



รายงานวิจัยฉบับสมบูรณ์

การค้นหาคำสำคัญโดยใช้ความหมาย (Semantic-Based Keyword Extraction)

โดย

รองศาสตราจารย์ ดร. โกสินทร์ จ่านงไทย

31 พฤษภาคม พ.ศ. 2551

รายงานวิจัยฉบับสมบูรณ์

โครงการการค้นหาคำสำคัญโดยใช้ความหมาย (Semantic-Based Keyword Extraction)

รองศาสตราจารย์ ดร. โกสินทร์ จันทน์ไทย

ภาควิชาวิศวกรรมอิเล็กทรอนิกส์และโทรคมนาคม
คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าธนบุรี

สนับสนุนโดยสำนักงานกองทุนสนับสนุนการวิจัย

(ความเห็นในรายงานนี้เป็นของผู้วิจัย สกว. ไม่จำเป็นต้องเห็นด้วยเสมอไป)

Project Code: RMU4880007
(รหัสโครงการ)

Project Title: โครงการการค้นหาคำสำคัญโดยใช้ความหมาย
(ชื่อโครงการ) (Semantic-Based Keyword Extraction)

Investigator: รองศาสตราจารย์ ดร. โกสินทร์ จำนงไทย
(ชื่อนักวิจัย) ภาควิชาวิศวกรรมอิเล็กทรอนิกส์และโทรคมนาคม
คณะวิศวกรรมศาสตร์
มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าธนบุรี

E-mail Address: kosin.cha@kmutt.ac.th, kosin.cha@gmail.com

Project Period: 3 ปี (พฤษภาคม 49 – พฤษภาคม 51)
(ระยะเวลาโครงการ)

In several keyword extraction systems based on semantic determination, the extracted keywords are selected from many candidates having similar meanings as ones in the knowledge base. To use only term similarity is limited because two candidates can be related without being similar. Keywords may not be considered only synonym but also antonym, hypernym and hyponym. Using term relatedness for considering candidates from a document can disclose more keywords than term similarity. In this research, we use three features as synonym (its similarity), hypernym (its generality) and hyponym (its particularity) formed as term relatedness for extracting keywords. Moreover, in previous researches, they used term relatedness level to extract a keyword. This judgment is limited to a single part-of-speech such as a pair of nouns or a pair of verbs. Since, each term has several senses of meaning; two related terms may be ignored if they are considered in different senses. Therefore, we need to know exactly which sense of term is considered. In this research, we use sentence-level relatedness determination instead of term relatedness to extract keywords. By using sentence-based relatedness, each sense of term can be known by its function in the sentence. Therefore, this research proposes a sentence relatedness measurement for determining keywords that can provide additional keywords. We experiment the proposed measurement to 100 technical papers from IEICE and ACM transactions in domain of computer science and engineering. The performance of the keyword extraction system is significantly improved.

Keywords— keyword extraction, synonym, hypernym, hyponym

สารบัญ

รายการ	หน้า
รายการรูปภาพ	v
รายการตาราง	vii
บทที่ 1 บทนำ	1
บทที่ 2 ทฤษฎีและงานวิจัยที่เกี่ยวข้อง	7
บทที่ 3 ระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย	18
บทที่ 4 ขั้นตอนการวิจัยและผลการวิจัย	33
บทที่ 5 การนำงานวิจัยไปใช้ประโยชน์	38
บรรณานุกรม	40
ภาคผนวก	43
[1] Chey, C., Kumhom, P. and Chamnongthai, K. (2006) Khmer Printed Character Recognition by using Wavelet Descriptors. International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems , 14(3), 337-350.	44
[2] Kongkachandra, R., Chamnongthai, K., and Kimpan, C., (2006), Intelligent Keyword Extraction for Digital Library, Information Technology and Libraries , Vol. 25, No. 2, (to be appeared).	59
[3] Abduction-Based Knowledge Revision for Key phrase Extraction System, Computational Intelligence , (to be submitted)	76

รายการรูปภาพ

รายการ	หน้า
รูปที่ 1.1 แสดงผลลัพธ์ของเครื่องมือค้นหา Google เมื่อใส่คำสอบถาม “Triangle” , “Bermuda” และ “Bermuda Triangle”	2
รูปที่ 1.2 กระบวนการในการ กำหนดคำสำคัญ (Keyword Assignment)	3
รูปที่ 1.3 แสดงถึงกระบวนการในการสกัดคำสำคัญ (Keyword Extraction)	5
รูปที่ 2.1 แสดงความสัมพันธ์ลักษณะต่าง ๆ ของ “car” กับคำอื่น ๆ	8
รูปที่ 2.2 ตัวอย่างการแสดงผลของเวิร์ดเน็ตจากการค้นหาคำว่า “match”	10
รูปที่ 2.3 แสดงตัวอย่าง Conceptual Graphs	11
รูปที่ 2.4 แสดงตัวอย่าง Conceptual Graphs	11
รูปที่ 2.5 แสดงระยะห่างระหว่าง “car” และ “boat” ในโครงข่ายเวิร์ดเน็ตเท่ากับ 6	12
รูปที่ 3.1 แผนภาพโดยรวมของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย	18
รูปที่ 3.2 ขั้นตอนการรู้จำตัวอักษร (Optical Character Recognition)	20
รูปที่ 3.3 แผนผังรวมของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย	22
รูปที่ 3.4 แสดงถึงการจัดเก็บข้อมูลภายในฐานความรู้ของคำสำคัญ	23
รูปที่ 3.5 การใช้กราฟเชิงความคิดแทนความหมายของประโยค “The dog scratches its ear with its paw”	24
รูปที่ 3.6 แสดงถึงการจัดเก็บข้อมูลภายในฐานความรู้คำสำคัญ	25
รูปที่ 3.7 หลักการของการดึงคำสำคัญเริ่มต้น (Initial Keyword Extraction)	26
รูปที่ 3.8 วิธีการสำหรับคัดเลือกนามวลี	27
รูปที่ 3.9 วิธีการในการแปลงจากประโยคภาษาอังกฤษ เป็น กราฟความหมาย	28
รูปที่ 3.10 กราฟความหมายของคำสำคัญ “UNIX” และ คำคู่แข่ง “LINUX”	29
รูปที่ 3.11 การวัดความเกี่ยวข้องทางความหมาย	29
รูปที่ 3.12 การเก็บค่าข้อมูลในแต่ละโนดของกราฟความหมาย	30
รูปที่ 3.13 แสดงถึงขั้นตอนการดึงคำสำคัญที่มีความเกี่ยวข้องทางความหมาย	31
รูปที่ 3.14 แสดงถึงขั้นตอนการขยายฐานความรู้คำสำคัญ	31
รูปที่ 3.15 ตัวอย่างของการปรับปรุงฐานความรู้คำสำคัญ	32
รูปที่ 4.1 หลักในการประเมินประสิทธิภาพของระบบ	33
รูปที่ 5.1 ระบบผู้เชี่ยวชาญสำหรับวินิจฉัยโรค	38

(Interactive Medical Diagnosis Expert System)	
รูปที่ 5.2 ระบบถาม-ตอบทางโทรศัพท์	38
(Intelligent Answering System for Tourist)	
รูปที่ 5.3 ระบบอัจฉริยะสำหรับจัดประเภทบทความวิชาการ	39
(Intelligent Technical Paper Classification)	

รายการตาราง

รายการ	หน้า
ตารางที่ 4.1 การกำหนดค่าต่างๆสำหรับงานวิจัย	34
ตารางที่ 4.2 ประสิทธิภาพของระบบที่นำเสนอเมื่อเทียบกับงานวิจัยเดิมที่มีอยู่	35
ตารางที่ 4.3 ตัวอย่างของผลที่ได้ของการวิจัยเมื่อเทียบกับงานวิจัยเดิมที่ใช้ คลังข้อมูล	36
ตารางที่ 4.4 ตัวอย่างของผลที่ได้ของการวิจัยเมื่อเทียบกับงานวิจัยเดิมที่ไม่ใช่ คลังข้อมูล	36
ตารางที่ 4.5 ประสิทธิภาพของการนำเอาฐานความรู้ของคำสำคัญมาใช้	37

บทที่ 1 บทนำ

1.1 ความสำคัญ

เมื่อ 20 กว่าปีที่แล้วถ้าเราต้องสืบค้นข้อมูลเกี่ยวกับเรื่องใดๆ เราจำเป็นต้องเดินทางไปยังห้องสมุดเพื่อค้นหาข้อมูลจากหนังสือบทความ นิตยสาร หรือหนังสือพิมพ์ต่างๆ อีกทั้งยังไม่สามารถรันตีได้ว่าเราจะเจอข้อมูลที่ต้องการจากการสืบค้นนั้นหรือไม่ ต่อมาอีก 10 ปี คอมพิวเตอร์ได้ถูกพัฒนาเพื่อใช้ในงานด้านการประมวลผลสารสนเทศมากขึ้น ตัวอย่างเช่น เครื่องมือค้นหา (Search Engine) ที่นิยมใช้กันในปัจจุบัน ได้แก่ Google, Yahoo, AltaVista หรือ InfoSeek ได้ช่วยให้ผู้ใช้สามารถสืบค้นข้อมูลได้สะดวกสบาย และรวดเร็วขึ้น

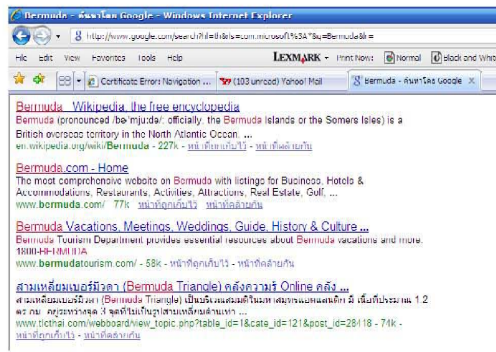
อย่างไรก็ตาม แม้ผู้ใช้จะสามารถสืบค้นข้อมูลได้รวดเร็วขึ้น แต่ก็ยังเจอกับปัญหาของข้อมูลที่ได้จากการสืบค้นมีมากเกินไป ผู้ใช้ต้องมาคัดแยกสารสนเทศที่ต้องการออกจากข้อมูลที่สืบค้นมาได้ อีกกรอบหนึ่งซึ่งอาจจะเสียเวลามากเช่นกัน นอกจากนี้ การค้นหาข้อมูลผู้ใช้จำเป็นต้องใส่ คำสอบถาม (query) ที่เหมาะสมเพื่อให้ตรงกับเนื้อหาของเอกสารที่ต้องการ รูปที่ 1.1 แสดงผลลัพธ์ของเครื่องมือค้นหา Google เมื่อใส่คำสอบถาม ดังนี้ “Triangle”, “Bermuda” และ “Bermuda Triangle” ในตัวอย่างถ้าผู้ใช้เจตนาจะค้นหาสารสนเทศเกี่ยวกับ “Bermuda Triangle” แต่ใส่คำสอบถามเพียง “Bermuda” หรือ “Triangle” ก็จะได้ข้อมูลอื่นๆที่ไม่ต้องการออกมาด้วย แต่ถ้าใส่ คำสอบถามเป็น “Bermuda Triangle” ก็จะได้สารสนเทศที่ตรงกับความต้องการ ปรากฏการณ์เช่นนี้สามารถบอกโดยนัยได้ว่า ถ้าเราใส่คำสอบถามที่เหมาะสม ก็จะได้รับข้อมูลที่ต้องการมากขึ้น เนื่องจากการทำงานโดยทั่วไปของเครื่องมือค้นหา คือ เปรียบเทียบ คำสอบถามกับคำสำคัญ (Keyword) ที่อยู่ในเอกสารซึ่งถ้าตรงกันก็จะแสดงเอกสารนั้นขึ้นมาให้กับผู้ใช้

เอกสารบางประเภท เช่น บทความวิชาการ หนังสือเรียน โดยส่วนใหญ่ผู้แต่งจะกำหนดคำสำคัญเอาไว้ในเอกสาร ทำให้เครื่องมือค้นหาสามารถเปรียบเทียบ คำสอบถาม กับ คำสำคัญเหล่านี้ได้ทันที แล้วได้ผลลัพธ์ ตามที่ต้องการ

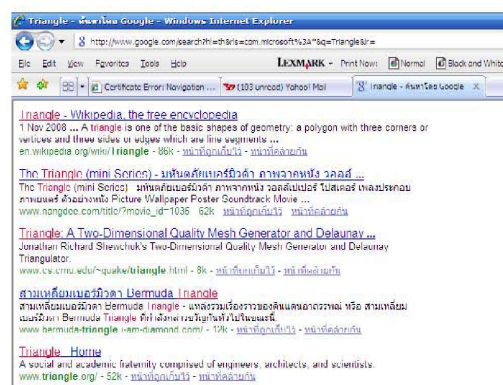
ความสำคัญของ คำสำคัญ (Keyword) สามารถสรุปได้ดังนี้

- 1) คำสำคัญ สามารถใช้เป็นบทสรุปของเนื้อหาเอกสารโดยรวมได้เราสามารถเดาเรื่องราวของเอกสารได้ด้วยการดูที่คำสำคัญ ซึ่งช่วยให้เข้าใจเอกสารนั้นได้เร็วขึ้น
- 2) คำสำคัญ สามารถใช้เป็น เครื่องบอกทาง ในการค้นหาข้อมูลที่เกี่ยวข้องอื่นๆได้
- 3) คำสำคัญ สามารถใช้ในการจัดแยกประเภทของเอกสารได้

Given: Bermuda



Given: Triangle



Given: Bermuda Triangle



รูปที่ 1.1 แสดงผลลัพธ์ของเครื่องมือค้นหา Google เมื่อใส่คำสอบถาม “Triangle”, “Bermuda” และ “Bermuda Triangle”

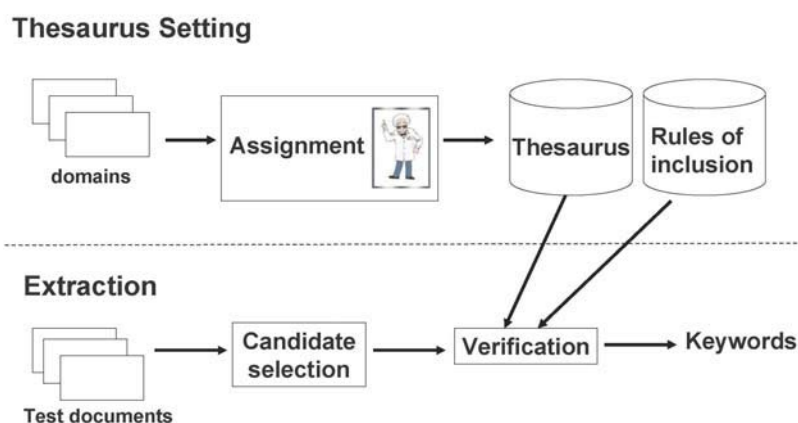
1.2 ที่มาของปัญหา

ถึงแม้ว่า คำสำคัญ จะเป็นที่ยอมรับว่าเป็นตัวแทนที่ดีของเอกสาร [1-11] แต่เอกสารส่วนใหญ่ ผู้เขียนไม่ได้กำหนด คำสำคัญ เอาไว้ จะมีก็เพียงเอกสารบางชนิด เช่น บทความหรือหนังสือวิชาการ เท่านั้น ดังนั้น การค้นหาคำสำคัญโดยอัตโนมัติ จึงเป็นสิ่งที่จำเป็นสำหรับการใช้งานด้านต่างๆ เช่น การค้นคืนเอกสาร การสรุปใจความสำคัญ การจัดประเภทเอกสาร และ การแปลเอกสาร เป็นต้น งานวิจัย [12] ได้กล่าวว่า เทคนิคที่ใช้ในการสร้างระบบค้นหาคำสำคัญแบบอัตโนมัติ สามารถแบ่งออกได้เป็น 2 กลุ่มคือ

1. การกำหนดคำสำคัญ (Keyword Assignment)
2. การสกัดคำสำคัญ (Keyword Extraction)

1.2.1 การกำหนดคำสำคัญ (Keyword Assignment)

เป็นกระบวนการในการดึงเอาคำสำคัญออกจากเอกสารแบบอัตโนมัติด้วยการจับคู่ ระหว่าง คำที่พิจารณา (Candidate) กับ คำสำคัญที่กำหนดเอาไว้ ซึ่งเรียกว่า คำศัพท์ควบคุม (Controlled Vocabulary) หรือ อรรถาภิธาน (Thesaurus) ซึ่งถูกกำหนดขึ้นมาจากการที่ผู้เชี่ยวชาญ อ่านเอกสาร ทั้งหมด ทำความเข้าใจ และ กำหนดคำศัพท์ควบคุมเหล่านี้ขึ้น นอกจากนี้ จะกำหนดกฎที่ใช้ในการ ผนวกคำศัพท์ควบคุมเข้ากับคำอื่นๆ (Rules of Inclusion) ด้วย ขั้นตอนนี้แสดงในรูปที่ 1.2 ในส่วน เตรียมคำศัพท์ควบคุม (Thesaurus Setting) จากนั้นเมื่อต้องการค้นหาคำสำคัญแบบอัตโนมัติ เอกสารที่ต้องการทดสอบจะถูกนำมาพิจารณาหา คำที่พิจารณา (Candidate) จากนั้นก็ทำการ เปรียบเทียบกับ คำศัพท์ควบคุมและกฎของการผนวกคำศัพท์ คำที่พิจารณา (Candidate) ใดเป็นไป ตามกฎที่ใช้ ก็จะนับว่าเป็น คำสำคัญ (Keyword) ของเอกสารนั้น ขั้นตอนนี้คือ ส่วนดึงคำสำคัญ (Extraction) ในรูปที่ 1.2



รูปที่ 1.2 กระบวนการในการ กำหนดคำสำคัญ (Keyword Assignment)

ตัวอย่างเช่น ผู้เชี่ยวชาญ กำหนดคำว่า “childbirth” ขึ้นมาเป็นคำศัพท์ควบคุม จากนั้นก็จะใช้กฎในการผนวกคำศัพท์ เชื่อมโยงไปยังคำอื่นที่เกี่ยวข้อง ได้แก่ “birth”, “labor”, “labour”, “delivery”, “forceps”, “baby” และ “born” เป็นต้น

เนื่องจาก เทคนิคการกำหนดคำสำคัญ ใช้ผู้เชี่ยวชาญในการกำหนดคำศัพท์ควบคุม เราจึง ถือได้ว่า คำสำคัญ ที่ได้เทคนิคนี้ เป็นตัวแทนทางความหมายที่ดีของเอกสาร แต่อย่างไร ก็ตามวิธีนี้มี ข้อจำกัด ดังนี้

- วิธีนี้ต้องอาศัยผู้เชี่ยวชาญในการกำหนดคำศัพท์ควบคุม ซึ่งหาได้ยากและต้องเสียค่าใช้จ่ายมาก
- วิธีนี้ ต้องใช้เวลาและความทุ่มเทอย่างมาก
- วิธีนี้ ใช้กฎในการผนวกคำศัพท์ที่เป็นแบบ Heuristic ซึ่งขึ้นอยู่กับผู้เชี่ยวชาญเฉพาะคน

