

ROUGH SET BASED TECHNIQUES FOR EXTRACTIVE TEXT SUMMARIZATION

NIDHIKA YADAV



DEPARTMENT OF MATHEMATICS
INDIAN INSTITUTE OF TECHNOLOGY DELHI
NOVEMBER 2020

ROUGH SET BASED TECHNIQUES FOR EXTRACTIVE TEXT SUMMARIZATION

by

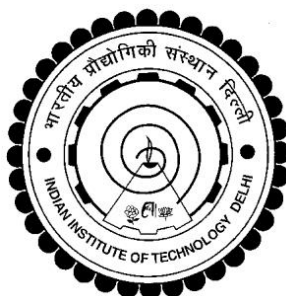
NIDHIKA YADAV

Department of Mathematics

Submitted

in fulfillment of the requirements of the degree of Doctor of Philosophy

to the



INDIAN INSTITUTE OF TECHNOLOGY DELHI
NOVEMBER 2020

CERTIFICATE

This is to certify that the thesis entitled ***Rough Set based techniques for Extractive Text Summarization*** submitted by ***Ms. Nidhika Yadav*** to the Indian Institute of Technology Delhi, for the award of the Degree of the **Doctor of Philosophy**, is a record of the original bona fide research work carried out by her under my supervision and guidance. The thesis has reached the standards fulfilling the requirements of the regulations relating to the degree.

The results contained in this thesis have not been submitted in part or full to any other university or institute for the award of any degree or diploma.

Place:

Professor Niladri Chatterjee

Date:

Department of Mathematics

Indian Institute of Technology Delhi

ACKNOWLEDGEMENTS

I am deeply thankful to Almighty for his immense and divine grace without which I would had been unable to carry out this research work. This thesis would not have been possible without the support of many people. It is impossible to acknowledge all of them here, but there are few people who deserve a special mention.

I am forever indebted to my Ph.D. supervisor **Professor Niladri Chatterjee** for teaching me basics of the research. I am grateful for his consistent support in guiding me till PhD Thesis submission. He encouraged me in dealing with the research and personal hurdles that I faced during this time. He taught me research writing, research methodologies and most importantly he helped me out in dealing with failures in experiments that came several times during this duration.

I would like to thank IIT Delhi for providing me necessary facilities to carry out my research work. I thank Head of Department of Mathematics, DRC Chairperson and my SRC members for their help at various stages of my research work.

I would like to thank all my family members for their unwandering support in carrying out my goal for completing the Thesis. I would like to thank them for their constant motivation and complete support in all aspects during this time which has been my true motivation. I would also like to thank all the people in my life who helped me in several aspects at various stages of Ph.D. I would like to thank to all my elders for their blessings.

I am thankful to the anonymous examiners and reviewers for their valuable comments and suggestions for improvising the papers and the thesis.

Nidhika Yadav

ABSTRACT

Rough Set (Pawlak, 1982) theory is a well-developed mathematical methodology to deal with uncertainty in data. The idea of Rough Set is to partition a given data of entities represented as attribute-value pairs into three classes: *positive region*, *negative region* and *boundary region*. The positive and negative regions separate the objects that belong to a specific class or outside it, while the boundary region consists of entities for which such a clear decision cannot be made with certainty.

Rough Set-based approaches have been used in different application areas including feature selection, data classification, knowledge acquisition, data mining, decision analysis among others. However, no significant work has been found for its application in the area of Extractive Text Summarization. Text Summarization, in general, condenses one or more collection of source text(s) into a single concise text that preserves the information content while reducing the size of the input text. In our approach we consider Extractive Text Summarization as a decision-making process, where corresponding to each input sentence the system takes a decision whether or not the sentence will be included in the summary. From this perspective we consider the process of sentence selection for generation of extractive summary as a process of decision making under uncertainty. The uncertainty comes from several aspects: the user-specified size of the summary, the relative importance of the sentences, presence of redundant information, the size and content of the document under consideration, among others.

We have proposed four different techniques for Rough Sets based Text Summarization. The schemes developed in the thesis are the following:

- (1) Novel concepts of Rough Set based *span* and *spanning sets* are proposed for generating a smaller representative representation of the Universe under consideration. Properties and computation of spanning sets are described. Further,

the unsupervised concept of spanning set is applied to generate extractive summary with a required number of important sentences, typically specified by the user, from the text.

