara, M., Araki, K., & Momouchi, Y. (2013). Affect analysis in context of characters in narratives. *Expert Systems with Applications*, 40(1), 168-176.
- Qiu, L., Lin, H., Ramsay, J., & Yang, F. (2012). You are what you tweet: Personality expression and perception on Twitter. *Journal of Research in Personality*, 46(6), 710-718.
- Quirk, C., Brockett, C., & Dolan, W. B. (2004). Monolingual Machine Translation for Paraphrase Generation. In *EMNLP* (pp. 142-149)
- Reisenzein, R., Hudlicka, E., Dastani, M., Gratch, J., Hindriks, K., Lorini, E., & Meyer, J. J. C. (2013). Computational modeling of emotion: Toward improving the inter-and intradisciplinary exchange. *IEEE Transactions on Affective Computing*, 4(3), 246-266.
- Roll, I., Aleven, V., McLaren, B. M., & Koedinger, K. R. (2011). Improving students' help-seeking skills using metacognitive feedback in an intelligent tutoring system. *Learning and Instruction*, 21(2), 267-280.

- Roll, I., Baker, R. S. D., Aleven, V., & Koedinger, K. R. (2014). On the benefits of seeking (and avoiding) help in online problem-solving environments. *Journal of the Learning Sciences*, 23(4), 537-560.
- Russell, J. A., (1980). A circumplex model of affect. *Journal of personality and social psychology*, 39(6), 1161.
- Sarsar, F., & Kisla, T. (2016). Emotional Presence in Online Learning Scale: A Scale Development Study. *Turkish Online Journal of Distance Education*.
- Scherer, K. R., & Wallbott, H. G. (1994). Evidence for universality and cultural variation of differential emotion response patterning. *Journal of personality and social psychology*, 66(2), 310.
- Schmid, H. (2013). Probabilistic part-of speech tagging using decision trees. In *New methods in language processing* (p. 154). Routledge.
- Shaheen, S., El-Hajj, W., Hajj, H., & Elbassuoni, S. (2014). Emotion recognition from text based on automatically generated rules. In *2014 IEEE International Conference on Data Mining Workshop* (pp. 383-392). IEEE.
- Shapiro, S. C. (1992). *Encyclopedia of artificial intelligence* second edition. New Jersey: A Wiley Interscience Publication.
- Shen, L., Wang, M., & Shen, R. (2009). Affective e-Learning: Using "Emotional" Data to Improve Learning in Pervasive Learning Environment. *Educational Technology & Society*, 12(2), 176-189.
- Shih, F. Y., Chuang, C. F., & Wang, P. S. (2008). Performance comparisons of facial expression recognition in JAFFE database. *International Journal of Pattern Recognition and Artificial Intelligence*, 22(03), 445-459.
- Shute, V. J., & Zapata-Rivera, D. (2010). Educational measurement and intelligent systems. *The International Encyclopedia of Education*. Oxford, UK: Elsevier Publishers.
- Shute, V.J. (2008) Focus on formative feedback, *Review of Educational Research*, Vol. 78, No. 1, pp.153–189.
- Sobel, I. (1990). An isotropic 3× 3 image gradient operator. *Machine Vision for three-dimensional Sciences*.
- Sokolova, M., & Lapalme, G. (2009). A systematic analysis of performance measures for classification tasks. *Information Processing & Management*, 45(4), 427-437.
- Stone, P. J., Dunphy, D. C., & Smith, M. S. (1966). *The General Inquirer: A Computer Approach to Content Analysis*.
- Strapparava, C., & Mihalcea, R. (2007). Semeval-2007 task 14: Affective text. In *Proceedings of the 4th International Workshop on Semantic Evaluations* (pp. 70-74). Association for Computational Linguistics.
- Strapparava, C., & Valitutti, A. (2004). WordNet Affect: an Affective Extension of WordNet. In *LREC* (Vol. 4, pp. 1083-1086).
- Thai, L. H., Nguyen, N. D. T., & Hai, T. S. (2011). A facial expression classification system integrating canny, principal component analysis and artificial neural network. *arXiv preprint arXiv:1111.4052*.