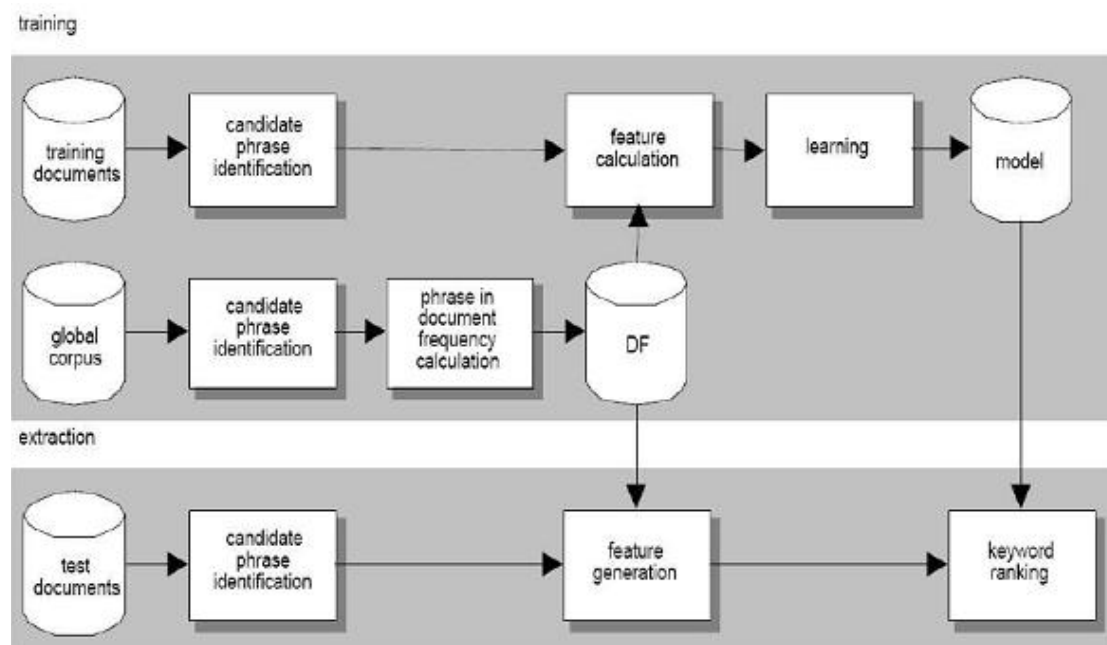
ได้มีงานวิจัยหลายงานที่พัฒนาขึ้นมาโดยใช้เทคนิคนี้ อาทิเช่น E.J. Schuegraf และ F. Van Bomme [14] เสนอวิธีการ โดยใช้ผู้เชี่ยวชาญหลายคนมาร่วมกันกำหนดกฎในการผนวกคำศัพท์ ซึ่งงานวิจัยนี้จะพิจารณาว่า คำใดที่มักเกิดขึ้นร่วมกัน (Co-occurrence) กับคำศัพท์ควบคุมบ่อยๆ เป็นคำที่สามารถนำมาผนวกเข้ากับคำศัพท์ควบคุมได้ งานวิจัยนี้สามารถลดเวลาในการค้นหาความสำคัญได้ แต่ยังคงต้องใช้ผู้เชี่ยวชาญอยู่ งานวิจัยอีกงานโดย C.H. Lueng และ W.K. Kan [15] เสนอที่จะไม่ใช้ผู้เชี่ยวชาญ แต่ใช้การฝึกฝนจากกลุ่มตัวอย่างที่ถูกต้อง (Positive example) และไม่ถูกต้อง (Negative example) ในการหาความสัมพันธ์ระหว่าง คำศัพท์ควบคุม กับ คำอื่นๆที่จะนำมาผนวกเข้ากับคำศัพท์ควบคุม ซึ่งความสัมพันธ์เหล่านี้คำนวณได้จากผลคูณของ TF x OSDF x CSIDF โดยที่ TF: Term Frequency, OSDF: Document Frequency in Original Set และ CSIDF: Inverse Document Frequency in Combined Set)

1.2.2 การสกัดคำสำคัญ (Keyword Extraction)

เทคนิค นี้เป็นอีกทางเลือกหนึ่งในการค้นหาคำสำคัญจากเอกสารโดยไม่ต้องใช้ผู้เชี่ยวชาญ วิธีนี้จะอาศัยหลักการของการฝึกฝนเพื่อให้รู้จำรูปแบบของคำสำคัญ เพื่อนำรูปแบบที่ได้นี้ ไปค้นหาคำสำคัญในเอกสารอื่นๆ ต่อไป ดังนั้น แทนที่จะใช้ผู้เชี่ยวชาญมาเป็นคนกำหนดรูปแบบของคำสำคัญ เราใช้ ตัวอย่างของเอกสาร (Training document) และ คลังข้อมูล (Corpus) จำนวนมากมาฝึกสอน เพื่อให้ระบบสามารถสร้างแบบจำลองของคำสำคัญ (Keyword Model) ได้ รูปที่ 1.3 แสดงถึงกระบวนการในการสกัดคำสำคัญ (Keyword Extraction) ซึ่งประกอบไปด้วย 2 ส่วนคือ ส่วนของการฝึกฝนเพื่อสร้างแบบจำลอง และ ส่วนของการสกัดคำสำคัญ

ในส่วนของการฝึกฝน เอกสารเพื่อฝึกฝน (Training document) และ คลังข้อมูล (Corpus) จะถูกนำมา พิจารณาเพื่อเลือกคำที่เราจะนำมาพิจารณาว่าเป็นคำสำคัญ (Candidate phrase identification) จากนั้น คำที่ถูกเลือกเหล่านี้จะถูกนำมาวิเคราะห์ว่ามีลักษณะเด่นใดบ้าง (Feature Extraction) จากนั้นนำเข้าสู่การเรียนรู้ เพื่อสร้างแบบจำลองของคำสำคัญ และเก็บไว้ในฐานข้อมูล เพื่อใช้ในส่วนของการสกัดคำสำคัญ ต่อไป

ส่วนของการสกัดคำสำคัญ จะใช้ เอกสารที่ต้องการทดสอบมา พิจารณาคำที่เราจะนำมาพิจารณาว่าเป็นคำสำคัญ จากนั้น คำที่ถูกเลือกเหล่านี้จะถูกนำมาวิเคราะห์ว่ามีลักษณะเด่นใดบ้าง เมื่อได้แล้วนำมาเปรียบเทียบกับ แบบจำลองของคำสำคัญที่เก็บไว้ คำที่พิจารณาใดมีใกล้เคียงกับแบบจำลอง N ตัวแรก จะถูกกำหนดว่าเป็น คำสำคัญของเอกสาร



รูปที่ 1.3 แสดงถึงกระบวนการในการสกัดคำสำคัญ (Keyword Extraction)

ถึงแม้ว่าวิธีการนี้จะเป็นที่ยอมรับและใช้กันมากในงานวิจัยทางด้านนี้ เนื่องจากสามารถทำให้เป็นระบบอัตโนมัติได้ง่ายกว่าเทคนิคแรก แต่ก็ยังมีข้อจำกัด ดังนี้

- ในส่วนของการฝึกฝน เพื่อให้ได้แบบจำลองของคำสำคัญที่ดี ต้องการจำนวนเอกสารจำนวนมาก
- ในการเปรียบเทียบความเหมือนกันระหว่างคำที่พิจารณากับแบบจำลองคำสำคัญ ส่วนใหญ่จะเป็นการเปรียบเทียบในแนวดื้น (Shallow Comparison) เช่น มีรูปแบบการเขียนเหมือนกันหรือไม่ ซึ่งในความเป็นจริงแล้ว คำที่เหมือนกันอาจจะมีรูปแบบการเขียนที่ต่างกัน แต่มีความหมายเหมือนกัน

จากข้อจำกัดนี้ ทำให้ คำสำคัญ ที่ได้ ไม่สามารถนำมาใช้เป็นตัวแทนของเอกสารได้ดี เท่ากับ เทคนิคของการกำหนดคำสำคัญ (Keyword Assignment)

1.3 เป้าหมายของการวิจัย

จากข้อจำกัดของ การค้นหาคำสำคัญในปัจจุบัน ที่กล่าวมาข้างต้น งานวิจัยนี้ จึงนำเสนอวิธีการในการปรับปรุงวิธีการในการค้นหาคำสำคัญ โดยมุ่งเน้นจะให้มีความชัดเจนใน 2 เรื่อง คือ

1. เพื่อให้ได้ คำสำคัญ ที่เป็นตัวแทนของเอกสารได้ดี โดยตั้งสมมติฐานว่า คำสำคัญที่เป็นตัวแทนเอกสารที่ดี จะมีความหมายสอดคล้องกับเนื้อหาโดยรวมของเอกสาร ดังนั้นงานวิจัยนี้เสนอการดึงคำสำคัญโดยพิจารณาความหมายของคำที่นำมาเปรียบเทียบกับระหว่างคำที่พิจารณา กับ แม่แบบของคำสำคัญที่กำหนด
2. เพื่อให้ใช้งานได้สะดวก โดยลดความต้องการเอกสารจำนวนมากที่นำมาใช้ในการฝึกฝนเพื่อให้ได้แบบจำลองคำสำคัญ งานวิจัยนี้เสนอ การดึงคำสำคัญ ที่ใช้ แม่แบบทางความหมายของคำสำคัญ ที่ได้มาจากเอกสารที่กำลังพิจารณาเท่านั้น ซึ่ง แม่แบบทางความหมายของคำสำคัญ สามารถหาได้จากเอกสารนั้นๆ โดยพิจารณาจาก ชื่อเรื่องของเอกสารนั้น โดยตั้งสมมติฐานว่า ผู้เขียนเอกสาร มักจะตั้งชื่อเอกสารโดยใช้ คำที่มีความหมายเกี่ยวข้องกับเนื้อหาเอกสาร

1.4 ประโยชน์ที่ได้รับจากงานวิจัย

ผลที่ได้รับจากงานวิจัยนี้ คือ วิธีการในการค้นหาความสำคัญแบบอัตโนมัติ โดยพิจารณาจากความหมายของคำ ว่าสอดคล้องกับแม่แบบคำสำคัญ ซึ่งมีประโยชน์ต่อการประยุกต์ใช้งานต่างๆ ในหัวข้อดังต่อไปนี้

1. รูปแบบจำลองของคำสำคัญเชิงความหมาย ที่ได้จากงานวิจัย สามารถนำไปประยุกต์ใช้กับงานที่มุ่งเน้นการเปรียบเทียบในเชิงความหมาย
2. ขั้นตอนวิธีในการเปรียบเทียบเชิงความหมาย ซึ่งสามารถประยุกต์ใช้ได้กับงานอื่นๆ ที่เป็นการประมวลผลสารสนเทศแบบข้อความ (Text-based Information Processing)
3. วิธีการในการ ปรับปรุงฐานความรู้ของระบบค้นหาคำสำคัญ ให้ทันสมัยอยู่ตลอดเวลา ซึ่งสามารถไปประยุกต์ใช้กับงานอื่นๆ ที่เกี่ยวกับการเรียนรู้เครื่องจักร (Machine Learning) และ การทำเหมืองข้อมูล (Data Mining)

บทที่ 2 ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

ทฤษฎีที่ใช้ในงานวิจัย

รายงานวิจัยฉบับนี้ นำเสนอการใช้ความหมายของคำในประโยคมาช่วยในการค้นหาคำสำคัญของเอกสาร ซึ่งจำเป็นต้องใช้ทฤษฎีที่เกี่ยวข้องดังนี้

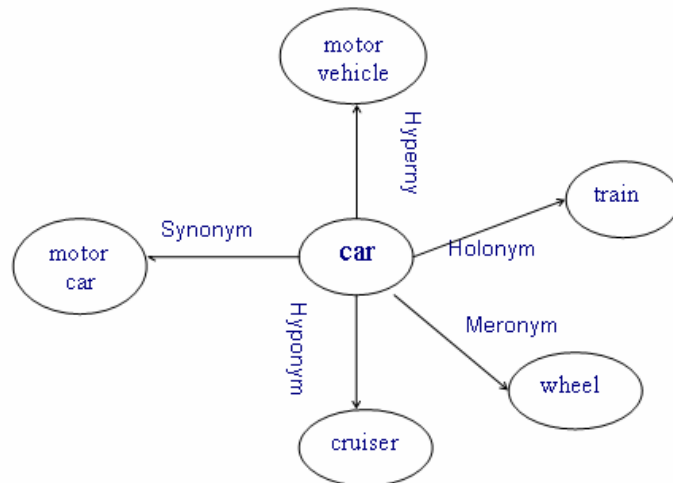
- 2.1.1 Semantic Similarity และ Semantic Relatedness
- 2.1.2 WordNet
- 2.1.3 Knowledge Representation

2.1.1 Semantic Similarity และ Semantic Relatedness

Semantic Similarity คือการเปรียบเทียบคำโดยพิจารณาว่าเหมือนกันในเชิงความหมายหรือไม่ เช่น “car” กับ “bicycle” จะมองว่ามีความหมายเหมือนกัน คือเป็นพาหนะที่มีล้อเหมือนกัน ในการเปรียบเทียบจะได้ผลเป็นมีความหมายเหมือนกันหากผลการเปรียบเทียบที่ได้อยู่ภายในค่า Threshold ที่กำหนด แต่หากเกินค่า Threshold จะถือว่าไม่มีความหมายเหมือนกัน

Semantic Relatedness คือการเปรียบเทียบคำในเชิงความหมาย โดยจะพิจารณาว่ามีความหมายที่มีความสัมพันธ์กันหรือไม่ ซึ่งการพิจารณาความสัมพันธ์เชิงความหมายนั้นเป็นแนวคิดที่เป็นทั่วไปมากกว่าการพิจารณาความหมายเหมือนกัน คือในกรณีที่คำที่ไม่มีความหมายเหมือนกันในเชิงความหมาย แต่อาจจะเป็นคำที่มีความสัมพันธ์กันในเชิงความหมาย เช่น “car” กับ “wheel” ไม่มีความหมายที่เหมือนกัน แต่มีความสัมพันธ์กันในเชิงความหมายในลักษณะที่ “car” มี “wheel” เป็นส่วนประกอบ (Meronym) ซึ่งความสัมพันธ์กันในเชิงความหมายนั้นสามารถแบ่งลักษณะความสัมพันธ์ได้ดังนี้

- Synonym มีความหมายเหมือนกัน
- Antonym มีความหมายที่ตรงกันข้ามกัน
- Hypernym มีความหมายที่เป็นแบบทั่วไปกว่า
- Hyponym มีความหมายที่เป็นแบบเฉพาะเจาะจงกว่า
- Holonym มีความสัมพันธ์ทางความหมายในลักษณะที่เป็นส่วนประกอบหรือสมาชิกของ
- Meronym มีความสัมพันธ์ทางความหมายในลักษณะที่มีเป็นส่วนประกอบหรือมีเป็นสมาชิก



รูปที่ 2.1 แสดงความสัมพันธ์ลักษณะต่าง ๆ ของ “car” กับคำอื่น ๆ

ซึ่งการเปรียบเทียบเพื่อหาความสัมพันธ์ในเชิงความหมายนี้ใช้ในการสรุปใจความเอกสาร (Text Summarization) และระบบค้นคืนสารสนเทศ (Information Retrieval) ซึ่งจะทำให้ระบบงานค้นคืนสารสนเทศสามารถค้นคืนข้อมูลหรือเอกสารที่มีความสัมพันธ์กัน ได้ครอบคลุมความต้องการมากขึ้น

ในการเปรียบเทียบเพื่อหาความสัมพันธ์ในเชิงความหมายจะพิจารณาจาก

- โครงข่ายของเวิร์ดเน็ต (WordNet) ซึ่งคำสองคำที่มีระยะห่างน้อยจะถือว่ามีค่าความสัมพันธ์กันมาก โดยหาความสัมพันธ์เชิงความหมายของคำจากการคำนวณจากเส้นทางระหว่างคำ (Synset) ที่สั้นที่สุดในเวิร์ดเน็ต งานวิจัยในกลุ่มนี้มีข้อดีคือง่าย แต่มีข้อจำกัดตรงที่คำที่นำมาเปรียบเทียบกันนั้นต้องมีอยู่ใน ฐานข้อมูล WordNet นอกจากนี้ คำที่เขียนเหมือนกันอาจอยู่ได้หลายที่ใน WordNet ขึ้นอยู่กับความหมาย เช่น คำว่า “mouse” อาจหมายถึง “หนูที่เป็นสัตว์” หรือ “อุปกรณ์คอมพิวเตอร์” ทำให้คำที่ได้อาจจะผิดพลาดได้
- พิจารณาความสัมพันธ์จากคำข้อมูล (Information Content หรือ Corpus Based) หรือข้อความที่เนื้อหาของคำทั้งสองสถิติอยู่ โดยวิธีนี้จะใช้การคำนวณทางสถิติ เช่น การกระจายของคำ, จำนวนครั้งที่เกิดขึ้น, ตำแหน่งของคำ ด้วยวิธีนี้สามารถแก้ปัญหาของคำที่ไม่มีใน WordNet ของกลุ่มแรกได้ แต่อย่างไรก็ตามวิธีนี้นี้จำเป็นต้องมีการสร้างคลังข้อมูลตัวอักษร (Corpus) และผลการคำนวณที่ได้ก็จะขึ้นอยู่กับขนาดและข้อมูลที่มีในคลังข้อมูลนั้นๆ
- พิจารณาจากการเปรียบเทียบคำจำกัดความหรือคำนิยามหรือคำแปลของคำทั้งคู่ (Gloss based) โดยดูคำนวณจากคำที่ใช้ร่วมกัน (Shared word) ในการอธิบายความหมายของคำทั้งคู่

ซึ่งในการหาค่าสัมพันธ์ในเชิงความหมายโดยพิจารณาจากข้อความที่เป็นเนื้อหาหรือการคำแปลที่มีการใช้คำร่วมกัน ต่อมาได้มีการปรับปรุงโดยในการเปรียบเทียบกันระหว่างเนื้อหาหรือคำแปลนั้นนอกจากจะพิจารณาความเหมือนของคำที่ใช้ร่วมกันแล้วยังมีการพิจารณาความสัมพันธ์ในเชิงความหมายที่เป็น Synonym, Hypernym และ Hyponym ด้วย โดยใช้เว็รคเน็ตเป็นฐานความรู้ (Knowledge-based) ในการพิจารณาเปรียบเทียบหาความสัมพันธ์ในระดับต่าง ๆ

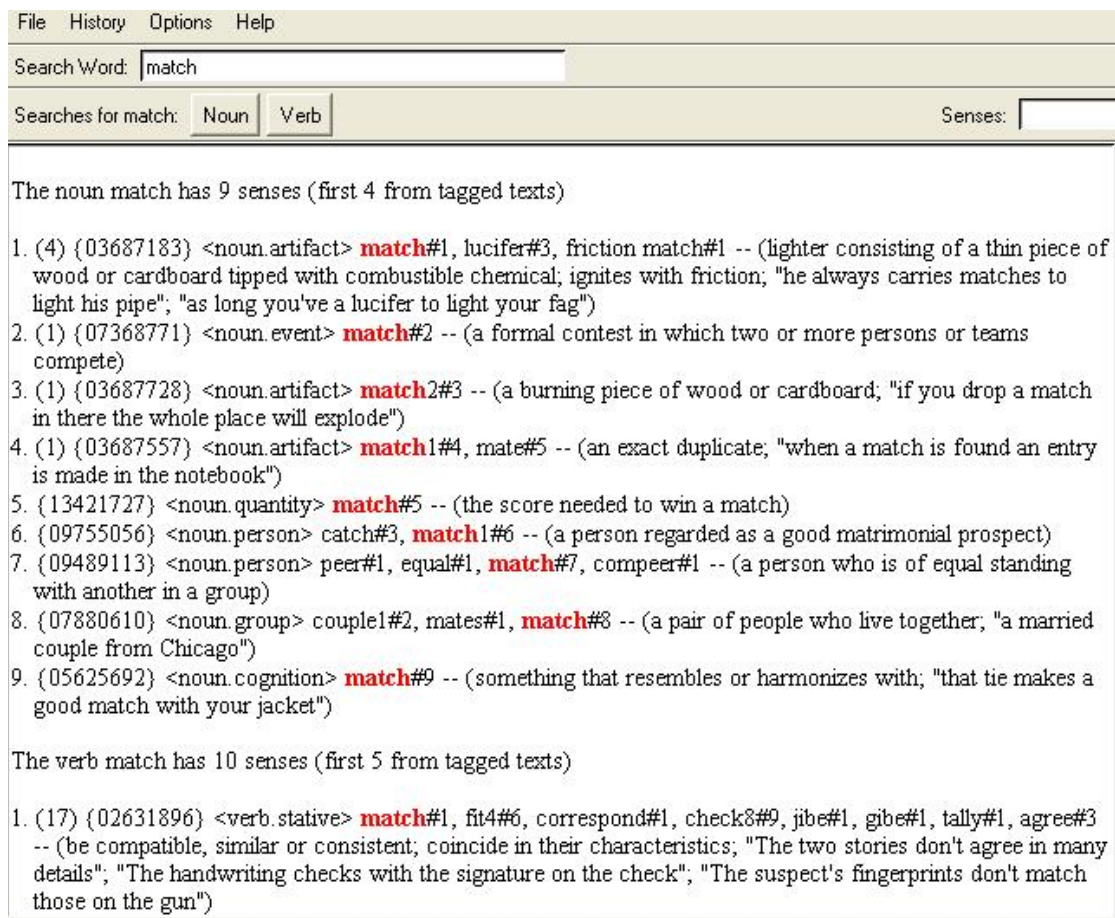
2.1.2 เว็รคเน็ต (WordNet)

เว็รคเน็ต (Cognitive Science Laboratory at Princeton University) เป็นฐานข้อมูลพจนานุกรมภาษาอังกฤษผู้ใช้สามารถdownload มาใช้หรือใช้งาน ในลักษณะออนไลน์ได้ โดยเว็รคเน็ตได้รวมคำในภาษาอังกฤษที่มีความหมาย(Sense) เหมือนกันเข้าเป็นกลุ่มเรียกว่า Synset เช่น “good” “right” และ “ripe” เป็น Synset เดียวกัน มีความหมายเหมือนกัน คือ “most suitable or right for a particular purpose”