- (2) Sentence to sentence similarity measures based on Fuzzy Rough Set are proposed and are further applied to unsupervised Single document Text Summarization. This new measure when applied to the problem of Text Summarization has been found to perform even better than the span-based approach with respect to certain sets of documents.
- (3) Suitability of learning from existing data (e.g. Document Understanding Conference 2002¹) using supervised Rough Set based techniques for Multi-document Text Summarization (MDTS) is analysed. Firstly, a Rough Set based ranking measure called *Rank-Measure* is proposed for post-processing the results obtained by supervised MDTS. The results show a significant improvement over results computed without using Rough Set based ranking measure. Further, a Neighborhood Rough Set based supervised learning algorithm is also proposed for MDTS and is found to improve the ROUGE scores further.
- (4) Rough Sets have been popularly used for feature selection. Selection of important features for Text Summarization using Rough Set based techniques is analysed. Further, two new algorithms for Rough Set based feature selection are proposed and their effect on performance for MDTS is analysed.
- (5) Finally, two more new techniques are developed during this research which are novel efforts in different areas to improve on the summarization results. The techniques are sentiment-based summarization and Random Indexing for MDTS. The suitability of these techniques is analysed for Text Summarization.

1. URL: <http://www-nlpir.nist.gov/projects/duc/guidelines/2001.html>

सार

रफ सेट (पावलक, 1982) डेटा में अनिश्चितता से निपटने के लिए एक अच्छी तरह से विकसित गणितीय पद्धति है। रफ सेट का विचार किसी दिए गए डेटा सेट को तीन वर्गों में विभाजित करना है: पॉजिटिव रीजन, नेगेटिव रीजन और बाउंड्री रीजन। पॉजिटिव रीजन और नेगेटिव रीजन उन डेटा पॉइंट्स को अलग करते हैं जो एक विशिष्ट वर्ग या उसके बाहर हैं, जबकि सीमा क्षेत्र में ऐसी संस्थाएं हैं जिनके लिए इस तरह का स्पष्ट निर्णय निश्चितता के साथ नहीं किया जा सकता है।

रफ सेट-आधारित दृष्टिकोणों का उपयोग विभिन्न एप्लिकेशन क्षेत्रों में किया गया है, जिसमें फीचर चयन, डेटा वर्गीकरण, ज्ञान अर्जन, डेटा माइनिंग, दूसरों के बीच निर्णय विश्लेषण शामिल हैं। हालांकि, एक्स्ट्रेक्टिव टेक्स्ट समराइजेशन के क्षेत्र में इसके आवेदन के लिए कोई महत्वपूर्ण काम नहीं मिला है। पाठ संक्षेपण, सामान्य रूप से, स्रोत पाठ के एक या अधिक संग्रह को एक संक्षिप्त पाठ में सम्मिलित करता है, जो इनपुट पाठ के आकार को कम करते हुए सूचना सामग्री को संरक्षित करता है। हमारे दृष्टिकोण में हम एक्स्ट्रेक्टिव टेक्स्ट समराइजेशन को निर्णय लेने की प्रक्रिया मानते हैं, जहाँ प्रत्येक इनपुट वाक्य के अनुसार सिस्टम निर्णय लेता है कि वाक्य को सारांश में शामिल किया जाएगा या नहीं। इस दृष्टिकोण से हम अर्क सारांश के निर्माण के लिए वाक्य चयन की प्रक्रिया को अनिश्चितता के तहत निर्णय लेने की प्रक्रिया मानते हैं। अनिश्चितता कई पहलुओं से आती है: सारांश का उपयोगकर्ता-निर्दिष्ट आकार, वाक्यों का सापेक्ष महत्व, निरर्थक जानकारी की उपस्थिति, विचाराधीन दस्तावेज़ का आकार और सामग्री, अन्य।

हमने रफ सेट्स आधारित टेक्स्ट सारांश के लिए चार अलग-अलग तकनीकों का प्रस्ताव दिया है। थ्रीसिस में विकसित की गई योजनाएं निम्नलिखित हैं:

(१) रफ सेट आधारित स्पैन और स्पैनिंग सेट्स की नोवेल अवधारणाएं विचाराधीन यूनिवर्स के एक छोटे प्रतिनिधि प्रतिनिधित्व को बनाने के लिए प्रस्तावित हैं। स्पैनिंग सेट्स के गणितीय गुण और गणना का वर्णन किया गया है। इसके अलावा, स्पैनिंग सेट्स की अवधारणा, आमतौर पर उपयोगकर्ता द्वारा निर्दिष्ट महत्वपूर्ण वाक्यों की एक आवश्यक संख्या के साथ निष्कर्षण सारांश उत्पन्न करने के लिए लागू किया जाता है।

(2) फजी रफ़ सेट के आधार पर वाक्या का वाक्या से मेल का प्रस्ताव किया गया हैं और आगे इसे एकल दस्तावेज़ पाठ संक्षेप के लिए लागू किया गया है।

यह नया उपाय जब टेक्स्ट समराइजेशन की समस्या पर लागू किया गया है, तो दस्तावेजों के कुछ सेटों के संबंध में स्पैन-आधारित दृष्टिकोण से भी बेहतर प्रदर्शन करने के लिए पाया गया है।

(3) मौजूदा डेटा (जैसे डॉक्यूमेंट अंडरस्टैंडिंग कॉन्फ्रेंस 2002) से सीखने की उपयुक्तता, मल्टी-डॉक्यूमेंट टेक्स्ट सुमरीकरण (MDTS) के लिए पर्यवेक्षित रफ़ सेट आधारित तकनीकों का उपयोग करके विश्लेषण किया गया हैं । सबसे पहले, एक रफ़ सेट आधारित रैंकिंग माप, पर्यवेक्षण MDTS द्वारा प्राप्त परिणामों के बाद के प्रसंस्करण के लिए प्रस्तावित है। परिणाम रफ़ सेट आधारित रैंकिंग उपाय का उपयोग किए बिना गणना किए गए परिणामों पर एक महत्वपूर्ण सुधार दिखाते हैं। इसके अलावा, एक नेइबोर्हूद रफ़ सेट आधारित सुपरवाइज्ड लर्निंग अल्गोरिथम भी MDTS के लिए प्रस्तावित किया गया हैं और ROUGE स्कोर को और बहुत बेहतर पाया जाता है।

(4) फीचर चयन के लिए रफ़ सेट्स का लोकप्रिय रूप से उपयोग किया गया है। रफ़ सेट आधारित तकनीकों का उपयोग करके पाठ संक्षेपण के लिए महत्वपूर्ण विशेषताओं का चयन किया जाता है। इसके अलावा, रफ़ सेट आधारित सुविधा चयन के लिए दो नए एल्गोरिदम प्रस्तावित हैं और एमडीटीएस के प्रदर्शन पर उनके प्रभाव का विश्लेषण किया गया है।

Table of Contents

CERTIFICATE	i
ACKNOWLEDGEMENTS	iii
ABSTRACT	v
Table of Contents	vii
List of Figures	xiv
List of Tables	xv
Chapter 1	1
Introduction	1
1.1 Rough Set	2
1.1.1 Motivation	3
1.1.2 Objectives of the Thesis	4
1.1.3 Organisation of the Thesis	4
Chapter 2	9
A Survey of Related Past Works	9
2.1 Text Summarization	9

2.1.1	Latent Semantic Analysis	9
2.1.2	Non-Negative Factorization	10
2.1.3	Graph Based Models for Text Summarization	11
2.1.4	Random Indexing based Text Summarization	13
2.1.5	WordNet based Summarization	14
2.1.6	Optimization Based Models for Text Summarization	15
2.1.7	Machine Learning Based Models for Text Summarization	18
2.1.8	Centroid Method for Text Summarization	20
2.1.9	Other Techniques for Text Summarization	21
2.2	Summary Evaluation	22
2.3	Rough Set.....	24
2.5.1	Rough Sets Theory	26
2.5.2	Variants of Rough Set Theory.....	29
2.5.3	Applications of Rough Sets	33
Chapter 3		37
Rough Sets based Span and its Application to Extractive Text Summarization		37
3.1	Introduction	37
3.2	Span and Spanning Sets	39