- Tian, Y., Kanade, T., & Cohn, J. F. (2011). Facial expression recognition. In *Handbook of face recognition* (pp. 487-519). Springer London.
- Tsitsoulis, A., & Bourbakis, N. (2013). A first stage comparative survey on human activity recognition methodologies. *International Journal on Artificial Intelligence Tools*, 22(06), 1350030.
- Valle, K. (2011). Graph-Based Representations for Textual Case-Based Reasoning.
- Van der Kleij, F. M., Eggen, T. J., Timmers, C. F., & Veldkamp, B. P. (2012). Effects of feedback in a computer-based assessment for learning. *Computers & Education*, 58(1), 263-272.
- VanLehn, K. (2006). The behavior of tutoring systems. *International journal of artificial intelligence in education*, 16(3), 227-265.
- VanLehn, K. (2011). The relative effectiveness of human tutoring, intelligent tutoring systems, and other tutoring systems. *Educational Psychologist*, 46(4), 197-221.
- VanLehn, K., Lynch, C., Schulze, K., Shapiro, J. A., Shelby, R., Taylor, L., & Wintersgill, M. (2005). The Andes physics tutoring system: Lessons learned. *International Journal of Artificial Intelligence in Education*, 15(3), 147-204.
- Vasilyeva, E., Pechenizkiy, M., & De Bra, P. (2008). Adaptation of elaborated feedback in e-learning. In *Adaptive hypermedia and adaptive web-based systems* (pp. 235-244). Springer Berlin Heidelberg.
- Verma, A., & Sharma, L. K. (2013). A comprehensive survey on human facial expression detection. *International Journal of Image Processing (IJIP)*, 7(2), 171.
- Vila, M., Martí, M. A., & Rodríguez, H. (2014). Is this a paraphrase? What kind? Paraphrase boundaries and typology. *Open Journal of Modern Linguistics*, 4(01), 205.
- Viera, A. J., & Garrett, J. M. (2005). Understanding interobserver agreement: the kappa statistic. *Fam Med*, 37(5), 360-363.
- Vinodhini, G., & Chandrasekaran, R. M. (2012). Sentiment analysis and opinion mining: a survey. *International Journal*, 2(6).
- Viola, P., & Jones, M. J. (2004). Robust real-time face detection. *International journal of computer vision*, 57(2), 137-154.
- Visutsak, P. (2012). Emotion classification using adaptive svms. *International Journal of Computer and Communication Engineering*, 1(3), 279.
- Visutsak, P. (2013). Emotion Classification through Lower Facial Expressions using Adaptive Support Vector Machines. *Journal of Man, Machine and Technology*, 2(1), 12-20.
- Wager, W., & Wager, S. (1985). Presenting questions, processing responses, and providing feedback in CAI. *Journal of Instructional Development*, 8(4), 2-8.
- Wang, G., Sun, J., Ma, J., Xu, K., & Gu, J. (2014). Sentiment classification: The contribution of ensemble learning. *Decision support systems*, 57, 77-93.
- Watson, I. (1999). Case-based reasoning is a methodology not a technology. *Knowledge-based systems*, 12(5), 303-308.

-
- Weston, J., Bordes, A., Chopra, S., & Mikolov, T. (2015). Towards ai-complete question answering: A set of prerequisite toy tasks. arXiv preprint arXiv:1502.05698.
- White, M., Steedman, M., Baldridge, J., & Kruijff, G. J. (2008). OpenCCG: The OpenNLP CCG Library.
- Wong, Y. W., & Mooney, R. J. (2007). Learning synchronous grammars for semantic parsing with lambdacalculus. In *Annual Meeting-Association for computational Linguistics* (Vol. 45, No. 1, p. 960).
- Worthington, T., & Wu, H. (2015, July). Time-shifted learning: Merging synchronous and asynchronous techniques for e-learning. In *Computer Science & Education (ICCSE), 2015 10th International Conference on* (pp. 434-437). IEEE.
- Xia, R., Zong, C., & Li, S. (2011). Ensemble of feature sets and classification algorithms for sentiment classification. *Information Sciences*, 181(6), 1138-1152.
- Yampolskiy, R. V. (2013). Turing test as a defining feature of AI-completeness. In *Artificial intelligence, evolutionary computing and metaheuristics* (pp. 3-17). Springer Berlin Heidelberg.
- Yampolskiy, R. V. (2012, April). AI-Complete, AI-Hard, or AI-Easy-Classification of Problems in AI. In *MAICS* (pp. 94-101).
- Yao, X., Van Durme, B., Callison-Burch, C., & Clark, P. (2013). Answer Extraction as Sequence Tagging with Tree Edit Distance. In *HLT-NAACL* (pp. 858-867).
- Zavaschi, T. H., Koerich, A. L., & Oliveira, L. E. S. (2011, May). Facial expression recognition using ensemble of classifiers. In *2011 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 1489-1492). IEEE.
- Zeng, Z., Pantic, M., Roisman, G. I., & Huang, T. S. (2009). A survey of affect recognition methods: Audio, visual, and spontaneous expressions. *IEEE transactions on pattern analysis and machine intelligence*, 31(1), 39-58.
- Zettlemoyer, L. S., & Collins, M. (2012). Learning to map sentences to logical form: Structured classification with probabilistic categorial grammars. arXiv preprint arXiv:1207.1420.
- Zhang, Z., Lyons, M., Schuster, M., & Akamatsu, S. (1998, April). Comparison between geometry-based and gabor-wavelets-based facial expression recognition using multi-layer perceptron. In *Automatic Face and Gesture Recognition, 1998. Proceedings. Third IEEE International Conference on* (pp. 454-459). IEEE.
- Zhao, S., Lan, X., Liu, T., & Li, S. (2009). Application-driven statistical paraphrase generation. In *Proceedings of the Joint Conference of the 47th Annual Meeting of the ACL and the 4th International Joint Conference on Natural Language Processing of the AFNLP: Volume 2-Volume 2* (pp. 834-842). Association for Computational Linguistics
- Zhao, S., Wang, H., Lan, X., & Liu, T. (2010). Leveraging multiple MT engines for paraphrase generation. In *Proceedings of the 23rd International Conference on Computational Linguistics* (pp. 1326-1334). Association for Computational Linguistics.
- Zhou, Y., Freedman, R., Glass, M., Michael, J. A., Rovick, A. A., & Evens, M. W. (1999). Delivering Hints in a Dialogue-Based Intelligent Tutoring System. In *AAAI/IAAI* (pp. 128-134).
-