ในระหว่าง Synset จะสัมพันธ์กันโดยมีความสัมพันธ์เชิงความหมายหลายรูปแบบคำ เช่น มีความหมายเหมือนกัน (Synonym) หรือมีความหมายตรงกันข้าม (Antonym) ซึ่งความสัมพันธ์จะขึ้นอยู่กับชนิดของคำโดยจะแยกความแตกต่างกันระหว่างคำที่เป็น Noun, Verb, Adjective และ Adverb ตามหลักไวยากรณ์ Synset ที่มีความสัมพันธ์กันเชิงความหมายจะมีหมายเลขที่เชื่อมต่อกันทุก Synset ที่เชื่อมต่อกันโดยตัวเลขของความสัมพันธ์เชิงความหมาย

คำ 1 คำในเว็รคเน็ตอาจมีได้หลายความหมาย นั่นคือคำ 1 คำ อาจปรากฏอยู่ในหลาย Synset โดยที่แต่ละ Synset มีหมายเลขกำกับที่แตกต่างกัน ในขณะที่เดียวกันอาจจะมีคำหลายคำที่มีความหมายเดียวกัน จากรูปที่ 2.1 จะเห็นได้ว่า “match” ที่เป็นคำนามมี 9 ความหมาย และ “match” ที่เป็นคำกริยามี 10 ความหมาย ในขณะที่บางความหมายมีคำที่มีความหมายเหมือนกัน (Synonym) ด้วย เช่น ว่า “match” ที่เป็นคำนามความหมายที่ 1 มี lucifer, friction, match เป็นคำที่มีความหมายเดียวกัน

เนื่องจากเว็รคเน็ตเป็นฐานข้อมูลพจนานุกรมซึ่งได้รวบรวมความหมายต่าง ๆ ของทั้งที่เป็น Noun, Verb, Adjective และ Adverb และสามารถเชื่อมโยงความสัมพันธ์เชิงความหมายระหว่างคำในแบบต่าง ๆ ได้ครอบคลุม รวมทั้งเว็รคเน็ตถูกพัฒนาและมีการเพิ่มคำศัพท์มาอย่างต่อเนื่องจนถึงปัจจุบัน ซึ่งประกอบไปด้วยข้อมูลของคำทั้งหมดประมาณ 90,000 คำ งานวิจัยนี้จึงเลือกใช้ เว็รคเน็ตเป็นฐานความรู้ในการเปรียบเทียบเพื่อคำนวณหาความสัมพันธ์



รูปที่ 2.2 ตัวอย่างการแสดงผลของเวิร์ดเน็ตจากการค้นหาคำว่า “match”

2.1.3 การแทนความรู้ (Knowledge Representation)

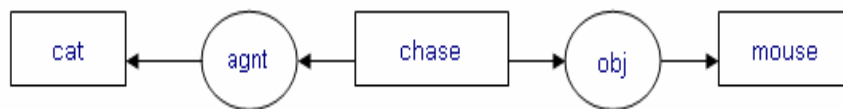
รายงานวิจัยฉบับนี้ใช้กราฟเชิงความคิด (Conceptual Graph) ในการแทนความรู้ กราฟเชิงความคิด (Chein et Marie-Laure Mugnier, M.,1992) คิดค้นโดย John F. Sowa เพื่อนำมาใช้เป็นรูปแบบการแทนความรู้ และนำมาประยุกต์ใช้กับการแสดงความหมายทางภาษาเพื่อบ่งชี้ความรู้ (Knowledge Acquisition) ในการวิจัยทางด้านปัญญาประดิษฐ์ (Artificial Intelligence) จุดเด่นที่สำคัญของกราฟเชิงความคิดคือสามารถอธิบายภาษาธรรมชาติ (Natural Language) ออกมาในรูปแบบของสัญลักษณ์เชิงตรรกศาสตร์และเป็นตัวกลางในการสื่อความหมายระหว่างภาษาของคอมพิวเตอร์และภาษาของมนุษย์ ทำให้การสื่อแทนความหมายโดยใช้กราฟเชิงความคิดนั้นเข้าใจได้ง่ายและพัฒนาต่อไปในการหาเหตุผลเชิงตรรกะได้ ในการแทนความรู้โดยใช้กราฟเชิงความคิดประกอบด้วย

1. Concept Nodes แทนด้วย สัญลักษณ์รูปสี่เหลี่ยม ซึ่งจะแสดง Entities, Actions และ Attributes ซึ่งมี 2 Attributes

- Type จะเป็นตัวบอก class ของ element ที่แทนด้วย concept
- Referent จะเป็นตัวบอกความเฉพาะเจาะจงของ class ที่ถูกอ้างโดย node

2. Conceptual Relation Nodes แทนด้วย สัญลักษณ์รูปวงกลม จะแสดงความสัมพันธ์ระหว่าง Concept node โดย Attributes ของ Relation Nodes จะมี 2 Attributes คือ

- Valence จะเป็นตัวบอกจำนวนของ concept ที่อยู่ข้าง ๆ relation
- Type แสดงความเกี่ยวข้องของตัวมันเอง เช่น subject, attribute, และ object



รูปที่ 2.3 แสดงตัวอย่าง Conceptual Graphs

ในรูปที่ 2.2 แสดง conceptual graph ของประโยค “Tom is chasing a brown mouse” ประกอบด้วย 3 concepts และ 2 relations ซึ่งสามารถเขียนในรูปแบบการแสดงผลที่ไม่เป็นกราฟิกได้ดังนี้

[cat: Tom]←-(Agnt)←-[chase]→(Ptnt)→[mouse]→(Attr)→[brown]

รูปที่ 2.4 แสดงตัวอย่าง Conceptual Graphs

ทิศทางของลูกศรในรูปด้านบนแสดงถึงประธาน (Subject) และกรรม (Object) ของความเกี่ยวข้อง (Relation) concept [cat: Tom] เป็น concept ที่เฉพาะเจาะจงของ Type cat ในขณะที่ [chase] และ [mouse] เป็น generic concept ความเกี่ยวข้อง (Attr) บอก attribute ของ mouseว่ามีสี brown

หลักการอีกส่วนหนึ่งที่สำคัญในการใช้กราฟเชิงความคิดในการแทนความรู้ คือการแสดงความเกี่ยวข้องระหว่างกลุ่มขององค์ประกอบแบบลำดับชั้น (Type hierarchy) หลักการดังกล่าวทำให้สามารถนำเอาคุณสมบัติการสืบทอด (Inheritance) มาประยุกต์ใช้ กล่าวคือองค์ประกอบที่มีลักษณะเฉพาะเรียกว่าเป็น Subtype ขององค์ประกอบที่มีลักษณะทั่วไปซึ่งเรียกว่า Supertype สิ่งนี้

เป็น Subtype ขององค์ประกอบหนึ่ง ๆ จะสืบทอดคุณสมบัติขององค์ประกอบนั้น ๆ นอกจากนี้หลักการของกราฟเชิงความคิดได้ถูกนำไปใช้ในระบบงานค้นคืนสารสนเทศ (Montes-y-Gomez, Lopez-Lopez, and Gelbukh, A., 2000) (Montes-y-Gomez, Lopez-Lopez, Gelbukh, and Baeza-Yates, 2001) โดยอาศัยหลักการเปรียบเทียบกันระหว่างกราฟ (Graph Matching)

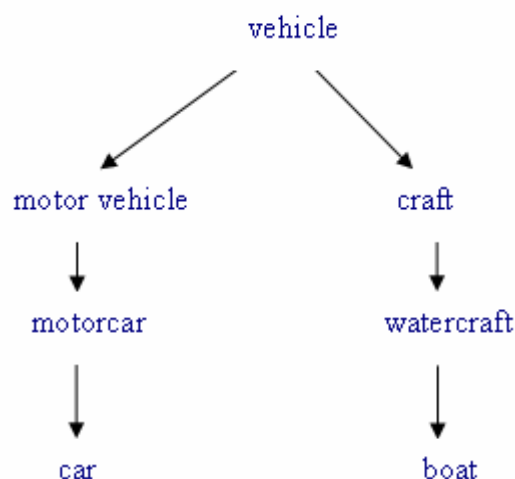
เนื่องจากกราฟเชิงความคิดสามารถแสดงส่วนประกอบ (concept) ในประโยค รวมทั้งแสดงความสัมพันธ์ของส่วนประกอบในประโยคนั้นด้วย ในงานวิจัยนี้จึงใช้กราฟเชิงความคิดในการแทนความหมายของประโยค

2.2 งานวิจัยอื่น ๆ ที่เกี่ยวข้อง

2.2.1 การเปรียบเทียบหาความเหมือนเชิงความหมาย

Resnik (1995) และ Jiang and Conrath (1997) เสนอแนวคิดที่ว่าคำจะมีความหมายเหมือนกันนั้นสามารถพิจารณาได้จากความเหมือนกันของข้อความที่เนื้อหาของคำทั้งสอง ซึ่งวิธีการนี้จำเป็นต้องมีสร้างฐานความรู้คลังข้อมูลตัวอักษร (Corpus) และผลการคำนวณที่ได้ก็จะขึ้นอยู่กับขนาดและข้อมูลที่มีใน คลังข้อมูลตัวอักษร (Corpus) ด้วย

Leacock and Chodorow (1998) เสนอการคำนวณหาความเหมือนกันทางความหมายของคำโดยคำนวณจากระยะห่างระหว่างคำ (Synset) ทั้งสองที่สั้นที่สุดโดยอาศัยโครงข่ายในเวิร์ดเน็ตโดยพิจารณาความสัมพันธ์แบบ “is-a”



รูปที่ 2.5 แสดงระยะห่างระหว่าง “car” และ “boat” ในโครงข่ายเวิร์ดเน็ตเท่ากับ 6

2.2.2 การเปรียบเทียบหาค่าความสัมพันธ์เชิงความหมาย

Hist and St-On (1998) คำนวณค่าความสัมพันธ์เชิงความหมายโดยอาศัย โครงข่ายในเวิร์ดเน็ต จะพิจารณาจำนวนครั้งของการเปลี่ยนทิศทางของเส้นทาง มีสูตรในการคำนวณ คือ

$$Rel(c_1, c_2) = C - path\ length - k * d \quad (2.1)$$

โดยที่ d จำนวนครั้งของการเปลี่ยนทิศทางใน path,

C และ k เป็นค่าคงที่

ซึ่งเส้นทางความสัมพันธ์ที่พิจารณานั้นประกอบด้วย Synonym, Hyponym, Hypernym และ Holonym โดยที่จะให้ความสำคัญในแต่ละรูปแบบความสัมพันธ์ที่เท่ากัน

Banerjee and Pederson (2003) เสนอวิธีการคำนวณหาความสัมพันธ์ในเชิงความหมายของคำซึ่งได้ประยุกต์ใช้หลักการของ Lesk (1986) ซึ่งจะนับจำนวนซึ่งสถิตในคำแปล (gloss) ของคำทั้งสองที่เหมือนกัน (overlap) โดยใช้ความคำแปลและโครงข่ายในเวิร์ดเน็ตเป็นฐานความรู้ แต่จะขยาย set ของคำแปลที่จะเปรียบเทียบโดยการนำเอาคำผลการเปรียบเทียบหาคำที่สถิตอยู่เหมือนกันในคำแปลของคำที่เป็น Hyponym และ Hypernym ของคำที่ต้องการเปรียบเทียบทั้งสองมารวมคำนวณค่าสัมพันธ์ด้วย ดังนี้

$$\begin{aligned} Relatedness(A,B) = & score(gloss(A), gloss(B)) + score(hype(A), hype(B)) + \\ & score(hypo(A), hypo(B)) + score(hype(A), gloss(B)) + \\ & score(gloss(A), hype(B)) \end{aligned} \quad (2.2)$$

นอกจากนี้ยังมีการให้น้ำหนักตามความยาวของวลีที่ปรากฏด้วย เช่น ปรากฏคำว่า “bank” เหมือนกันจะได้คะแนนเป็น 1 ในขณะที่ปรากฏคำว่า “bank account” เหมือนกันคะแนนที่ได้จะเป็น 2^2 เท่ากับ 4 เป็นต้น

Strube and Ponzetto (2006) เสนอวิธีการ WikiRelate โดยใช้หลักการเดียวกันกับ Lesk (1986) โดยใช้คำแปลหรือคำนิยามของคำใน Wikipedia ซึ่งเป็นพจนานุกรมภาษาอังกฤษแบบออนไลน์ในการคำนวณหาความสัมพันธ์ในเชิงความหมาย โดยคำแปลที่นำมาคำนวณค่าความสัมพันธ์นั้นจะใช้ข้อความจากย่อหน้าแรกใน Wikipedia ในการเปรียบเทียบคำที่สถิตอยู่

เหมือนกันนั้นมีการให้น้ำหนัก ตามความยาวของวลีที่ปรากฏเช่นเดียวกับ Banerjee and Pederson (2003) โดยมีสูตรในการคำนวณหาความสัมพันธ์ดังนี้

$$Related_{gloss}(t_1, t_2) = \tanh(\text{overlap}(t_1, t_2) / \text{length}(t_1) + \text{length}(t_2)) \quad (2.3)$$

ซึ่งเมื่อเปรียบเทียบผลที่ได้กับการคำนวณหาความสัมพันธ์กันโดยใช้นิยามของคำจากเวิร์ดเน็ต ปรากฏว่าให้ผลที่ดีกว่าการใช้คำนิยามจากเวิร์ดเน็ต

Zhang, Sun, Wang, and Bai, (2005) นำเสนอการวัดความสัมพันธ์ของประโยคกับหัวข้อ เอกสารโดยการเปรียบเทียบหาความสัมพันธ์ในเชิงความหมายของประโยค ค่าความสัมพันธ์จะคำนวณจากผลรวมของค่าความสัมพันธ์ของคำในประโยค A กับประโยค B รวมกับผลรวมของค่าความสัมพันธ์ของคำในประโยค B กับประโยค A หาดด้วยจำนวนคำทั้งหมดในประโยคทั้งสอง โดยค่าความสัมพันธ์ของคำในประโยค A กับประโยค B คือระยะห่างระหว่างคำในประโยค A กับคำในประโยค B ที่สั้นที่สุด ซึ่งในการคำนวณหาระยะห่างระหว่างคำที่สั้นที่สุดนั้นจะพิจารณาความสัมพันธ์จากฐานข้อมูลในเวิร์ดเน็ตระดับ Synonym Hyponym Hypernym รวมทั้ง Antonym และ derivation ด้วย โดยจะให้น้ำในความสัมพันธ์ระดับ Antonym และ derivation ในระดับต่ำสุด ถึงแม้ว่าจะพิจารณาความสัมพันธ์เชิงความหมายในระดับประโยคแต่ก็ไม่ได้พิจารณาความสัมพันธ์กันของคำในประโยคนั้น ๆ แต่อย่างใด

Yang and Powers (2005) ได้เสนอวิธีการในการวัดค่าความสัมพันธ์ทางความหมายระหว่างคำจากระยะห่างโหนดโดยอาศัยโครงข่ายในเวิร์ดเน็ต โดยความสัมพันธ์ที่พิจารณาได้แก่ Synonym, Antonym, Hypernym, Hyponym, Antonym, Holonym และ Meronym คือนอกเหนือจากพิจารณาความสัมพันธ์ “is-a” แล้ว ยังพิจารณาความสัมพันธ์ “part-of”, อีกด้วย มีการให้น้ำหนักความสัมพันธ์ในแบบ Synonym และ Antonym มากกว่าความสัมพันธ์ในแบบอื่นๆ ซึ่งผลการทดลองกับชุดคู่ลำดับของคำมาตรฐาน 28 คู่ โดยการพิจารณาของมนุษย์ปรากฏว่าให้ผลที่ดีกว่าเมื่อเปรียบเทียบกับงานวิจัยที่ผ่านมา

ตารางที่ 2.1
เปรียบเทียบงานวิจัยที่เกี่ยวข้อง

ผู้เขียน	เสนอ	ข้อดี	ข้อเสีย
Hist and St-Ong (1998)	คำนวณค่าความสัมพันธ์เชิงความหมายโดยอาศัยโครงข่ายในเวิร์ดเน็ต จะพิจารณาจำนวนครั้งของการเปลี่ยนทิศทางของเส้นทาง	ทำได้ง่าย	พิจารณาแค่ความสัมพันธ์ในแบบ Synonym และ Hyponym/Hypernym เท่านั้น
Lesk M. (1986)	นับจำนวนซึ่งปรากฏในคำแปล (gloss) ของคำทั้งสองที่เหมือนกัน (overlap)	ทำได้ง่าย	พิจารณาเฉพาะคำที่เหมือนกันใน gloss เท่านั้น รวมทั้งมีข้อจำกัดในเรื่องคำแปล (gloss) ที่มีอยู่ใน dictionary อาจให้ข้อมูลที่ไมเพียงพอ
Banerjee and Pederson (2003)	นับจำนวนซึ่งปรากฏในคำแปล (gloss) ของคำทั้งสองที่เหมือนกัน (overlap) โดยมีการให้น้ำหนักตามความยาวของวลีที่ปรากฏด้วย เช่น ปรากฏคำว่า “bank account” เหมือนกัน คะแนนที่ได้จะเป็น 2 2	ไม่พิจารณาเฉพาะ gloss ของ คำนั้น ๆ เท่านั้น แต่ยังพิจารณา gloss ของคำที่เป็น Hyponym, Hypernym, Meronym Holonym and Troponym Synsets ของคำกำหนดอีกด้วย	ไม่ได้พิจารณาคำอื่น ๆ ที่อยู่ในประโยคหรือข้อความที่อยู่รอบ ๆ

ตารางที่ 2.1 (ต่อ)

ผู้เขียน	เสนอ	ข้อดี	ข้อเสีย
Strube and Ponzetto (2006)	ใช้ความหมายหรือคำนิยามของคำ ใน Wikipedia ซึ่งเป็นพจนานุกรมภาษาอังกฤษแบบออนไลน์ในการคำนวณหาความสัมพันธ์ในเชิงความหมาย ซึ่งอาศัยหลักการของ Lesk และ Banerjee, Pederson เช่นกัน โดยมีการ normalize ค่าที่คำนวณได้อีกที	เนื่องจาก Wikipedia เป็น on-line encyclopedia database ที่มีขนาดใหญ่ ทำให้สามารถคำนวณค่าความสัมพันธ์จาก คำแปล (gloss) หรือ full text page ใน Wikipedia ได้	เปรียบเทียบคำที่เหมือนกันเท่านั้น
Yang and Powers (2005)	เสนอการพิจารณาความสัมพันธ์แบบ Synonym, Antonym, Hypernym, Hyponym, Antonym, Holonym และ Meronym คือ นอกเหนือจากพิจารณาความสัมพันธ์ “is-a” แล้ว ยังพิจารณาความสัมพันธ์ “part-of”, อีกด้วย มีการให้น้ำหนักความสัมพันธ์	ครอบคลุมความสัมพันธ์ในเชิงความหมายมากขึ้น	เป็นการพิจารณาความสัมพันธ์ในระดับคำเท่านั้น

ตารางที่ 2.1 (ต่อ)

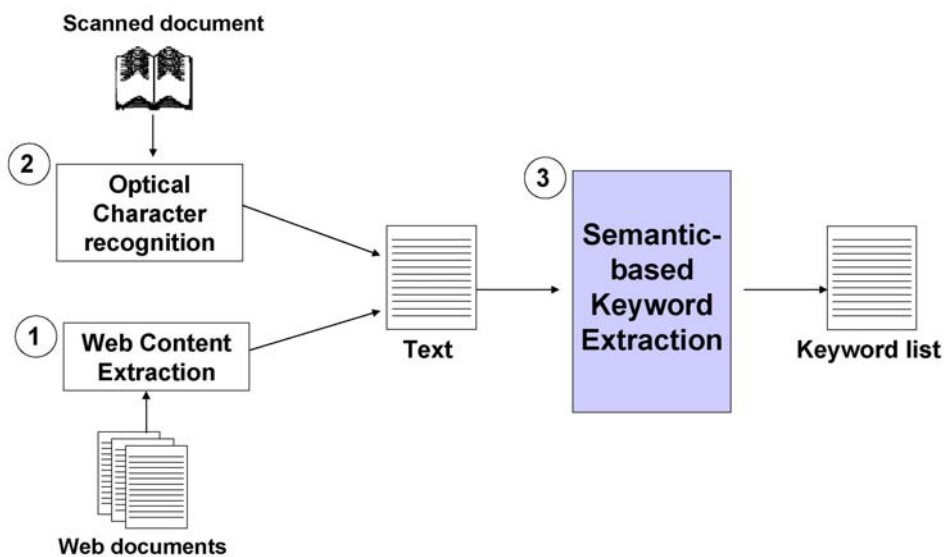
ผู้เขียน	เสนอ	ข้อดี	ข้อเสีย
Zhang, Sun, Wang and Bai (2005)	เสนอการหาประโยชน์จะ สัมพันธ์กับหัวข้อ เอกสารโดยการ เปรียบเทียบเชิง ความหมายของประโยค คำนวณค่าความสัมพันธ์ จากผลรวมของค่า ความสัมพันธ์ของคำใน ประโยค A กับประโยค B รวมกัน แล้วหารด้วย จำนวนคำทั้งหมดใน ประโยคทั้งสอง	เป็นการพิจารณา ความสัมพันธ์ในระดับ ประโยค	ในการเปรียบเทียบ ระหว่างคำภายในประโยค พิจารณาแค่ความสัมพันธ์ ในแบบ Synonym และ Hyponym/Hypernym เท่านั้น และไม่ได้ พิจารณาความสัมพันธ์ ระหว่างคำในประโยค