3.3	Computing Spanning Sets	46
3.4	Span based Text Summarization	48
3.4.1	Determining Extract of Text using Span	48
3.4.2	Text Pre-processing	48
3.4.3	Optimization Problem.....	50
3.5	Proposed Solution for Computing Spanning Sets of Text Documents	50
3.6	Results and Observations	51
3.6.1	Experiments	51
3.7	Concluding Remarks.....	56
Chapter 4.....		58
Application of Fuzzy Rough Set Based Sentence Similarity Measure in Text Summarization		58
4.1	Fuzzy Rough Set.....	58
4.2	Fuzzy Rough Sets Theory	61
4.3	Proposed Fuzzy Rough Set based Sentence Similarity Measure	63
4.4	Proposed Text Summarization Technique	65
4.5	Results & Discussions.....	68
4.6	Concluding Remarks.....	73

Chapter 5	75
Rough Set based Supervised Techniques for Multi-Document Summarization	75
5.1 Introduction	75
5.2 Sentiment as a Feature in Summarization.....	77
5.2.1 Sentiment Analysis	78
5.2.2 Proposed Summarization Model.....	81
5.2.3 Results & Analysis	82
5.3 Features for Supervised Summarization	83
5.4 Supervised Learning Techniques	85
5.4.1 Support Vector Machine	85
5.4.2 Neural Networks	87
5.4.3 Random Forest	88
5.4.4 <i>K</i> -Nearest Neighbor (KNN).....	90
5.4.5 Fuzzy Nearest Neighbour	91
5.4.6 Logistic Regression	92
5.4.7 Naïve Bayes.....	92
5.5 Rough Set based Supervised Techniques	93
ii) Fuzzy Rough Nearest Neighbour Algorithm	94

5.6	Comparison of Rough Set based methods with other Supervised Techniques for MDTs	95
5.7	Proposed Rough Set based Rank-Measure	96
5.7.1	Motivation.....	97
5.7.2	Proposed Rank-Measure.....	98
5.7.3	Proposed Aggregate-Rank-Measure	99
5.7.4	Properties of Rough Set based Rank-Measure and Aggregate-Rank-Measure.....	101
5.7.5	Results for Rough Set based Rank-Measure based Multi-document Text Summarization	104
5.8	The Proposed Neighborhood LER Algorithm	108
5.9	Concluding Experiments & Results of Proposed Techniques	110
5.10	Conclusion.....	114
Chapter 6.....		116
Rough Set based Feature Selection & Its Application to Text Summarization.....		116
6.1	Introduction	116
6.2	Rough Set based Feature Selection	118
6.2.1	Quick Reduct Algorithm.....	118
6.3	Rough-Cluster Reduct	120

6.4	Best-First Reduct	125
6.5	Experiments and Results	128
6.5.1	Conclusion.....	130
Chapter 7.....		131
Conclusion of Thesis		131
7.1	Summary of the Research Work.....	132
7.1.1	Single-Document Summarization	132
7.1.2	Multi-document Text Summarization (MDTS)	136
7.2	Future Research Directions	137
7.3	Concluding Remarks.....	139
References.....		141
Appendix - I.....		156
	Random Indexing and Centroid based Summarization	156
	Proposed Algorithm.....	156
	Concluding Remarks	159
Appendix - II.....		160
	Rough-Cluster based Feature Selection: Other Datasets	160
	Results of Rough-Cluster on other Machine Learning Datasets.....	161