จากงานวิจัยที่เกี่ยวข้องกับการเปรียบเทียบค่าความสัมพันธ์ในเชิงความหมายนั้นยังมีข้อจำกัด ซึ่งสรุปได้ดังนี้

1. ระดับของการเปรียบเทียบ งานวิจัยส่วนใหญ่ใช้การเปรียบเทียบในระดับคำ เพื่อเปรียบเทียบความเหมือนกันระหว่างคำ แต่เนื่องจากว่าคำหนึ่งคำอาจแทนความหมายได้หลายความหมาย อาจได้ผลลัพธ์ที่ไม่ถูกต้อง
2. ความสัมพันธ์ของคำ งานวิจัยโดยทั่วไปจะเปรียบเทียบเฉพาะคำที่มีความหมายเหมือนกัน (Synonym) ในความสัมพันธ์ “is-a” ซึ่งในการใช้งานจริงๆ แล้ว ยังมีความสัมพันธ์อื่นๆ มาใช้ในการเปรียบเทียบได้ด้วย เช่น ความสัมพันธ์ “part-of”, ความสัมพันธ์ “member-is-a-kind-of”, ความสัมพันธ์ “member-is-a-part-of” เป็นต้น

บทที่ 3 ระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย

ในบทนี้จะกล่าวถึงโครงสร้างของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมายที่ได้นำเสนอ ซึ่งนำเอาหลักการของการเปรียบเทียบเชิงความหมายมาประยุกต์ใช้ในการค้นหาคำสำคัญ รูปที่ 3.1 แสดงแผนภาพโดยรวมของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย โดยระบบที่นำเสนอนี้สามารถค้นหาคำสำคัญจากเอกสารได้ 2 ลักษณะคือ เอกสารที่อยู่ในรูปของข้อความบนเว็บ และเอกสารที่เป็น Hard copy ที่ได้ผ่านการสแกนจากเครื่องตรวจกราด



รูปที่ 3.1 แผนภาพโดยรวมของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย

เมื่อระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมายได้รับเอกสารสำหรับค้นหาคำสำคัญจากผู้ใช้ ก็จะทำเตรียมเอกสารเหล่านั้นให้อยู่ในรูปของเอกสารข้อความ โดยผ่านกระบวนการ (1) การเลือกเฉพาะเนื้อหา (Web Content Extraction) ถ้าเอกสารที่เข้ามาเป็นข้อความบนเว็บ หรือ กระบวนการ (2) การรู้จำตัวอักษร (Optical Character Recognition) ถ้าเอกสารที่เข้ามาเป็น Hard copy ที่ได้ผ่านการสแกนจากเครื่องตรวจกราด เมื่อข้อมูลเข้าถูกแปลงเป็นเอกสารข้อความเรียบร้อยแล้ว ก็จะถูกส่งเข้ายังกระบวนการ (3) การค้นหาคำสำคัญโดยใช้ความหมาย (Semantic-based Keyword Extraction) ซึ่งผลลัพธ์ที่ได้จะเป็นรายการของคำสำคัญของเอกสารนั้น โดยคำสำคัญเหล่านี้คาดหมายว่าจะมีความหมายสอดคล้องกับเนื้อหาโดยรวมของเอกสาร

ในส่วนต่อไปจะกล่าวถึงรายละเอียดของกระบวนการต่างๆของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย

(1) การเลือกเฉพาะเนื้อหา (Web Content Extraction)

ในปัจจุบัน เอกสารที่สำคัญต่างๆสามารถถูกเรียกสืบค้นได้โดยผ่านทางเครือข่ายโดยโปรแกรมเว็บเบราว์เซอร์ (Web Browser) ซึ่งเอกสารเหล่านั้นส่วนใหญ่มักจะอยู่ในรูปของ HTML (Hyper Text Mark-up Language) ซึ่งประกอบด้วย แถบป้าย (tag) มากมาย ดังแสดงเป็นตัวอย่างตามตาราง 3.1

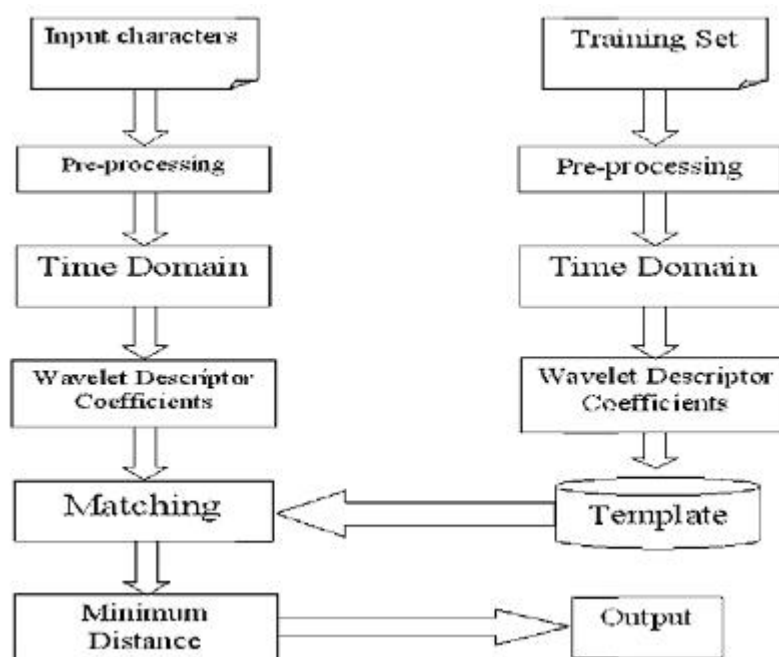
ตารางที่ 3.1 ตัวอย่างของแถบป้ายที่ใช้ในแฟ้มข้อมูลประเภท HTML

Tag	Description	DTD
<u><!--...--></u>	Defines a comment	STF
<u><!DOCTYPE></u>	Defines the document type	STF
<u><abbr></u>	Defines an abbreviation	STF
<u><acronym></u>	Defines an acronym	STF
<u><applet></u>	Defines an applet	TF
<u><area></u>	Defines an area inside an image map	STF
<u><basefont></u>	Defines a base font	TF
<u><bdo></u>	Defines the direction of text display	STF
<u><blockquote></u>	Defines a long quotation	STF
<u><body></u>	Defines the body element	STF
<u>
</u>	Inserts a single line break	STF
<u><caption></u>	Defines a table caption	STF
<u><center></u>	Defines centered text	TF
<u><colgroup></u>	Defines groups of table columns	STF
<u><div></u>	Defines a section in a document	STF
<u><dl></u>	Defines a definition list	STF
<u><dt></u>	Defines a definition term	STF
<u></u>	Defines text font, size, and color	TF
<u><h1> to <h6></u>	Defines header 1 to header 6	STF
<u><head></u>	Defines information about the document	STF
<u><hr></u>	Defines a horizontal rule	STF
<u><map></u>	Defines an image map	STF
<u><menu></u>	Defines a menu list	TF
<u><object></u>	Defines an embedded object	STF
<u><p></u>	Defines a paragraph	STF
<u><q></u>	Defines a short quotation	STF
<u><script></u>	Defines a script	STF
<u><select></u>	Defines a selectable list	STF
<u><small></u>	Defines small text	STF
<u></u>	Defines a section in a document	STF
<u><style></u>	Defines a style definition	STF
<u><table></u>	Defines a table	STF
<u><tbody></u>	Defines a table body	STF
<u><td></u>	Defines a table cell	STF
<u><textarea></u>	Defines a text area	STF
<u><tfoot></u>	Defines a table footer	STF
<u><th></u>	Defines a table header	STF
<u><thead></u>	Defines a table header	STF
<u><title></u>	Defines the document title	STF
<u><tr></u>	Defines a table row	STF

ดังนั้น กระบวนการนี้จะเป็นการเลือกเอาเฉพาะส่วนของเนื้อหาออกมาจากเอกสาร กระบวนการนี้จะอ่านเอกสารจำนวน 2 รอบ โดยรอบที่ 1 จะอ่านเอกสารเพื่อข้ามส่วนของ ตาราง เมนู รูปภาพ และ การกำหนดลักษณะพิเศษ จากนั้นบันทึกข้อมูลเก็บไว้ในรอบที่ 2 จะเป็นการอ่าน เพื่อตัดแถบป้ายออก ให้เหลือเฉพาะ ส่วนของเนื้อหา และจัดเก็บเป็นแฟ้มข้อมูลชนิดที่เป็น ข้อความล้วน (Text file)

(2) การรู้จำตัวอักษร (Optical Character Recognition)

ในกรณีที่เอกสารที่รับเข้ามาอยู่ในรูปของ Hard Copy กระบวนการนี้จะทำหน้าที่ในการแปลง เอกสารที่อยู่ในรูปของภาพให้เป็นเอกสารข้อความ โดยจะใช้หลักการของรู้จักตัวอักษรบนโดเมน ความถี่ กระบวนการนี้ได้ประยุกต์ใช้วิธีการของ (Chey, C. et al. 2006) ที่เลือกใช้ Wavelet descriptor มาเป็นเทคนิคสำคัญในการสร้างแม่แบบ (template) ของแต่ละตัวอักษร รายละเอียดของ ขั้นตอนมีดังนี้



รูปที่ 3.2 ขั้นตอนการรู้จำตัวอักษร (Optical Character Recognition)

รูปที่ 3.2 แสดงถึงแผนภาพของระบบรู้จำตัวอักษร ประยุกต์มาจากการรู้จำภาษาเขมรโดยใช้ Wavelet Descriptor (Percival, D. B. and Walden, A. T.: 2000) ทางด้านขวามือของรูป ชุดของตัวอักษรสำหรับการฝึกฝน (Training Set) ถูกนำมาใช้ในการสร้างแม่แบบ (Template) ซึ่ง

จัดเก็บไว้ในรูปแบบของค่าเฉลี่ยของสัมประสิทธิ์ของ Wavelet Descriptor ซึ่งคำนวณมาจากข้อมูลประเภทที่สามารถคำนวณได้ภายในช่วงเวลา (Temporal Data) ของแต่ละตัวอักษรของชุดของการฝึกฝน เมื่อแม่แบบของตัวอักษรแต่ละตัวได้ถูกสร้างขึ้น ก็จะถูกนำมาเปรียบเทียบกับ สัมประสิทธิ์ของ Wavelet Descriptor ของภาพตัวอักษรที่ได้รับการอินพุตเข้ามา ถ้าแม่แบบของตัวอักษรตัวใดมีค่าระยะห่างของความแตกต่าง (Minimum Distance) น้อยที่สุด ก็จะเป็นคำตอบของการรู้จำตัวอักษรนั้นๆ รายละเอียดของระบบรู้จำตัวอักษร มีดังนี้

(1) การเตรียมข้อมูล (Pre-processing)

ในขั้นตอนนี้ ภาพตัวอักษรที่ได้มาจากการสแกนจะถูกนำมาเตรียมด้วยการแปลงเป็นรูปขาวดำ โดยมีการกำหนดระดับเงื่อนไข (Threshold) ค่าความเข้มของแต่ละจุดของภาพ (Pixel) จะถูกนำมาเปรียบเทียบกับระดับเงื่อนไขนี้ ถ้าค่าความเข้มของแต่ละจุดของภาพ สูงกว่า ระดับเงื่อนไขจะถูกกำหนดให้เป็นสีดำ และในทางตรงกันข้ามจะเป็นสีขาว เมื่อได้รูปภาพตัวอักษรเป็นสีขาว-ดำแล้ว รูปภาพเหล่านี้จะถูกแยกออกเป็นตัวอักษรแต่ละตัว โดยการหาโครงร่างของแต่ละตัวอักษร เรียกว่า Skeletonization เมื่อแยกตัวอักษรออกเป็นแต่ละตัวได้แล้ว ก็จะถูกนำไปหารูปร่างของตัวอักษรด้วยการทำให้บางลง (Thinning) รูปภาพตัวอักษรแต่ละตัวจะถูกส่งไปยังขั้นตอนการดึงเอาลักษณะเด่นของตัวอักษร ต่อไป

(2) การดึงเอาลักษณะเด่นของตัวอักษร (Feature Extraction)

ขั้นตอนนี้ถือว่าสำคัญที่สุดในการรู้จำตัวอักษร เพราะลักษณะเด่นของตัวอักษรจะนำไปสู่การรู้จำตัวอักษรที่มีความถูกต้องสูง งานวิจัยนี้จะเริ่มต้นด้วย

- การแปลงโครงร่างของตัวอักษรที่ได้จากขั้นตอนการเตรียมข้อมูลให้อยู่ในรูปของ Temporal Domain ด้วยสมการที่ (3-1)

$$r(t) = \sqrt{(X(t) - cg_x)^2 + (Y(t) - cg_y)^2} \quad (3-1)$$

เมื่อ $r(t)$ คือ โครงร่างของตัวอักษรแต่ละตัวที่อยู่ในรูปของ Temporal domain

$X(t)$, $Y(t)$ คือ โครงร่างของตัวอักษรแต่ละตัวที่อยู่ในรูปของ Spatial domain

cg_x , cg_y คือ ตำแหน่งศูนย์กลางของตัวอักษร

- นำเอาโครงร่างของตัวอักษรทั้งหมดที่อยู่ในรูปของ Temporal Domain มาเข้าฟังก์ชัน Wavelet Descriptor (Chey, C. et al. 2006) เพื่อหาลักษณะเด่นของแต่ละตัวอักษรซึ่งอยู่ในรูปของสัมประสิทธิ์เวฟเล็ต ที่สามารถหาได้ทั้งแบบ global และ local และเก็บเป็นแม่แบบ (template) ของแต่ละตัวอักษร

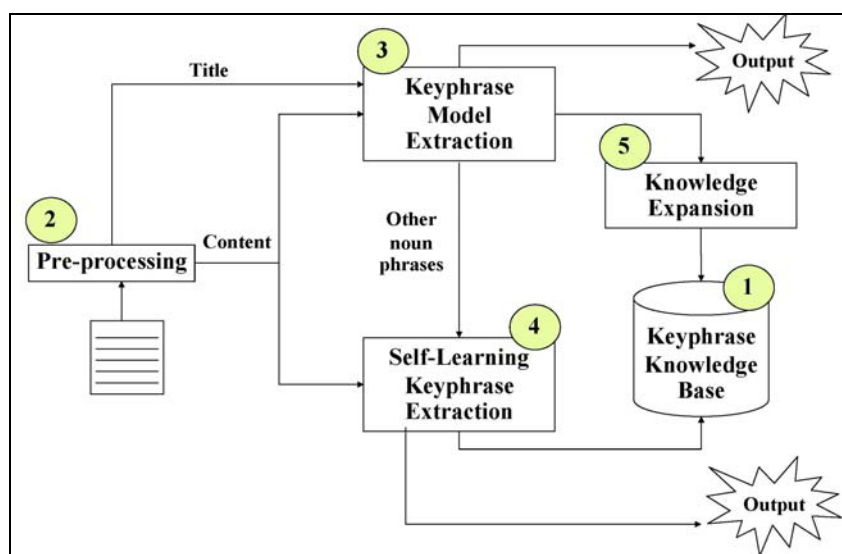
(3) การเปรียบเทียบ (Matching Process)

ขั้นตอนนี้จะเป็นการรู้จำตัวอักษรที่ใช้ในการทดสอบ ซึ่งผ่านการเตรียมข้อมูล และ ดึงลักษณะเด่นของตัวอักษร เหมือนกับตัวอักษรที่ใช้ฝึกฝน จะถูกนำมาเปรียบเทียบกับ

แม่แบบที่ผ่านการฝึกฝนมาก่อนหน้านี้ การเปรียบเทียบใช้หลักการของ Minimum Euclidean Distance แม่แบบของตัวอักษรที่ให้ค่าระยะทางน้อยที่สุด คือคำตอบของตัวอักษรนั้นด้วยวิธีการรู้จำตัวอักษรที่ใช้ในงานวิจัยนี้ สามารถให้ความถูกต้องได้สูงสุดถึง 92.99%

(3) การค้นหาคำสำคัญโดยใช้ความหมาย (Semantic-based Keyword Extraction)

ขั้นตอนนี้อถือว่าเป็นหัวใจของงานวิจัยนี้ ซึ่งเสนอการค้นหาคำสำคัญจากตัวเอกสารเพียงอย่างเดียว ไม่มีการใช้คลังข้อมูลสำหรับฝึกฝน (Kongkachandra, R., et al.: 2006) ดังนั้นความหมายของเอกสารที่งานวิจัยนี้กำหนด จะมาจากประโยคต่างๆที่มี คำสำคัญ ปรากฏอยู่ คำสำคัญจะถูกดึงออกมาจากเอกสารด้วยการเปรียบเทียบระหว่างความหมายของ คำที่พิจารณา (Candidate) กับ คำที่เป็นตัวแทนของเอกสารซึ่งจะได้มาจากการพิจารณาจากหัวเรื่อง ดังนั้น ในงานวิจัยนี้จะมีข้อจำกัดตรงที่เอกสารที่นำมาใช้สำหรับค้นหาคำสำคัญจะต้องมีส่วนของหัวเรื่องที่สื่อความหมายของเอกสารทั้งหมดได้ รูปที่ 3.3 แสดงถึงแผนผังรวมของการค้นหาคำสำคัญโดยใช้ความหมาย



รูปที่ 3.3 แผนผังรวมของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย

ส่วนประกอบต่างของระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมาย ที่แสดงในรูปที่ 3.3 มีดังนี้

1. ฐานความรู้ของคำสำคัญ (Keyword Knowledge Base)

ฐานความรู้ของคำสำคัญ ถูกออกแบบขึ้นมาเพื่อใช้สำหรับเก็บคำสำคัญ ความหมาย และกลุ่มของเอกสาร (Domain) ที่คำสำคัญนั้นสังกัดอยู่ ฐานความรู้ของคำสำคัญที่ใช้ในงานวิจัยนี้เป็น

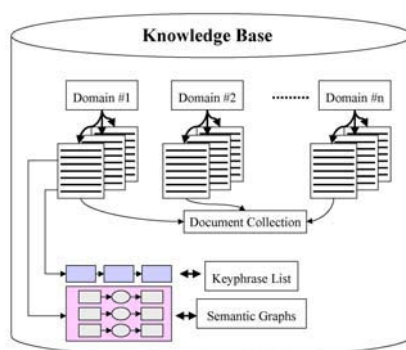
แบบที่ไม่เฉพาะเจาะจงกับกลุ่มของเอกสารกลุ่มใดกลุ่มหนึ่ง และสามารถเพิ่มเติมได้แบบอัตโนมัติ แต่อย่างไรก็ตาม ตอนที่เริ่มต้นสร้างฐานความรู้ของคำสำคัญนั้น ระบบจำเป็นต้องให้ผู้ใช้เป็นผู้กำหนดสิ่งต่างๆ เหล่านี้

(1) จำนวนกลุ่มของเอกสาร (Number of Domains)

(2) คำศัพท์ควบคุม (Controlled Vocabulary) ของแต่ละกลุ่ม ประมาณ 8-10 คำศัพท์ต่อกลุ่ม

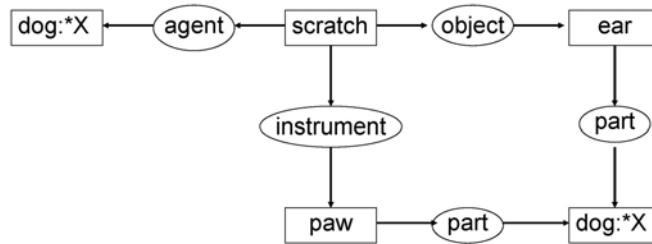
(3) โครงข่ายทางความหมายของกลุ่มของเอกสาร (Domain Hierarchy)

เมื่อระบบอัตโนมัติค้นหาคำสำคัญเชิงความหมายทำงานในแต่ละรอบ ฐานความรู้ของคำสำคัญจะถูกอัปเดตด้วยคำสำคัญ 2 ประเภทด้วยกัน ชนิดแรกคือ คำสำคัญเริ่มต้น (Initial keyword) ที่ได้จากการพิจารณาที่หัวเรื่องของเอกสาร ซึ่งได้มาจากระบบการดึงคำสำคัญเริ่มต้น (Initial Keyword Extraction) และ ชนิดที่สองคือ คำสำคัญที่มีความหมายเกี่ยวข้อง (Relative keyword) ที่ได้จากการพิจารณาที่เนื้อหาของเอกสาร ซึ่งได้มาจากระบบการดึงคำสำคัญที่เกี่ยวข้องเชิงความหมาย (Semantic-Relatedness Keyword Extraction)



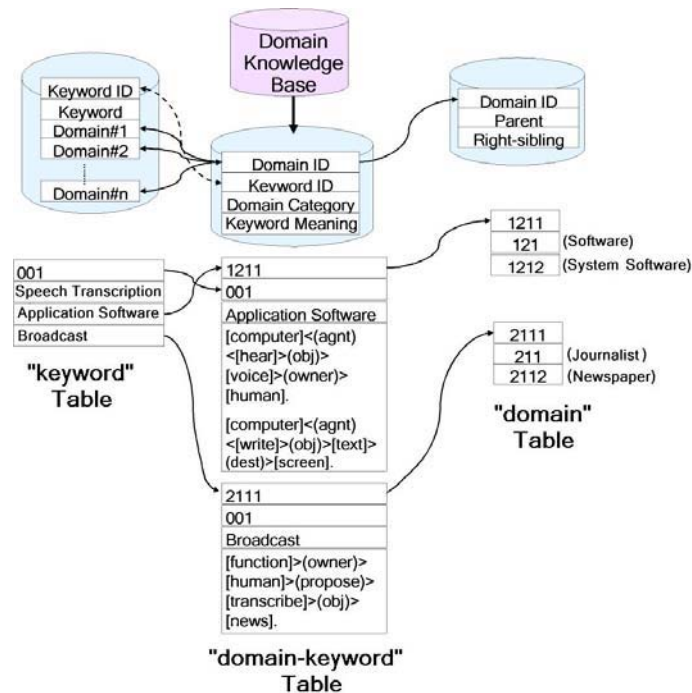
รูปที่ 3.4 แสดงถึงการจัดเก็บข้อมูลภายในฐานความรู้ของคำสำคัญ

รูปที่ 3.4 แสดงถึงการจัดเก็บข้อมูลภายในฐานความรู้ของคำสำคัญ จะเห็นว่าภายในฐานความรู้ของคำสำคัญจะแบ่งออกเป็น กลุ่มของเอกสาร ซึ่งภายในแต่ละกลุ่มจะรวบรวมเอกสารที่เกี่ยวข้องกับกลุ่มนั้นๆเข้าด้วยกัน และภายในแต่ละเอกสารจะเก็บ รายการคำสำคัญ ที่ได้จากการทำงานของระบบอัตโนมัติค้นหาคำสำคัญก่อนหน้านี้ พร้อมกันนี้ แต่ละคำสำคัญจะเก็บความหมายเอาไว้ด้วย ในงานวิจัยนี้ ความหมายของคำสำคัญ จะถูกเก็บในรูปแบบของ กราฟเชิงความคิด (Conceptual Graph) ซึ่งถูกกำหนดขึ้นเมื่อปี ค.ศ. 1984 โดยนักคณิตศาสตร์ชื่อ จอห์น โซวะ (Sowa, J.F.: 1984) กราฟเชิงความคิด เป็นกราฟชนิด Bipartite ที่ประกอบไปด้วยโนด 2 ประเภทคือ Conceptual Entity และ Conceptual Relation รูปที่ 3.5 แสดงถึงการใช้กราฟเชิงความคิดแทนความหมายของประโยค “The dog scratches its ear with its paw”

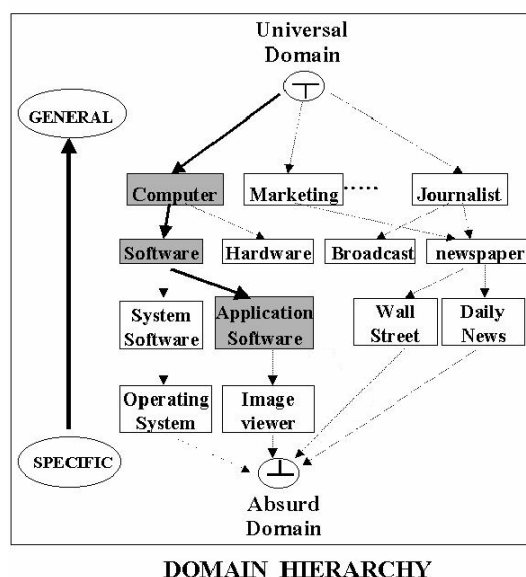


รูปที่ 3.5 การใช้กราฟเชิงความคิดแทนความหมายของประโยค “The dog scratches its ear with its paw”

ข้อมูลที่อยู่ในกรอบสี่เหลี่ยมเรียกว่า Concept Node ซึ่งจะเก็บคำ โดยที่คำเหล่านี้อาจจะเป็น คำนาม คำกริยา หรือ คำคุณศัพท์ ส่วนข้อมูลที่อยู่ในวงกลมเรียกว่า Conceptual Relation Node ซึ่งจะแสดงความสัมพันธ์ระหว่าง Concept Node ที่เชื่อมกัน ความหมายของกราฟนี้ คือ “สุนัข (dog) เป็นประธาน (agent) ของการเกา (scratch) ซึ่งกรรม (object) ของการเกา (scratch) คือหู (ear) เครื่องมือ (instrument) ที่ใช้ในการเกา (scratch) คือเล็บ (paw) และ หูก็เป็นส่วนหนึ่ง (part) ของสุนัข”



(ก)



(ข)

รูปที่ 3.6 แสดงถึงการจัดเก็บข้อมูลภายในฐานความรู้สำคัญ

รูปที่ 3.4 แสดงถึงโครงร่างโดยคร่าวของฐานความรู้ของคำสำคัญ รายละเอียดในการจัดเก็บภายในฐานความรู้ของคำสำคัญ ดังแสดงในรูปที่ 3.6 ซึ่งประกอบไปด้วย ตารางความสัมพันธ์ 3 ตารางคือ ตารางคำสำคัญ (Keyword Table) ตารางกลุ่มของเอกสาร-คำสำคัญ (Domain-Keyword Table) และ ตารางกลุ่มของเอกสาร (Domain Table) ซึ่งสอดคล้องกับโครงข่ายของกลุ่มเอกสารในรูปที่ 3.6 (ข) ตัวอย่างคือ ในตารางคำสำคัญทางซ้ายมือสุดของรูปที่ 3.6 ก. คำสำคัญคือ “Speech Technology” (มีหมายเลข 001) ซึ่งสามารถปรากฏอยู่ในกลุ่มเอกสารได้ 2 ด้านคือ “Application Software” และ “Broadcast” ซึ่งมีลิงค์เชื่อมไปที่ ตารางกลุ่มของเอกสาร-คำสำคัญที่อยู่ตรงกลาง เนื่องจากการวิจัยนี้เสนอการค้นหาคำสำคัญได้ในหลายกลุ่มของเอกสาร ดังนั้นเราจำเป็นต้องจัดระเบียบความสัมพันธ์ของกลุ่มต่างๆเอาไว้ ซึ่งเป็นไปตาม Semantic Web ที่ชื่อว่า WordNet (Teich, E. and Fankhauser, P.: 2004) และ (Yang, D., and Powers, D.: 2005) ดังแสดงเป็นบางส่วนในรูปที่ 3.6 (ข)

2. ส่วนเตรียมข้อมูล (Preprocessing)

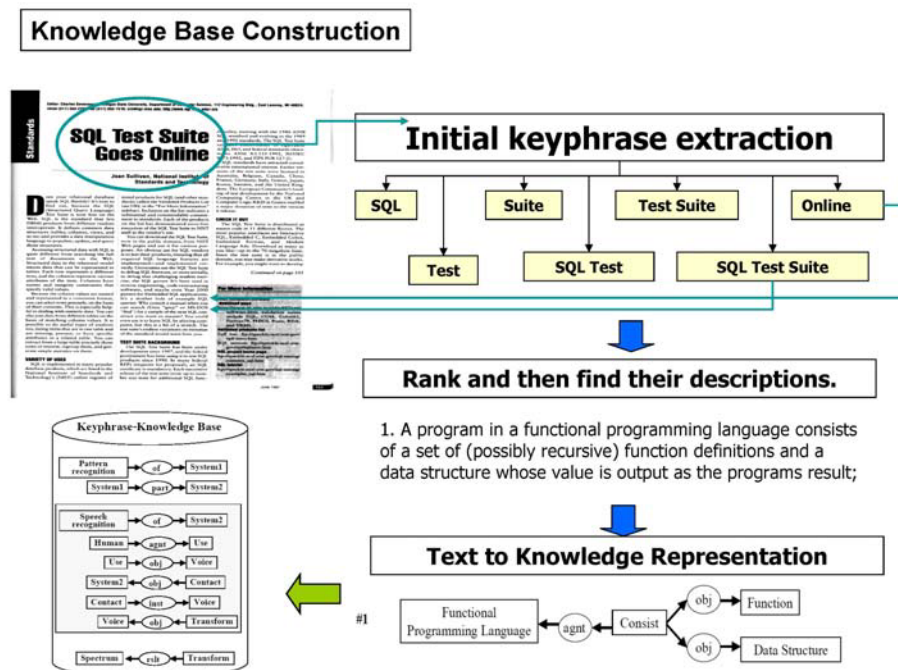
จากที่กล่าวไว้ว่า ฐานความรู้ของคำสำคัญ จะถูกอัพเดทด้วย คำสำคัญ 2 ประเภท ซึ่งได้มาจากการพิจารณาเอกสารคนละส่วน นั่นคือ ส่วนของหัวเรื่อง และ ส่วนของเนื้อความ ดังนั้น ในขั้นตอนของการเตรียมข้อมูล จะดำเนินการดังนี้

- (1) แยกส่วนของหัวเรื่องออกจากเนื้อความ โดยใช้กฎที่ระบุว่า ตำแหน่งของประโยคในบรรทัดแรกของเอกสารจะเป็นหัวเรื่องของเอกสารนั้น

- (2) ในส่วนของเนื้อความ ระบบจะแบ่งเอกสารออกเป็น หนึ่งประโยค/หนึ่งบรรทัด ซึ่งเกณฑ์ที่ใช้ในการตัดประโยคคือพิจารณาที่ ตัวแบ่งประโยค (Punctuation) เช่น จุลภาค เครื่องหมายหยุด เครื่องหมายคำถาม เป็นต้น

3. ส่วนของการดึงคำสำคัญเริ่มต้น (Initial Keyword Extraction)

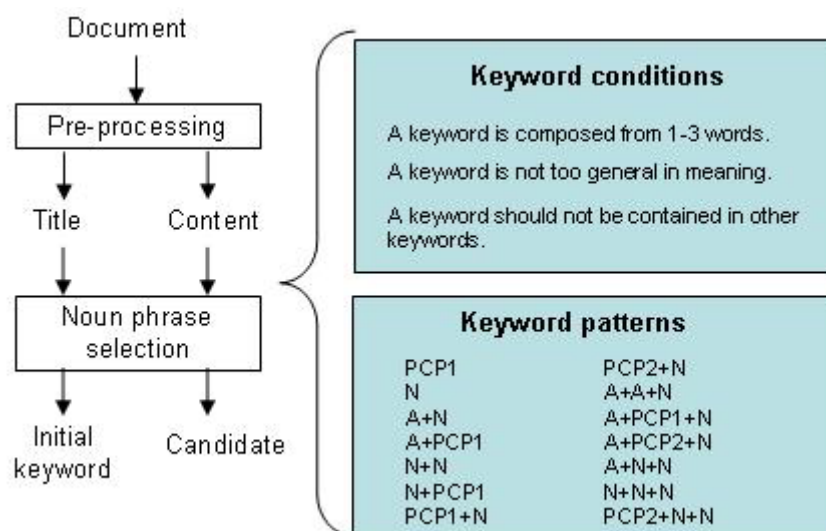
งานวิจัยนี้เสนอการดึงความสำคัญจากเอกสารโดยการเปรียบเทียบกับความหมายของเอกสาร แต่เนื่องจากงานวิจัยนี้เสนอที่จะพิจารณาเพียงเอกสารนั้นๆ ในการดึงคำสำคัญออกมา ดังนั้นเราจำเป็นต้องเตรียมข้อมูลที่เป็นความหมายของเอกสาร M. Montes-y Gme และคณะ (M. Montes-y Gmez et al., 1999) ได้วิจัยไว้ว่า หัวเรื่องของเอกสาร ถือเป็นแหล่งอธิบายความหมายที่ดีของเอกสาร ดังนั้น หัวเรื่องของเอกสารจึงถูกใช้เป็นแหล่งข้อมูลที่จะใช้ในการดึงคำสำคัญเริ่มต้น โดยมีหลักการในการดำเนินการ ดังแสดงในรูป 3.7



รูปที่ 3.7 หลักการของการดึงคำสำคัญเริ่มต้น (Initial Keyword Extraction)

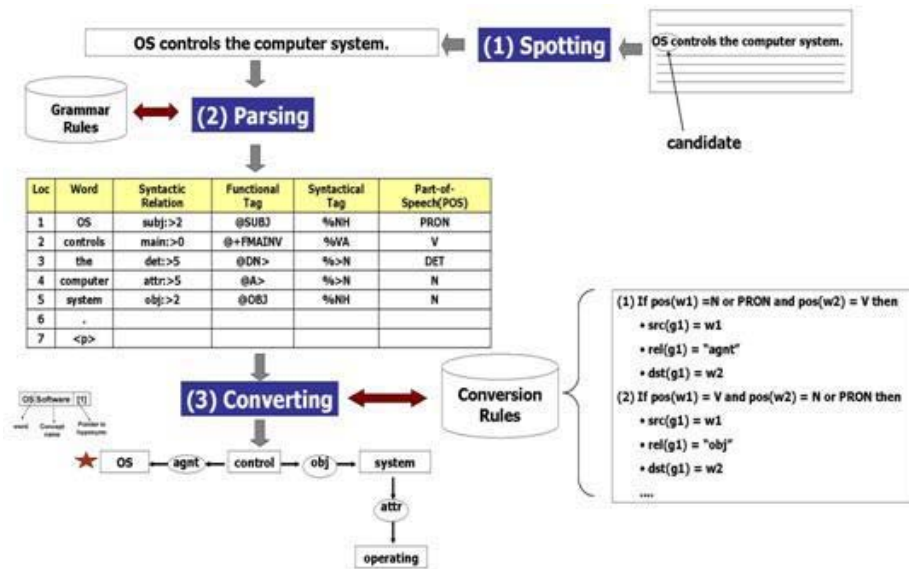
- (1) หัวเรื่องของเอกสารจะถูกส่งไปยังขั้นตอนคัดเลือกคำคู่แข่ง (Candidate Selecting) เพื่อค้นหาคำสำคัญเริ่มต้น ในขั้นตอนนี้ หัวเรื่องของเอกสารจะถูกแจกแจงไวยากรณ์ เพื่อให้ทราบหน้าที่ของคำ (Part-of-speech) จากนั้นจะคัดเลือกเฉพาะคำที่ไม่อยู่ในรายการของคำหยุด (Stop-word list) และมีบทบาทเป็น นามวลี เท่านั้น โดยเปรียบเทียบกับรูปแบบของนามวลี ดังแสดงในรูปที่ 3.8 และรายการของคำหยุดที่กำหนดไว้อย่างไรก็ตาม จากหัวเรื่องเดียวกัน เราสามารถดึงคำสำคัญเริ่มต้น ได้หลายคำ ซึ่งอาจจะมีส่วนที่ซ้อนทับกันอยู่ เช่น

“SQL”, “Test”, “Suite”, “Test Suite”, “SQL Test” และ “SQL Test Suite” (ดังตัวอย่างในรูป 3.7) งานวิจัยนี้ จะเลือกนามวลีที่ประกอบด้วยคำเดี่ยวตั้งแต่ 2 คำขึ้นไป ดังนั้น “Test Suite”, “SQL Test” และ “SQL Test Suite” จะถูกเลือกเป็นคำสำคัญเริ่มต้น สำหรับ “Online” นั้นก็จะถูกเลือกด้วย เนื่องจากไม่ซ้อนทับกับคำใด



รูปที่ 3.8 วิธีการสำหรับคัดเลือกนามวลี

จากข้อ (1) เราจะได้คำสำคัญเริ่มต้นออกมาจากเอกสาร ในขั้นตอนนี้ จะดึงเอาความหมายของแต่ละคำสำคัญออกมา แล้วจัดเก็บใน ฐานความรู้คำสำคัญ เพื่อใช้งานในระยะต่อไป ความหมายของคำสำคัญ ที่ระบุในงานวิจัยนี้ คือ ประโยคต่างๆ ที่อยู่ในเอกสาร ที่มีคำสำคัญปรากฏอยู่ ดังนั้น ส่วนของเนื้อหาเอกสารจะถูกส่งมาเพื่อเข้าสู่ขั้นตอนการสร้างฐานความรู้ (Knowledge Generating) โดยใช้ คำสำคัญ ที่ได้จากข้อที่ (1) เป็นตัวชี้ตำแหน่งของประโยคเหล่านั้น จากนั้น ประโยคที่ดึงออกมาจะถูกนำมาแจกแจงไวยากรณ์ (Parsing) และแปลงให้อยู่ในรูปของกราฟเชิงความคิดโดยกฎในการแปลงที่กำหนด (Conversion Rules) ซึ่งกราฟเชิงความคิดที่ได้นี้ ต่อไปเรียกว่า **กราฟความหมาย (Semantic Graph)** วิธีการในการแปลงจากประโยคภาษาอังกฤษ เป็น กราฟความหมาย ดังแสดงในรูปที่ 3.9



รูปที่ 3.9 วิธีการในการแปลงจากประโยคภาษาอังกฤษ เป็น กราฟความหมาย

- (2) คำสำคัญที่ได้ พร้อมกับกราฟความหมายที่ได้จาก (2) และ (3) จะถูก ส่วนของการขยายฐานความรู้คำสำคัญ (Knowledge Expansion) เพื่อเพิ่มเติมเข้าไปในฐานความรู้ของคำสำคัญ

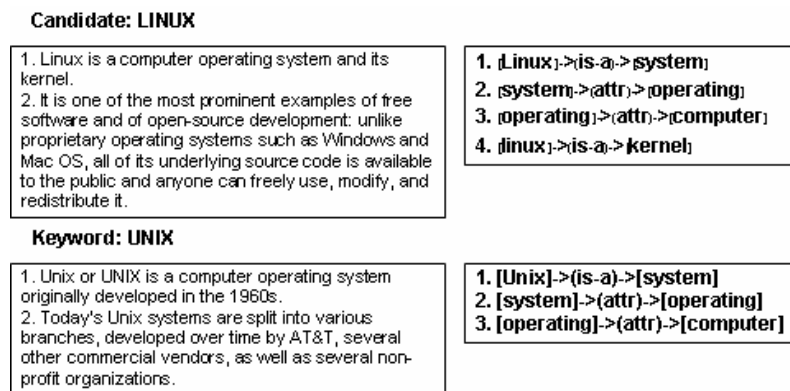
4. ส่วนของการดึงคำสำคัญที่เกี่ยวข้องเชิงความหมาย (Semantic-Relatedness Keyword Extraction)

ฐานความรู้ของคำสำคัญที่ได้จากส่วนของการดึงคำสำคัญเริ่มต้น มีจำนวนไม่มากพอที่จะใช้ในการประยุกต์ใช้งาน ดังนั้น ขั้นตอนในส่วนนี้เป็นส่วนที่ดึงคำสำคัญเพิ่มเติม โดยพิจารณาจากคำคู่แข่งที่มีความหมายเกี่ยวข้องกับคำสำคัญเริ่มต้นที่อยู่ในฐานความรู้คำสำคัญ ขั้นตอนนี้จะเรียนรู้ด้วยตนเองด้วยการเปรียบเทียบความหมายของคำคู่แข่งแต่ละคำกับคำสำคัญเริ่มต้น คำคู่แข่งใดที่มีคะแนนที่ได้จากการวัดความเหมือนกันทางความหมายกับคำสำคัญเริ่มต้น แล้วมีค่าสูงกว่า ค่าเงื่อนไข (Threshold) ที่กำหนด เราจะถือว่าคำนั้นเป็นคำสำคัญด้วย เรียกว่า คำสำคัญที่เกี่ยวข้องเชิงความหมาย (Semantic-Related Keyword)