Publications.....	166
About the Author	168

List of Figures

Fig. 2.1. Positive, Negative and Boundary Regions	25
Fig. 3.1. Pre-processed text example	49
Fig. 4.1. Algorithm to compute extract using Fuzzy Rough Set.....	66
Fig. 5.1. Two Class Classification by Hyperplane.....	86
Fig. 5.2. Architecture of a Neural Network based training algorithm.	88
Fig. 5.3. Illustration of workflow of Decision Tree.....	89
Fig. 5.4. Algorithm for K-Nearest Neighbor	90
Fig. 5.5. Algorithm for Fuzzy K-Nearest Neighbor.....	91
Fig. 5.6. Algorithm to compute Rank-Measure and Extract.....	104
Fig 5.7. Flowchart of overall flow of the proposed Supervised MDTS	105
Fig. 5.8. Algorithm for NLER based MDTS	109
Fig. 6.1. Quick Reduct Algorithm	125
Fig. 6.2. Rough-Cluster Reduct Algorithm.....	128
Fig. 6.3. Example Information System	129
Fig. 6.4. Proposed Best First-Reduct	133
Fig. 6.5. Proposed A* Reduct	134

List of Tables

Table 2.1: Knowledge Representation System for Example 2.2	28
Table 3.1. Example on span and spanning set computation	46
Table 3.2. Minimal Complete Spanning set for Example 1	47
Table 3.3. Description of datasets.....	51
Table 3.5. Average of Recall Results for Enron Personal Email Data	55
Table 3.4. Average of Recall Results DUC2001	54
Table 3.7. Average of 109 ROUGE Results Enron Corporate Email Data	56
Table 3.6. Average of Recall Results DUC2002	Error! Bookmark not defined.
Table 4.1: Fuzzy Rough Set based lower & upper approximations for Sentence 1	67
Table. 4.2. Fuzzy Rough Set based lower & upper approximations for Sentence 2	68
Table 4.3. MSE and correlation score of various approaches of sentence similarity on SICK2014 dataset	70
Table 4.4: ROUGE scores for summarization using proposed and standard summarization techniques for 50%, 25% and 10% compression degrees.....	73
Table 5.1: Various Sentiments associated with word “unable”	79
Table 5.2: Various Sentiments associated with different POS tags.....	80
Table 5.3: Example of sentiment scores in a sentence.....	81
Table 5.4. Results for sentiment-based summarization	82
Table 5.5. Results of DUC2003 dataset for supervised techniques and Rough Set based techniques.	95
Table 5.6. Results of DUC2005 dataset for supervised techniques and Rough Set based techniques.	96

Table 5.7. Information System for Example 1	99
Table 5.8. Rank-Measure and Aggregate Rank-Measure of Example 1	100
Table 5.9. Results of DUC2003 for various supervised learning algorithms with Aggregate-Rank-Measure	105
Table 5.10. Percent increase in DUC2003 results for various supervised learning algorithms with Aggregate-Rank-Measure	106
Table 5.11. Results of DUC2005 for various supervised learning algorithms with Aggregate-Rank-Measure	107
Table 5.12. Percent increase in DUC2005 results for various supervised learning algorithms with aggregate Rank-Measure.	107
Table 5.13. Results of DUC2003 for various supervised learning algorithms and Neighborhood Rough Sets without Aggregate-Rank-Measure	111
Table 5.14. Results of DUC2003 for various supervised learning algorithms and Neighborhood Rough Sets with classical membership functions.....	111
Table 5.16. Results of DUC2003 for various supervised learning algorithms and Neighborhood Rough Sets with Aggregate-Rank-Measure	112
Table 5.15. Results of DUC2005 for various supervised learning algorithms and Neighborhood Rough Sets without Aggregate-Rank-Measure	112
Table 5.18. Results of DUC2005 for various supervised learning algorithms with Aggregate-Rank-Measure with Neighborhood Rough Sets	113
Table 5.17. Results of DUC2005 for various supervised learning algorithms and.....	113
Neighborhood Rough Sets with classical membership.....	113
Table. 6.1. ROUGE Scores of DUC2002 dataset	129
Table 7.1. Average of 90 DUC 2002 Dataset	134
Table 7.2. Average of ROUGE Scores DUC2001	134
Table 7.3. Average ROUGE Scores of Enron Personal Email Dataset	135
Table 7.4. Average ROUGE Scores of Enron Corporate Email Dataset.....	135

Table A1.1: ROUGE scores for MDTS with the proposed models.....	159
Table A2.1: Properties of datasets used.....	162
Table A3.3: Results for 5-folds with KNN classifier.	163
Table A2.2: Results for 5-folds with SVM classifier.	163
Table A2.4: Average number of features selected.....	164