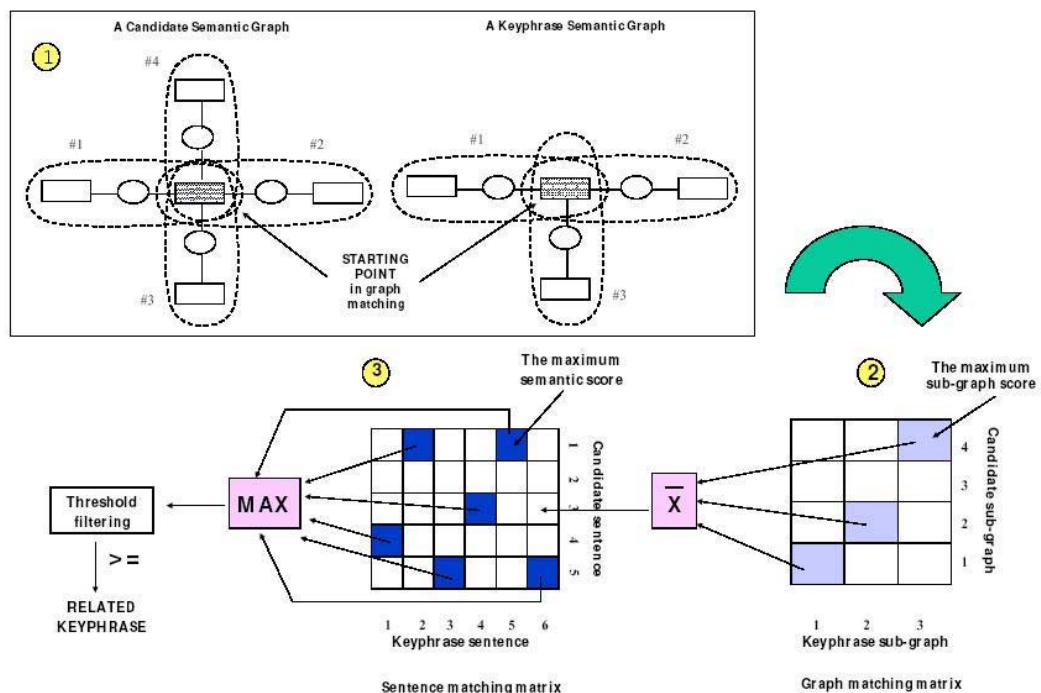
ขั้นตอนการทำงานใน 2 ขั้นตอนแรก ของส่วนนี้จะคล้ายกับของส่วนของการดึงความสำคัญเริ่มต้น มีที่แตกต่างกันคือในขั้นตอนที่ (1) ดังนี้

- (1) เปลี่ยนจากหัวเรื่องของเอกสารเป็นเนื้อความของเอกสาร จะถูกส่งไปยังขั้นตอนคัดเลือกคำคู่แข่ง (Candidate Selecting) เพื่อค้นหาคำวลี ซึ่งคำวลีใดมีการซ้อนทับกันแล้ว คำวลีที่มีจำนวนคำเดียวมากกว่า 2 คำขึ้นไปจะถูกเลือกไว้พิจารณา
- (2) ทำในลักษณะเดียวกันกับส่วนของการดึงคำสำคัญเริ่มต้น

- (3) คำคู่แข่งพร้อมความหมายทั้งหมดที่ได้จากขั้นตอนข้างบนจะถูกส่งเข้าไปในขั้นตอนการวัดความเหมือนกันทางความหมาย (Semantic Similarity Scoring) คำคู่แข่งใดมีคะแนนสูงกว่าค่าเงื่อนไข (Threshold) ที่กำหนด เราจะถือว่าคำนั้นเป็นคำสำคัญด้วย เรียกว่า คำสำคัญที่เกี่ยวข้องเชิงความหมาย (Semantic-Related Keyword) วิธีการในการวัดความเหมือนกันทางความหมาย แสดงในรูปที่ 3.10 และ 3.11



รูปที่ 3.10 กราฟความหมายของคำสำคัญ “UNIX” และ คำคู่แข่ง “LINUX”

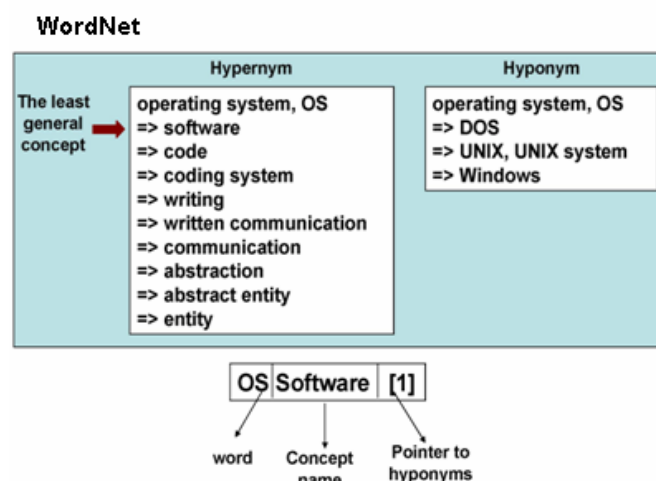


รูปที่ 3.11 การวัดความเกี่ยวข้องทางความหมาย

ในรูปที่ 3.10 คำสำคัญ คือ “UNIX” และมีประโยคที่มีคำว่า “UNIX” ปรากฏอยู่ 2 ประโยค ซึ่งเมื่อผ่านขั้นตอนการแปลงเป็นกราฟความหมายแล้ว จะได้กราฟความหมายทั้งหมด 3 กราฟ ดังแสดงทางขวามือของรูป ส่วนคำคู่แข่ง ที่นำมาพิจารณา คือคำว่า “LINUX” ซึ่งมีประโยคที่คำนี้

ปรากฏอยู่ 2 ประโยค และเมื่อผ่านขั้นตอนการแปลงเป็นกราฟความหมายแล้ว จะได้กราฟความหมายทั้งหมด 4 กราฟ

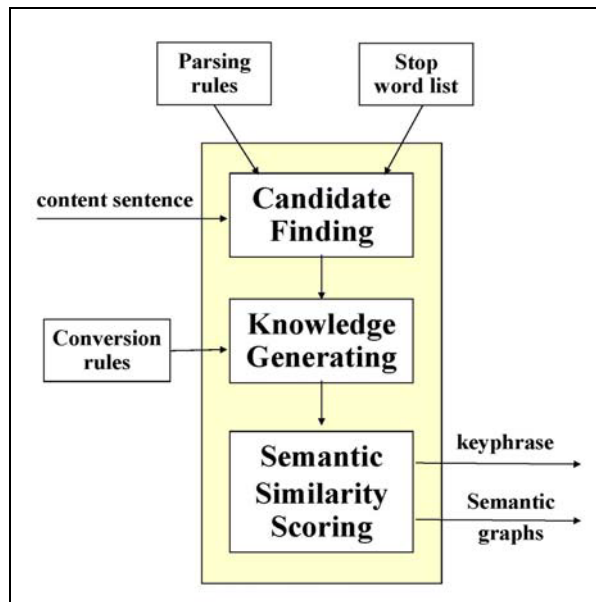
กราฟความหมายของคำสำคัญ 3 กราฟ และ กราฟความหมายของคำคู่แข่ง 4 กราฟ จะถูกนำมาเปรียบเทียบดังแสดงในรูปที่ 3.11 การเปรียบเทียบเริ่มที่ โหนด ที่มีจำนวนลิงค์มากที่สุด ซึ่งเรียกว่า Core Node ของทั้งสองฝ่าย การเปรียบเทียบจะทำให้ละกราฟย่อย ซึ่งประกอบไปด้วย 2 Concept node และ 1 Conceptual Relation node ซึ่งข้อมูลในโหนดประกอบด้วย คำ (word) ชื่อกลุ่มของคำ (concept name) และ ลิงค์ไปยังกลุ่มของคำที่เฉพาะเจาะจงกว่า (Pointer to hyponyms) ตัวอย่างของการเก็บข้อมูลในแต่ละโหนด ดังแสดงในรูปที่ 3.12 โหนดที่แสดงเป็นตัวอย่างเก็บคำว่า “OS” ซึ่งมี ชื่อกลุ่มของคำเป็น “Software” และ ลิงค์ไปยังกลุ่มของคำที่เฉพาะเจาะจงกว่า ซึ่งได้แก่ “DOS”, “UNIX” และ “Windows”



รูปที่ 3.12 การเก็บค่าข้อมูลในแต่ละโหนดของกราฟความหมาย

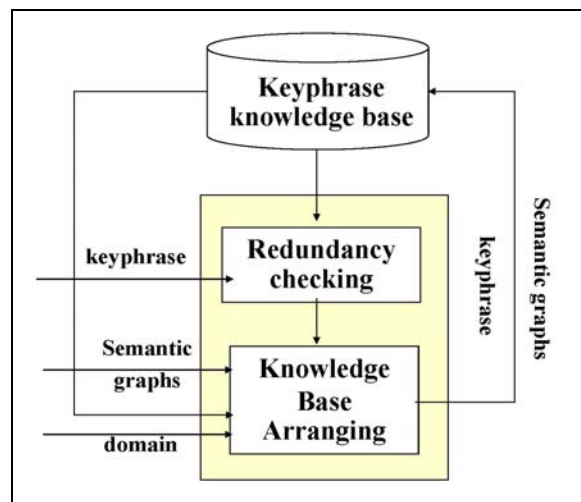
ในการเปรียบเทียบแต่ละกราฟย่อยจะได้ กราฟย่อยที่มีคะแนนสูงสุด ซึ่งจะเก็บไว้ใน เมตริกซ์ของการเปรียบเทียบกราฟ (Graph Matching Matrix) จากนั้น คะแนนของกราฟย่อยทั้งหมดของแต่ละประโยค จะถูกนำมาหาค่าเฉลี่ย และเก็บไว้ใน เมตริกซ์ของการเปรียบเทียบประโยค (Sentence Matching Matrix) คำคู่แข่งใดที่มีคะแนนของการเปรียบเทียบประโยคที่มากที่สุด จะถูกนำมาเปรียบเทียบกับค่าเงื่อนไข ซึ่งค่าที่มีค่าสูงกว่าค่าเงื่อนไข ก็จะถูกยอมรับว่าเป็นคำสำคัญที่เกี่ยวข้องทางความหมาย

- (4) คำสำคัญที่ได้ พร้อมกับกราฟความหมายที่ได้จาก (3) จะถูกส่งไปยัง ส่วนของการขยายฐานความรู้คำสำคัญ (Knowledge Expansion) เพื่อเพิ่มเติมเข้าไปในฐานความรู้ของคำสำคัญ รูปที่ 3.13 แสดงถึงขั้นตอนการดึงคำสำคัญที่มีความเกี่ยวข้องทางความหมาย



รูปที่ 3.13 แสดงถึงขั้นตอนการดึงคำสำคัญที่มีความเกี่ยวข้องทางความหมาย

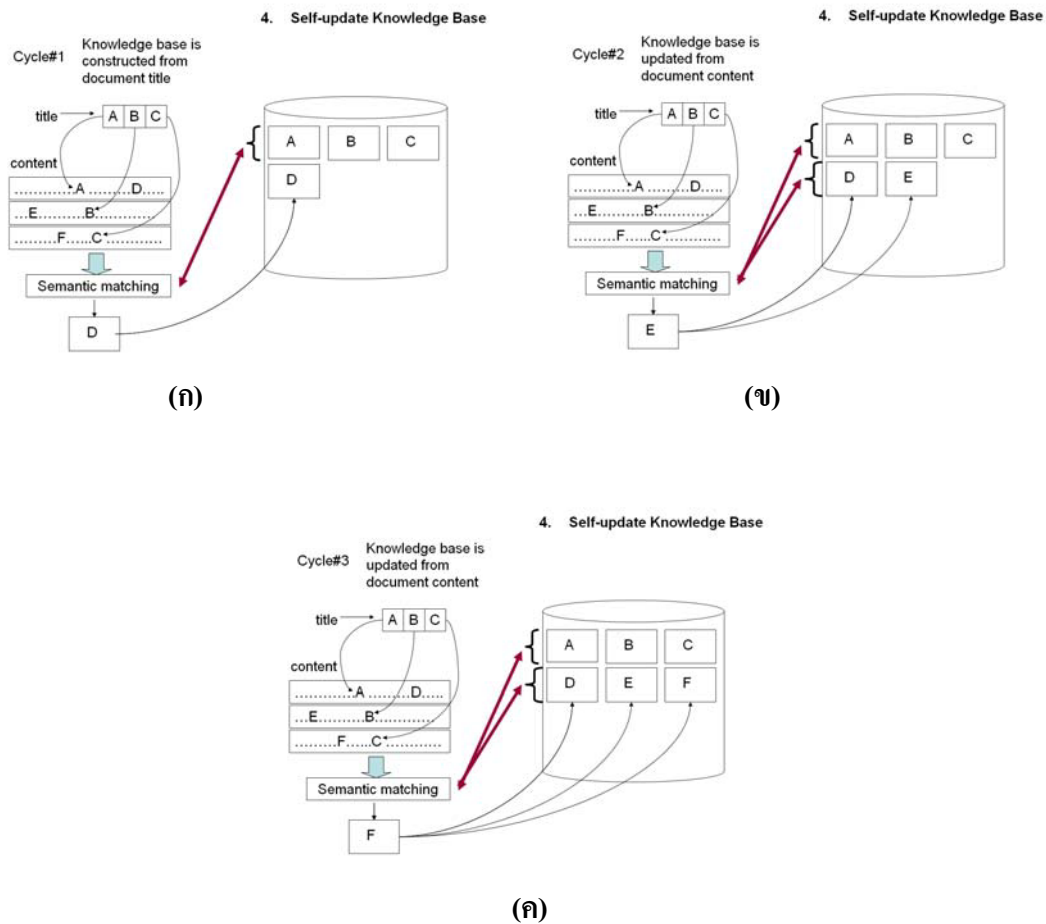
5. ส่วนของการขยายฐานความรู้คำสำคัญ (Knowledge Expansion)



รูปที่ 3.14 แสดงถึงขั้นตอนการขยายฐานความรู้คำสำคัญ

รูปที่ 3.14 แสดงถึงขั้นตอนการขยายฐานความรู้คำสำคัญซึ่งเป็นจุดเด่นของงานวิจัยนี้ นั่นคือ การใช้ฐานความรู้คำสำคัญที่ถูกสร้างและปรับปรุงให้ทันสมัยอย่างอัตโนมัติโดยไม่ใช่คลังข้อมูลขนาดใหญ่ ภายในขั้นตอนนี้ คำสำคัญ และความหมายที่ได้มาจากการดึงคำสำคัญเริ่มต้น และส่วนของการดึงคำสำคัญที่เกี่ยวข้องเชิงความหมาย จะถูกส่งเข้าไปตรวจสอบใน Redundancy checking เพื่อพิจารณาว่าคำสำคัญเคยถูกกำหนดไว้ในฐานความรู้ของคำสำคัญหรือไม่ ถ้ายังไม่เคยถูกกำหนดมาก่อน ทั้งคำสำคัญและกราฟความหมาย จะถูกเพิ่มเข้าไปในฐานความรู้ของคำสำคัญ

โดย Knowledge Base Arranging ซึ่งจะปรับปรุงตารางทั้ง 3 ตารางของฐานความรู้ แต่ถ้าเคยถูกกำหนดไว้แล้ว ระบบจะเพิ่มเฉพาะส่วนของความหมายเข้าไป โดยจะปรับปรุงเฉพาะ ตารางคำสำคัญ (Keyword Table) เท่านั้น



รูปที่ 3.15 ตัวอย่างของการปรับปรุงฐานความรู้คำสำคัญ

ตัวอย่างการปรับปรุงฐานความรู้ดังแสดงในรูป 3.15 (ก) - (ค) โดยเริ่มต้นฐานความรู้จะถูกสร้างจาก คำสำคัญเริ่มต้นที่ได้จากส่วนของหัวเรื่อง ในตัวอย่างคือ A, B และ C ดังแสดงในรูป 3.15 (ก) จากนั้นระบบเข้าสู่ส่วนของการดึงคำสำคัญที่เกี่ยวข้องทางความหมายด้วยการพิจารณาส่วนของ เนื้อความเอกสาร ซึ่งจะได้อ D, E และ F ดังแสดงในรูปที่ 3.15 (ก)-(ค) ตามลำดับ

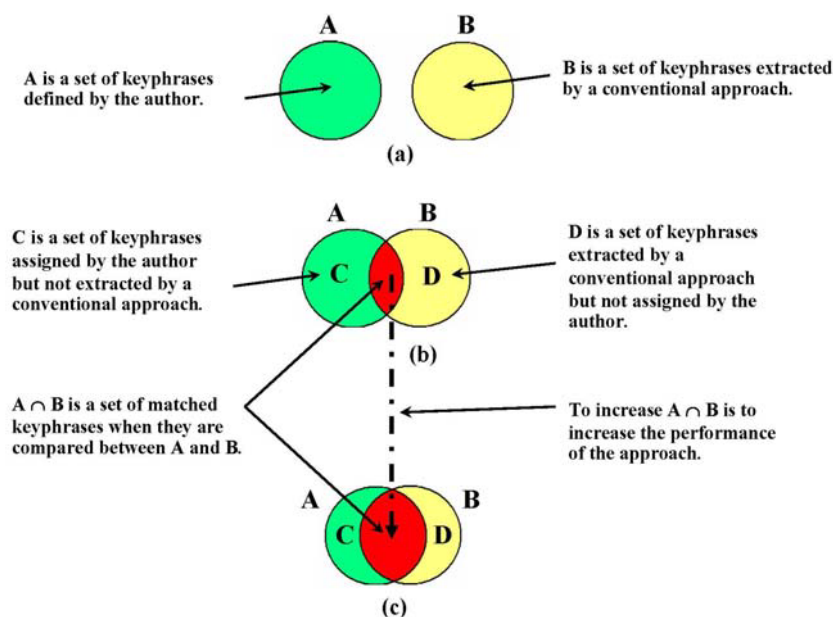
บทที่ 4 ขั้นตอนการวิจัยและผลการวิจัย

งานวิจัยนี้ได้นำเสนอหัวข้อในการวิจัยไว้ 3 ส่วนด้วยกันคือ

1. งานวิจัยนี้สามารถค้นหาคำสำคัญได้ โดยไม่จำเป็นต้องใช้คลังข้อมูลอื่นๆ ในการค้นหาคำสำคัญ นอกจากเอกสารนั้นๆ และ ฐานความรู้ของคำสำคัญที่ได้สร้างขึ้น
2. งานวิจัยนี้สามารถค้นหาคำสำคัญที่มีความหมายสอดคล้องกับความหมายของเอกสารได้ดีกว่าวิธีอื่น
3. งานวิจัยนี้เสนอการนำเอาฐานความรู้มาเพิ่มประสิทธิภาพในการค้นหาคำสำคัญ

ดังนั้น เพื่อพิสูจน์ประสิทธิภาพของงานวิจัย ในส่วนต่อไปนี้จะกล่าวถึง วิธีการประเมินประสิทธิภาพของผลงานวิจัยที่ได้, การตั้งสมมุติฐานและออกแบบการวิจัย, ผลการทำวิจัย และวิเคราะห์ผลการวิจัย

4.1 วิธีการประเมินประสิทธิภาพ



รูปที่ 4.1 หลักในการประเมินประสิทธิภาพของระบบ

รูปที่ 4.1 ประกอบไปด้วย วงกลม A และ B ซึ่งแสดงถึงเซตของคำสำคัญที่กำหนดโดยผู้เขียนและคำสำคัญที่ได้จากระบบอัตโนมัติค้นหาคำสำคัญ ตามลำดับ ในทางอุดมคติ ระบบอัตโนมัติค้นหาคำสำคัญจะต้องผลิตเซตของคำสำคัญ B ที่มีขนาดเท่ากับ เซตของคำสำคัญ A แต่ในทางปฏิบัติระบบ

อัตโนมัติสำหรับดึงคำสำคัญที่มีอยู่ ผลที่ได้เฉพาะส่วนที่เป็น $A \cap B$ ยังมีส่วนที่เป็นข้อผิดพลาดอยู่ 2 แบบด้วยกันคือ

1. ข้อผิดพลาดที่สำคัญที่อยู่เขตของคำสำคัญที่กำหนดโดยผู้เขียน แต่ไม่ถูกดึงออกมาด้วยระบบอัตโนมัติค้นหาคำสำคัญ ได้แก่ เขตของ C
2. ข้อผิดพลาดที่สำคัญที่ดึงออกมาด้วยระบบอัตโนมัติค้นหาคำสำคัญ แต่ไม่ได้ถูกกำหนดไว้โดยผู้เขียน ได้แก่ เขตของ D

ในการเพิ่มประสิทธิภาพของระบบอัตโนมัติค้นหาคำสำคัญ สามารถทำได้ด้วยการลดจำนวนข้อผิดพลาดทั้งสอง นั่นคือ เพิ่มขนาดของ $A \cap B$ สำหรับงานวิจัยนี้ จะเพิ่มประสิทธิภาพของระบบด้วยการลดขนาดของ เขต C และใช้ F-measure ดังแสดงในสมการ (6.3) เป็นตัววัดประสิทธิภาพของระบบ โดยใช้คำสำคัญที่ผู้เขียนกำหนดไว้ในแต่ละเอกสารเป็นมาตรฐานในการเปรียบเทียบ

$$Precision = \frac{A \cap B}{B} \quad (6.1)$$

$$Recall = \frac{A \cap B}{A} \quad (6.2)$$

$$F - measure = \frac{2 \times (Precision \times Recall)}{(Precision + Recall)} \quad (6.3)$$

4.2 การตั้งสมมติฐานและออกแบบการวิจัย

ตารางที่ 4.1 การกำหนดค่าต่างๆสำหรับงานวิจัย

รายการที่กำหนด	ข้อมูลที่ใช้
ระบบอัตโนมัติค้นหาคำสำคัญที่มีอยู่	1. EXTRACTOR 2. KEA 3. Matsuo's System
จำนวนของกลุ่มเอกสาร	5 กลุ่ม: ธุรกิจ , คอมพิวเตอร์, การศึกษา, การแพทย์ และ จิตวิทยา
จำนวนเอกสารที่ใช้ทดสอบ	100 เอกสาร
จำนวนคำสำคัญที่ต้องการ	5, 10 and 15 คำสำคัญ
เครื่องมือทางภาษา	Machine Syntax
ระดับเงื่อนไขในการพิจารณาคำสำคัญ	0.5, 0.8 และ 1.0

4.3 ผลการทำวิจัย

ในงานวิจัยนี้ ได้ออกแบบการทดลองเพื่อประเมินประสิทธิภาพของระบบใน 3 แนวทางด้วยกัน คือ

1. งานวิจัยนี้สามารถค้นหาคำสำคัญได้ โดยไม่จำเป็นต้องใช้คลังข้อมูลอื่นๆ ในการค้นหาคำสำคัญ นอกจากเอกสารนั้นๆ และ ฐานความรู้ของคำสำคัญที่ได้สร้างขึ้น โดยการเปรียบเทียบผลการวิจัยกับ งานวิจัยเดิมที่ใช้ คลังข้อมูล ได้แก่ EXTRACTOR [] และ KEA []
2. งานวิจัยนี้สามารถค้นหาคำสำคัญที่มีความหมายสอดคล้องกับความหมายของเอกสาร ได้ดีกว่าวิธีอื่น โดยการเปรียบเทียบผลการวิจัยกับ งานวิจัยเดิมที่ใช้ คลังความรู้และไม่ใช้คลังความรู้ ได้แก่ EXTRACTOR, KEA และ ระบบค้นหาคำสำคัญของ Matsuo []
3. งานวิจัยนี้เสนอการนำเอาฐานความรู้มาเพิ่มประสิทธิภาพในการค้นหาคำสำคัญ โดยการเปรียบเทียบประสิทธิภาพก่อนและใช้ฐานความรู้

4.3.1 เปรียบเทียบผลการวิจัยกับ งานวิจัยเดิมที่ใช้ คลังข้อมูล

ตารางที่ 4.2 แสดงให้เห็นถึงการเปรียบเทียบผลการวิจัยเมื่อนำเอาการพิจารณาทางความหมายเข้ามาใช้ เราได้ทดลองด้วยการนำเอา ระบบอัตโนมัติค้นหาคำสำคัญเดิม มา เพิ่มส่วนของการพิจารณาความหมายเข้าไป ขั้นตอนการทำการทดลองคือ นำเอาคำคู่แข่งที่ถูกปฏิเสธโดย EXTRACTOR, KEA และ Database Mapping มาพิจารณาอีกครั้งโดยใช้ความหมาย ผลการทดลองแสดงให้เห็นได้ว่า เมื่อเพิ่มการพิจารณาทางความหมายเข้าไปทำให้ระบบเดิมที่มีอยู่ทำงานได้ดีขึ้นอย่างมีนัยสำคัญ

ตารางที่ 4.2 ประสิทธิภาพของระบบที่นำเสนอเมื่อเทียบกับงานวิจัยเดิมที่มีอยู่

Methods	%Precision	%Recall	F-Measure
1. EXTRACTOR	0.53	0.70	0.60
2. EXTRACTOR + proposed function	0.56	0.80	0.66
3. KEA	0.53	0.72	0.61
4. KEA + proposed function	0.58	0.83	0.68
5. Database Mapping	0.52	0.67	0.58
6. Database Mapping + proposed function	0.61	0.74	0.67

4.3.2 เปรียบเทียบผลการวิจัยที่ใช้ความหมายกับงานวิจัยเดิมที่ไม่ใช้ความหมาย

จากการทดลองในข้อ 4.3.1 พิสูจน์ให้เห็นได้ว่าการพิจารณาความหมายช่วยเพิ่มประสิทธิภาพของการทำงานของระบบได้ ด้วยการเพิ่มฟังก์ชันนี้เข้าไปในระบบเหล่านั้น แต่อย่างไรก็ตาม ผลการทดลองที่ได้ขึ้นอยู่กับประสิทธิภาพของระบบเดิมด้วย ดังนั้น การทดลองในส่วนนี้ จะเป็นการพิสูจน์ว่า ระบบอัตโนมัติค้นหาคำสำคัญที่ใช้การพิจารณาทางความหมาย มีประสิทธิภาพที่ดีกว่าระบบเดิมที่อยู่

ตาราง 4.3 และ 4.4 แสดงผลการทดลองที่ได้จากการเปรียบเทียบกับ EXTRACTOR (Turney, P. D.: 1997,2000,2001,2003) และ KEA (Witten, I.H., et al.: 1999) ที่เป็นตัวแทนของระบบอัตโนมัติค้นหาคำสำคัญ ที่ใช้คลังข้อมูล และเปรียบกับ ระบบค้นหาคำสำคัญของ Matsuo (Matsuo, Y., and Ishizuka, M.: 2004) ที่เป็นตัวแทนระบบที่ไม่ใช้คลังข้อมูล ในตารางที่ 4.3 และ 4.4 คำที่เป็นตัวเอียงและหนา คือคำที่เหมือนกับคำสำคัญที่ผู้เขียนกำหนดไว้ ซึ่งจะเห็นว่าระบบที่นำเสนอให้ คำสำคัญที่ตรงกับผู้เขียนมากกว่า

ตารางที่ 4.3 ตัวอย่างของผลที่ได้ของการวิจัยเมื่อเทียบกับงานวิจัยเดิม ที่ใช้คลังข้อมูล

EXTRACTOR	KEA	KEA2	The proposed system
exhibit	Baton Based	Design Patterns	<i>interface design</i>
instruments	<i>interactive exhibits</i>	input device	<i>worldbeat system</i>
user feedback	<i>interface design</i>	user interface	<i>worldbeat exhibit</i>
user interface design	user interface	user <i>interface design</i>	User interface
<i>WorldBeat system</i>	<i>WorldBeat system</i>	virtual realities	<i>music</i>

ตารางที่ 4.4 ตัวอย่างของผลที่ได้ของการวิจัยเมื่อเทียบกับงานวิจัยเดิม ที่ไม่ใช้คลังข้อมูล

Matsuo's System	The proposed system
<i>digital computer</i>	<i>digital computer</i>
<i>storage capacity</i>	<i>storage capacity</i>
<i>imitation game</i>	<i>imitation game</i>
machine	<i>discrete-state machine</i>
human mind	<i>intelligence</i>
universality	compute machinery
logic	computation
property	Object
mimic	sense
<i>discrete-state machine</i>	man

4.3.3 เปรียบเทียบประสิทธิภาพของการนำเอาฐานความรู้ของคำสำคัญมาใช้

จากข้อเสนอของการวิจัยที่บอกว่า ฐานความรู้ที่สร้างและปรับปรุงตัวเองได้โดยอัตโนมัติ ช่วยเพิ่มประสิทธิภาพของการค้นหาคำสำคัญได้ ตารางที่ 3.5 แสดงให้เห็นว่า ในช่วงที่ 1-3 ของการทำงานของระบบอัตโนมัติค้นหาคำสำคัญ เปอร์เซ็นต์ของคำสำคัญที่ตรงกับคำสำคัญที่ผู้เขียนกำหนดมีมากขึ้นอย่างมีนัยสำคัญ แต่อย่างไรก็ตาม เนื่องจากงานวิจัยนี้เสนอว่าใช้เพียงเอกสารนั้นๆ ในการพิจารณาความหมาย ซึ่งส่วนใหญ่เอกสารที่นำมาจะเป็นเอกสารวิชาการที่มีข้อมูลไม่มากนัก ทำให้เมื่อทดลองไปอีก 2 รอบ ประสิทธิภาพของระบบยังคงเท่าเดิม

ตารางที่ 4.5 ประสิทธิภาพของการนำเอาฐานความรู้ของคำสำคัญมาใช้

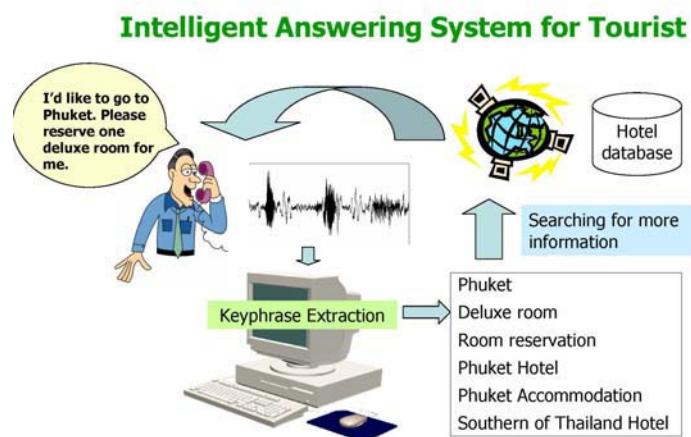
# of extraction cycle	%of matched keyphrases
1	54.90
2	66.98
3	74.12
4	75.55
5	75.55

บทที่ 5 การนำงานวิจัยไปใช้ประโยชน์

รูปที่ 5.1 – 5.3 แสดงถึงตัวอย่างการประยุกต์ใช้ การค้นหาคำสำคัญโดยใช้ความหมาย กับ การประมวลผลสารสนเทศ ต่างๆ ได้แก่ ระบบถาม-ตอบ (Question/Answering System) การจัดประเภทเอกสาร (Text Classification) และ การสรุปใจความสำคัญของเอกสาร (Text Summarization)



รูปที่ 5.1 ระบบผู้เชี่ยวชาญสำหรับวินิจฉัยโรค
(Interactive Medical Diagnosis Expert System)



รูปที่ 5.2 ระบบถาม-ตอบทางโทรศัพท์ (Intelligent Answering System for Tourist)

จากรูปที่ 5.1 และ 5.2 เราสามารถประยุกต์ใช้ระบบอัตโนมัติค้นหาคำสำคัญ กับระบบสนทนาอัตโนมัติได้ ในรูปที่ 5.1 เป็นระบบผู้เชี่ยวชาญสำหรับวินิจฉัยโรคได้ และ รูปที่ 5.2 เป็นระบบถาม-

ตอบสำหรับการท่องเที่ยว ทั้งสองระบบผู้ใช้จะป้อนข้อมูลเข้ามาในลักษณะที่เป็นภาษาธรรมชาติ ซึ่งมีรูปแบบของการถามหลากหลาย ถ้าต้องให้ระบบทั้งสอง มาวิเคราะห์คำทุกคำในประโยค ก็จะทำได้ยาก ดังนั้น ถ้ามีระบบอัตโนมัติค้นหาคำสำคัญ ดึงเฉพาะคำที่สำคัญ ออกมา ก็ทำให้ระบบทำงานได้รวดเร็วขึ้น



รูปที่ 5.3 ระบบอัจฉริยะสำหรับจัดประเภทบทความวิชาการ
(Intelligent Technical Paper Classification)

สำหรับรูปที่ 5.3 จะเป็นการประยุกต์ใช้ ระบบอัตโนมัติค้นหาคำสำคัญกับการแยกประเภทเอกสาร ตัวอย่างการประยุกต์ใช้คือ ระบบสำหรับจัดประเภทของบทความให้ตรงกับ Profile ของผู้เชี่ยวชาญ ระบบที่ใช้อยู่ปัจจุบันจะอาศัยมนุษย์ในการอ่านเพื่อทำความเข้าใจ และเลือกบทความให้ตรงกับ ความถนัดของผู้เชี่ยวชาญแต่ละท่าน แต่ถ้าระบบมีการจัดเก็บความถนัดของผู้เชี่ยวชาญแต่ละท่าน ด้วย คำสำคัญ ซึ่งอาจจะเป็นชื่อสาขางานวิจัยที่สนใจ จากนั้นระบบอัตโนมัติค้นหาคำสำคัญ ก็จะอ่านบทความและดึงเอาเฉพาะคำสำคัญออกมา จากนั้น คำสำคัญเหล่านี้ จะถูกนำไปเปรียบเทียบกับ ความถนัดที่ผู้เชี่ยวชาญให้ไว้ ก็จะสามารถจัดประเภทบทความได้ใกล้เคียงกับความต้องการของ ผู้เชี่ยวชาญได้

บรรณานุกรม

- Banerjee, S., Pederson T. (2003) "Extended gloss overlap as a measure of semantic relatedness." In Proc. Of IJCAI-03:805-810.
- Breiman, L., (1996), "Bagging Predictors", Machine Learning, Vol. 24, No. 2, pp. 123-140.
- Budanitsky, A. and Graeme H., (2001), "Semantic Distance in Wordnet: An Experimental, Application-Oriented Evaluation of Five Measures", The NAACL 2001 Workshop on WordNet and Other Lexical Resources [Electronic], Pittsburgh, USA, Available: <http://citeseer.ist.psu.edu/budanitsky01semantic.html> [2005, December 10].
- Chein et Marie-Laure Mugnier, M. (1992) "Conceptual Graphs fundamental notions." Intelligence Artificial, Vol.6,:365-406
- Chen, K.H. and Wu, C.T., (1999) "Automatically Controlled-Vocabulary Indexing for Text Retrieval", The 12th Research on Computational Linguistics Conference, Hsinchu, Taiwan, pp. 171-185.
- Chey, C., Kumhom, P. and Chamnongthai, K. (2006) Khmer Printed Character Recognition by using Wavelet Descriptors. International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems, 14(3), 337-350.
- Cognitive Science Laboratory at Princeton University, "WordNet 2.1 Browser", Available: <http://wordnet.princeton.edu>
- Connexor Oy Inc., (2005) Machine Syntax [Online], Available: <http://www.connexor.com/demo/syntax/>
- Frank, E., Paynter, G.W., Witten, I.H., Gutwin, C. and Nevill-Manning, C.G., (1999) "Domain-specific keyphrase extraction", The 16th International Joint Conference on Artificial Intelligence (IJCAI-99), July 31 - August 6, Stockholm, Sweden, pp. 668--673.
- Gutwin, C., Paynter, G.W., Witten, I.H., NevillManning, C.G. and Frank, E., (1998) Improving Browsing in Digital Libraries with Keyphrase Indexes, Technical Report, Department of Computer Science, University of Saskatchewan, Canada.
- Hist, G., St-Onge, D. (1998) "Lexical Chains as Representations." In Fellbaum:305-332
- Jiang, J. J., and Conrath, W. D., (1997) "Semantic Similarity Based on Corpus Statistic and Lexical Taxonomy" In Proceedings of International Conference on Research in Computational Linguistics, Taiwan.
- Jones, S. and Paynter, G.W., (2003), "An Evaluation of Document Keyphrase Sets", Journal of Digital Information [Electronic], Vol. 4, No. 1, Available: <http://jodi.ecs.soton.ac.uk/Articles/v04/i01/Jones/> [2005, September 9].
- Kongkachandra, R., Chamnongthai, K., and Kimpan, C., (2006), Intelligent Keyphrase Extraction for Digital Library, **Information Technology and Libraries**, Vol. 25, No. 2, (to be appeared).
- Leacock, C., Chodorow, (1998) M. "Combining Local Context and WordNetSimilarity for Word Sense Identification." In Fellbaum:265-283.
- Lesk, M., (1986) "Automatic Sense Disambiguation Using Machine Readable Dictionaries: How to tell a Pine Cone from a ice Cream Cone." In Proceedings

of the 5th Annual Conference on Systems Documentation, Toronto, Ontario, Canada, 24-26.

Matsuo, Y., and Ishizuka, M., (2004) "Keyphrase Extraction from a Single Document Using Word Co-Occurrence Statistical Information", *International Journal on Artificial Intelligence Tools*, Vol. 13, No. 1, pp. 157-169.

Montejo-Ráez and Steinberger, R., (2004) "Why Keywording Matters", *High Energy Physics Libraries Webzine [Electronic]*, Issue 10, Available: <http://library.cern.ch/HEPLW/10/papers/2/> [2005, September 1].

Montejo-Ráez, A., (2002) "Towards Conceptual Indexing Using Automatic Assignment of Descriptors", *Workshop in Personalization Techniques in Electronic Publishing on the Web: Trends and Perspectives*, May 28, Málaga, Spain.

Montes-y-Gomez, M., Lopez-Lopez, A. and Gelbukh, A.F., (1999) "Document Title Patterns in Information Retrieval", *Lectures Notes in Computer Science*, Vol. 1692, pp. 364-367.

Montes-y-Gomez, M., and Lopez-Lopez A., (2000) "Comparison of Conceptual Graphs." *Proc. MICAI-2000*, 1st Mexican International Conference on Artificial Intelligence, Acapulco, Mexico, April

Montes-y-Gomez, M., Lopez-Lopez, A., and Gelbukh, A. (2000) "Information Retrieval with Conceptual Graph Matching." *Proc. DEXA-2000*, 11 th International Conference and Workshop on Database and Expert Systems Application, Greenwich England

Montes-y-Gomez, M., Lopez-Lopez, A., Gelbukh, A., and Baeza-Yates, R. (2001). *Flexible Comparison of Conceptual Graphs*, *Proc DEXA-2001*, 12th International Conference and workshop on Database and Expert System Applications, *Lecture Notes in Computer Science*, N2113, Springer-Verlag, 2001, pp 102-111.

Percival, D. B. and Walden, A. T. (2000). *Wavelet Methods for Time Series Analysis*. Cambridge University Press.

Resnik, P., (1995) "Using Information Content to Evaluate Semantic Similarity in a Taxonomy." *Proceeding of The 14th International Joint Conference on Artificial Intelligence*. Vol. 1:448-453.

Sowa, J.F. (1984) *Conceptual Structures: Information Processing, Mind And Machine*, Addison-Wesley, Reading, MA, USA.

Teich, E. and Fankhauser, P. (2004) "WordNet for Lexical Cohesion Analysis", *The Second Global WordNet Conference*, January 21-23, Masaryk University, Brno, Czech Republic, pp. 326-331.

Turney, P. D. (1997) *Extraction of Keyphrases from Text: Evaluation of Four Algorithms*, Technical Report NRC 41550, National Research council.

Turney, P. D. (2000) "Learning Algorithms for Keyphrase Extraction", *Information Retrieval*, Vol. 2, No. 4, pp. 303-336.

Turney, P. (2001) "Mining the Web for Synonyms: PMI-IR Versus LSA on TOEFL", *The Twelfth European Conference on Machine Learning (ECML-2001)*, Freiburg, Germany, pp. 491-502.

Turney, P. (2003) "Coherent Keyphrase Extraction via Web Mining", *The Eighteenth International Joint Conference on Artificial Intelligence (IJCAI-03)*, Acapulco, Mexico, pp. 434-439.

University of Ottawa (2000) DIPETT on the Web [Online], Available: <http://www.site.uottawa.ca/tanka/dipett-on-the-Web/frames.html> [September 9].

Witten, I.H., Paynter, G.W., Frank, E, Gutwin, C., and NevillManning, C.G. (1999) "Kea: Practical Automatic Keyphrase Extraction", The Fourth ACM Conference on Digital Libraries (DL'99), August 11-14, University of California, Berkeley, USA, pp. 254-256.

Yang, D., and Powers, D., (2005) "Measuring Semantic Similarity in the Taxonomy of WordNet." Flinders University of South Australia.

Zhang, H., Sun, J., Wang, B., and Bai, Shou. (2005) "Computation on Sentence Semantic Distance for Novelty Detection" Inst. of Computation Tech., The Chinese Academy of Science , Beijing, China.

-ภาคผนวก-

ผลงานที่ได้รับการตีพิมพ์

Chey, C., Kumhom, P. and Chamnongthai, K. (2006) Khmer Printed Character Recognition by using Wavelet Descriptors. **International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems**, 14(3), 337-350.

Kongkachandra, R., Chamnongthai, K., and Kimpan, C., (2006), Intelligent Keyphrase Extraction for Digital Library, **Information Technology and Libraries**, Vol. 25, No. 2, (to be appeared).

Abductive Reasoning for Keyword Recovering in Semantic-based Keyword Extraction, **Informatica**, (submitted)

KHMER PRINTED CHARACTER RECOGNITION BY USING WAVELET DESCRIPTORS

C. CHEY, P. KUMHOM, and K. CHAMNONGTHAI

*King Mongkut's University of Technology Thonburi
91 Prachautit Rd., Bangmod, Tungkru, Bangkok 10140, THAILAND*

Received 25 October 2005

Revised 3 May 2006

In Khmer printed characters, same character has various shapes according to the fonts and some characters are very similar in shape. In this paper we try to solve these problems, and propose a method of Khmer printed character recognition by using Wavelet Descriptors. In the recognition, firstly the Khmer printed character images are converted to skeleton forms, then skeletons of Khmer character are converted to temporal domain. The templates are obtained by wavelet coefficients from the character training set. To match the input characters with templates, the character recognition method using deformable wavelet descriptor is adapted by using fixed template and Euclidean distance classifier for matching. The smallest distance is the recognition result of the proposed method. As a result, the deformation can be skipped because it might get low recognition rate of similar characters. The experiment consists of two parts. The first part is to evaluate the overall recognition rate of input characters with three different sizes (22-point, 18-point and 12-point) from 10 different fonts of Khmer printed character. Twenty styles of characters are used as the training set. The results show 92.85, 91.66, and 89.27 percent for 22-point, 18-point, and 12-point respectively. The second part is to specifically evaluate the system, testing with one document that has 21 pages of Khmer printed character with different resolutions from a scanner and facsimile (fax). The document is initially printed with 300 dpi (dots per inch), then scanned with three different resolutions, 600 dpi, 300 dpi and 150 dpi. The document that received from fax machine is scanned by 300dpi. The results show 92.99, 88.61, and 80.05 percent recognition rate for 300, 150 dpi resolutions, and input from fax respectively.

Keywords: Optical Character Recognition; Wavelet Descriptor; Deformable Template Matching; Template Matching.

1. Introduction

Optical Character recognition (OCR) becomes more and more important in the modern world. This technology enables the use of a pen device as a keyboard and, thus, reduces the need for a bulky keyboard as standard equipment in palm (short note) computing devices. On-line hand writing character recognition is also important as a means of data entry for those that do not wish to type on the keyboard, for whatever their reason may be. For postal system, OCR can fasten the process of address classification by finding the zip code and separating the mail into different distribution piles for each destination. Also hard copies and many old books are needed to represent in electronic form that need OCR machine or OCR software package for converting. It also plays important role in many applications such as the

information retrieval system, knowledge-base-updating system, and e-library. However, an OCR can not be applied into all languages. There are many researches that proposed the methods for optical character recognition (OCR) system in different languages. These methods can be categorized into two groups, including statistical classification [5, 6], and template matching [1, 2, 3, 4].

In the first group, Kijirikul, Sinthupinyo and Supanwansa [5] proposed Thai printed recognition using combination inductive logic programming with back propagation neural network (BNN). They present a new approach that combines two learning algorithms that are Inductive Logic Programming (ILP) and Back Propagation Neural Network (BNN). The results show high recognition rate with time consuming. Thammano and Ruxpakawon [6] proposed an approach to recognize the Thai printed characters using the hybrid of global feature, local feature, fuzzy membership function and neural network. The methods in this group give high recognition rate, but they require much more computation time.

For the second group in template matching approach, in general, a pre-processed image of an input character is matched with the templates representing some kind characters' features. Then, the input image is recognized as the character whose template match the image the most. There are two kinds of templates, fixed template and deformable template. The following two researches use the fixed templates. Liao and Lu [1] proposed a method for Chinese character recognition by using moment functions as the features of the template for recognizing Chinese characters, especially for characters that very similar in shapes. Euclidean distance classifier is used to classify all Chinese characters. The results show those features can represent all Chinese characters reasonably well, especially for Legendre moment. The methods mentioned above work well with the fixed fonts, they do not work well for the OCR system with various forms of inputs such as printed characters with many fonts and handwritten characters. Chiang, Liao, Lu and Pawlak [2] proposed a method for Chinese character recognition, especially for the characters closed in shapes, using Gegenbauer Moments as the features of the Chinese characters. The results show that Gegenbauer Moment-based features extracted from each Chinese character can represent all Chinese characters reasonably well. The Gegenbauer moment with smaller values of parameter would provide more recognition power to the recognition system. The methods mentioned above work well with the fixed fonts, they do not work well for the OCR system with various forms of inputs such as printed characters with many fonts and handwritten characters. The deformable template is more appropriate for such inputs. For the deformable template, Phaopanus, Arunrungrusmi and Chamnongthai [4] proposed deformable wavelet descriptors for the Thai handwritten character recognition. In the recognition, contour of Thai characters are utilized to determine template in term of temporal domain. Then, range of deformation is obtained by standard deviation of Wavelet descriptor coefficients from characters training set. To fit the templates with input characters, all templates are deformed to obtain the best fit by determining deformation range. The best fit of each character template is represented

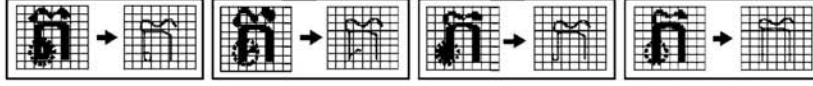


Fig. 4. character /tor/ in skeleton form Fig. 5. character /kor/ in skeleton form Fig. 6. character /phor/ in skeleton form Fig. 7. character /gar/ in skeleton form

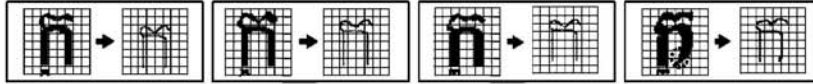


Fig. 8. character /gar/ in skeleton form Fig. 9. character /gar/ in skeleton form Fig. 10. character /gar/ from different fonts represented in skeleton form Fig. 11. characters /gar/ from different fonts represented in skeleton form

the flesh of same characters from different fonts are different according to the fonts. However, the skeletons of the similar characters are differentiated as shown in Fig. 4, Fig. 5, Fig. 6 and Fig. 7, and the skeletons of the same characters from different fonts are mostly same as shown in Fig. 8, Fig. 9, Fig. 10 and Fig. 11.

Since same characters from different fonts are varied, it makes the features of same characters varied. Some characters are very similar in shapes, it causes confusion with other characters. When we analyze them into skeleton form, we found that different characters in similar fonts are obviously able to differentiate as samples shown in Fig. 4-6, and same characters in shape-varied fonts are simultaneously recognized as the same result as shown in Fig. 7-11. The skeletonization technique is used to normalize the shapes of the characters. Wavelet descriptor uses the basic functions with local support and multiscale dilation, thus it can model local features as well as global ones. These make system flexible for the template, and robust against changes. Normally, information exists in the low frequency ranges[7]. At the same time, most of the methods would cut the signals in the high frequency range because they are usually noises. Therefore, Wavelet descriptor coefficient at the low frequency range is used as the feature extraction of all characters. The proposed method uses fixed template and Euclidean distance classifier as the matching scores instead in order to reduce the computation time. The reason is that, in most cases, the relative Euclidean distance should represent the relative correlation. In other words, if we compare the Euclidean distance of two templates from an input, the template with shorter distance would give a higher correlation score.

3. Proposed Method

In Fig. 12, a common setup of proposed method is illustrated. The first step in the process is to digitize the analog document using an optical scanner or digital camera. The extracted characters may then be preprocessed to facilitate the extraction of

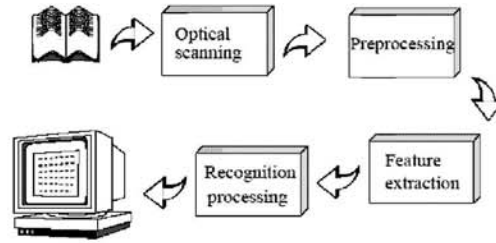


Fig. 12. OCR system

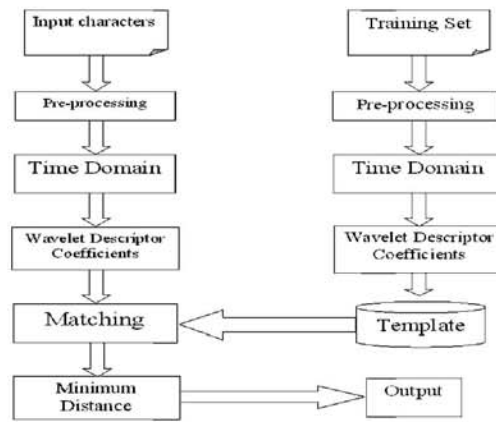


Fig. 13. block diagram of proposed method

features in the next step. The identity of each symbol is found by comparing the extracted features with descriptions of the character classes obtained through a previous learning phase. Finally the recognition process is used to recognize all characters.

3.1. System

Fig. 13 shows the system block diagram of Khmer character recognition. The set of characters called training set is used to construct the templates storing in terms of the means values of wavelet descriptor coefficients from their temporal data. During the running time, an input image is pre-processed and converted to the temporal domain, then, to wavelet coefficients. It is matched with all the templates using the Euclidean distances classifier. They find distance between input and every template then compare all distances. The smallest distance is classified to be the output of proposed method. The proposed method is separated into three parts, including Preprocessing, Feature extraction, and Classification process.

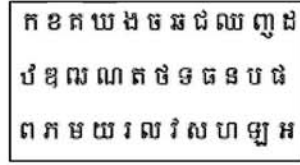


Fig. 14. character images before Thresholding

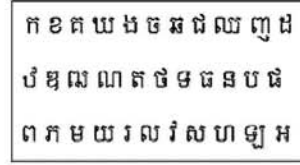


Fig. 15. binarized character images after Thresholding

3.2. Pre-Processing

In the pre-processing, we extract gray scale from character images obtained from scanner, then convert it into binary images by using thresholding. All the intensity values above the threshold intensity are converted to one high intensity value, representing black. All intensity values below the threshold are converted to another low intensity value representing white. The binary character images, then, are transformed to skeletonization form, applying the thinning algorithm to normalize the shape of the characters.

3.2.1. Binarization

In this paper, we assumed that the document does not have high noise content, in which case the image can be binarized directly. Fig. 14 shows a set of 33 Khmer consonants, representing in gray scale image. Fig. 15 is the character mages, representing in binary images that converted from gray scale image from Fig. 14.

3.2.2. Thinning

The skeleton image is usually smaller than its original image. With thinning algorithm, it not only increases the speed of recognition, but also improves the performance of recognition. Skeletonization is the process of peeling off of a pattern as many pixels as possible without affecting the general shape of the pattern. In other words, after pixels have been peeled off, the pattern should still be recognized. The obtained skeleton must have the following properties:

- as thin as possible
- connected
- centered

When these properties are satisfied, the algorithm must stop. A number of thinning algorithms have been proposed and are being used. The most common used algorithm is the classical Hilditch algorithm [8]. A simple algorithm is used for thinning images [9, 10]. The algorithm makes two passes: one pass on the actual image and another pass over intermediate images, which are constructed from the original image during the first pass. The result after the two passes decides if a pixel is to be removed to get the skeleton of the object of interest.

The thinning algorithm works as follows. Fig. 16 shows an input image on

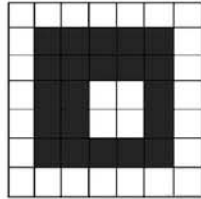


Fig. 16. a sample of character to be thinned

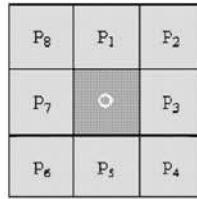


Fig. 17. mask for thinning

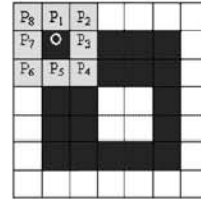


Fig. 18. using thinning mask

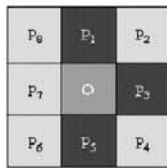


Fig. 19. transition thinning masks(a)

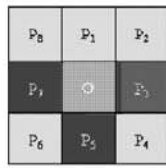


Fig. 20. transition thinning masks(b)

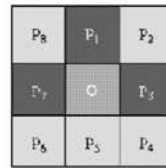


Fig. 21. transition thinning masks(c)

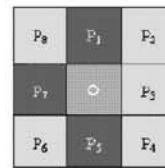


Fig. 22. transition thinning masks(d)

which the thinning algorithm is to be applied. Each box represents a pixel. Fig. 17 shows a 3x3 mask that is used to move over the input image during the first pass and then on the intermediate image during the second pass. The center of the mask is positioned on the pixel of interest during the scanning process and the numbered boxes represent the corresponding neighbor pixels. The pixel of interest, shown as "o", is marked for deletion when all the following conditions are satisfied. These conditions are checked only for a "dark pixel".

- If the number of neighboring 'dark pixels' of "o" is between 2 and 6.
- If the number of transitions from 'bright pixel' to 'dark pixel' in the order shown is one.
- If the logical function $P1.P3.P5 = 0$ and also $P3.P5.P7 = 0$, as shown by the darkened portions of the masks in the Fig. 19 and Fig. 20. These conditions are repeated for each pixel in the image as it is scanned. The intermediate result is an image, which consists of pixels that are marked for deletion and other pixels that are left as in the original input image. The second pass is performed on this intermediate image and now the conditions of deletion on the pixel are as follows:
- If the number of neighboring 'dark pixels' of "o" is between 2 and 6.



Fig. 23. original image

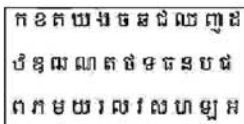


Fig. 24. original image after Thresholding

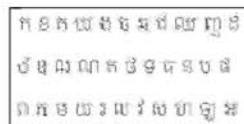


Fig. 25. thinned image

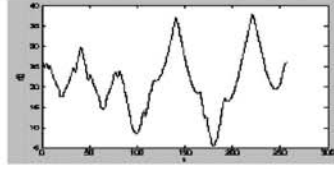


Fig. 26. temporal domain of /gar/ that similar with /phor/

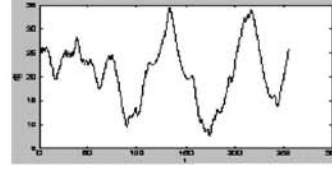


Fig. 27. temporal domain of /phor/ that similar with /gar/

- If the number of transitions from 'bright pixel' to "dark pixel" in the order shown is one.

- If the logical function $P1.P3.P7 = 0$ and also $P1.P5.P7 = 0$, as shown by the darkened portions of the masks in fig. 21 and fig. 22. As mentioned above, the pixel is deleted only when all the conditions are satisfied. Some results using this thinning algorithm are shown in fig. 25. It is evident from these results that input characters have been reduced to one-pixel width.

3.3. Feature Extraction

Feature extraction is the most important part in character recognition because features are main keys for recognizing unknown characters. The objective of feature extraction is to capture the essential characteristics of the characters, and it is generally accepted that this is one of the most difficult problems of pattern recognition. In our method, the characters after spatial domain are converted to wavelet descriptor coefficients, are called feature extractions.

3.3.1. Temporal Domain

First, we convert all of skeleton characters into temporal domain as reference pattern of each character. Then, the training character sets of Khmer printed characters are performed in the same way. In temporal domain, we want to extract only the wanted objects and remove the unwanted objects. So, the temporal algorithm does not only increase the speed of recognition, but also improves the performance of recognition. equation (1) gives the definition of the temporal domain for representing all skeleton characters. Representing all skeleton characters in temporal domain is giving by

$$r(t) = \sqrt{(X(t) - cg_x)^2 + (Y(t) - cg_y)^2} \quad (1)$$

where $r(t)$ represents the skeleton characters in temporal domain, $[(X(t), Y(t))]$ is skeleton characters coordinate in spatial domain, and $[cg_x, cg_y]$ is center of mass of the character.

3.3.2. Wavelet Descriptors

Wavelet based approaches become increasingly popular in pattern recognition, and recently have been applied for character recognition [3, 7, 4]. After all skeleton char-

acters are converted into temporal domain, the skeleton characters of training set in temporal domain are presented by wavelet descriptor to obtain wavelet coefficients. Wavelet descriptor uses a set of basic functions with local support and multiscale dilation, thus, it can model local as well as global features. Wavelet descriptor [7, 11] offers natural multi-resolution representation of the signal. Wavelet decomposes 1-D signal into multi-resolution scales. The wavelet equations are given as seen in appendix A.

3.4. Matching Process

Recognition techniques based on matching represent each class by a prototype pattern vector. An unknown pattern is assigned to the class to which it is closest in term of a predefined metric. In this paper, the minimum distance classifier, which, as its name implies, computes the (Euclidean) distance between unknown and each of prototype vectors called template, is applied. It chooses the smallest distance to make a decision. In matching process, firstly, the templates are obtained from wavelet coefficients of the characters training set. Then, wavelet coefficients of all input characters are matched with every template. To match the input characters and template the Euclidean distance classifier is used. We calculate the distance between input character and the templates. From matching, then, the template that gives the minimum distance is the output to be classified. The minimum classifier approach is given as seen in appendix B [11].

4. Experimental Results

In order to test the efficiency of the proposed algorithm, we test all Khmer characters. Each input character is binarized with '1' denoting the object and '0' the background. The input characters are segmented into isolated characters. Twenty font styles (LIMON S1, LIMON S2, LIMON S3, LIMON S4, LIMON S5, LIMON S6, LIMON S7, ABC-TEXT-11, ABC-TEXT-12, ABC-TEXT-13, ABC-TEXT-14, ABC-TEXT-15, ABC-TEXT-16, ABC-TEXT-17, ABC-TEXT-19, ABC-TEXT-20, Limon S2D, SOVAN 1, TAPROM and EKREACH) are used as training set. For testing, the first part is to evaluate the overall recognition rate of input characters with three different sizes (22-point, 18-point and 12-point) from 10 different fonts (ABC-TEXT-01, ABC-TEXT-02, ABC-TEXT-03, ABC-TEXT-04, ABC-TEXT-05, ABC-TEXT-60, ABC-TEXT-07, ABC-TEXT-08, ABC-TEXT-09 and ABC-TEXT-10) of Khmer printed character. The results show 92.85 percent recognition rate for the 22-point, 91.66 percent recognition rate for the 18-point, and 89.27 percent recognition rate for the 12-point. The second part is to specifically evaluate the system, testing with one document that has 21 pages from font ABC-TEXT-18 of Khmer printed character with different resolutions from scanner and fax. The document is printed with 300 dpi (dots per inch) then scanned with three different resolution, 600 dpi, 300 dpi and 150 dpi. For the document that received from fax machine is scanned with 300 dpi resolution. The results show 92.99 percent

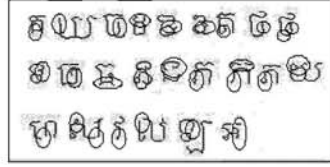


Fig. 28. misclassified characters caused by skeleton

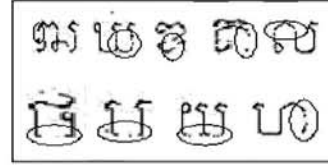


Fig. 29. misclassified characters caused by noise

recognition rate for 600 dpi resolution, 88.61 percent recognition rate for 300 dpi resolution, 80.05 percent recognition rate for 150 dpi resolution and 71.68 percent recognition rate for the input from fax machine. The testing characters are shown in appendix-A. All the experimental results are shown in Table 1 and Table 2.

5. Discussion

The first experiment, the proposed algorithm can receive high recognition rate. In this case, we also see that the system has misclassified some characters in some cases. These misclassifications are probably caused by preprocessing errors. During converting binary character image to skeleton characters form, some skeleton characters shape could not be obtained as the original one. It loses some parts of the characters as shown in Fig. 28.

The proposed algorithm also received 92.99 % recognition rate for the input data scanned with 600 dpi resolution, 88.61% recognition rate for the input data scanned with 300 dpi resolution, 80.05% recognition rate for the input data scanned with 150 dpi resolution and 71.68% recognition rate for the input data received from the Fax machine that scanned with 300 dpi resolution. In these cases, we also see that the system has misclassified some characters in some cases. These are probably caused by noise that makes some characters not fully connected or maybe make it looks like other characters. Especially, in case of the input characters from Fax machine the system has trouble to recognize most of the characters. From experiments, we realized that the system still has troubles identifying the characters that have low resolution, such as the input characters from Fax machine. Fig. 29 shows that some input characters are misclassified with others because some characters are not fully connected making them looked like other characters. In the case of similar

Table 1: Recognition results of input Images from 10 different fonts scanned by 300 dpi

Font size	Recognition rate
22-point	92.85%
18-point	91.66%
12-point	89.27%