

การเพิ่มประสิทธิภาพระบบเครือข่ายสังคมด้านรูปภาพด้วยเทคนิคการจัดอันดับใหม่แบบรวม

Enhancement image social networking by using Combination Re-ranking

พิจิตรา จอมศรี (Pijitra Jomsri)*

บทคัดย่อ

เทคโนโลยีเว็บ และระบบโซเชียลบุคมาร์กด้านรูปภาพถือเป็นอีกทางเลือกหนึ่งของผู้ใช้สำหรับการค้นหาข้อมูลภาพ ด้วยระบบดังกล่าวสามารถบริการผู้ใช้งาน ในด้านการค้นหารูปภาพ การแบ่งปัน การจัดการ และการให้บริการต่าง ๆ ผ่านทางระบบโซเชียลแท็กและเทคนิคอื่นๆที่เกี่ยวข้อง ทั้งนี้มีระบบโซเชียลบุคมาร์ก และเครือข่ายสังคมทางด้านรูปภาพที่ได้รับความนิยมในปัจจุบันจำนวนมากเช่น Flickr

งานวิจัยนี้มีวัตถุประสงค์เพื่อเพิ่มประสิทธิภาพผลการค้นหารูปภาพโดยทำการศึกษาถึงเทคนิคการจัดเรียงอันดับผลลัพธ์ของรูปภาพให้สอดคล้องกับความต้องการของผู้ใช้ ทั้งนี้จึงได้กำหนดการทดลองเพื่อทดสอบสมมติฐาน โดยการพัฒนาเทคนิคการจัดอันดับใหม่ จำนวน 4 เทคนิค คือ การจัดอันดับตามความเหมือน (SimRank) การจัดอันดับตามความเหมือนร่วมกับจำนวนผู้เข้าชมรูปภาพ (DTVRank) การจัดอันดับตามความเหมือนร่วมกับจำนวนผู้ที่ชื่นชอบรูปภาพ (DTFRank) การจัดอันดับตามความเหมือนร่วมกับจำนวนผู้เข้าชมรูปภาพ และจำนวนผู้ที่ชื่นชอบรูปภาพ (DTVFRank) โดยใช้ดัชนีแบบรวมของคำอธิบายภาพรวมกับแท็ก (DT) เพื่อทำการศึกษาลงถึงเทคนิคที่มีประสิทธิภาพมากที่สุด

ทั้งนี้ประสิทธิภาพของงานวิจัยนี้ได้ถูกประเมินด้วยเทคนิค NDCG และการทดสอบทางสถิติที่ระดับนัยสำคัญ 0.05 พบว่าเทคนิคการจัดอันดับใหม่แบบ DTVF(50:50) มีประสิทธิภาพดีที่สุดเมื่อเปรียบเทียบกับเทคนิคอื่นที่ทำการทดลอง ดังนั้นเทคนิคดังกล่าวจึงสามารถเพิ่มประสิทธิภาพผลการค้นหารูปภาพบนระบบโซเชียลบุคมาร์กทางด้านรูปภาพได้

คำสำคัญ: การค้นคืนเนื้อหาภาพ ,จัดอันดับใหม่, เครือข่ายสังคม,โซเชียลบุคมาร์ก

* อาจารย์ ดร. พิจิตรา จอมศรี ประจำสาขาวิชาเทคโนโลยีสารสนเทศ คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยราชภัฏสวนสุนันทา

Pijitra Jomsri, Ph.D. in Department of Information technology, Faculty of science and technology, Suan Sunandha Rajabhat University E-mail: pijitra.jo@ssru.ac.th, pijitra.jo@gmail.com

Abstract

Web Technology and image social bookmarking are alternatives for image searching. Web system allows users to search for images, sharing, organize, information service sources through social tagging and other method activities. One of image social bookmarking is such as Flickr.

This research is examined to increase efficiency of image search results by study for re-ranking image results relate to user needs. Therefore the experiment was created for test hypothesis. Four re-ranking techniques were created which are Similarity Ranking (*SimRank*), combination ranking of number of views (*DTVRank*), combination ranking of number of favorites (*DTFRank*), and combination ranking of number of favorite with number of views (*DTVFRank*) based on Description with Tag” or *DT* indexer.

The retrieval performances of this research are evaluated using mean values of Normalized Discount Cumulative Gain (NDCG) and statistic testing ($\alpha=0.05$). The result illustrate that *DTVFRank* (50:50) is better than other techniques. This primary evaluation in the experiments proves that the chosen heuristic re-ranking model can improve the efficiency of web image searching on social bookmarking websites.

Keywords: Content-Based Image Retrieval, Re- Ranking, Social network, Social bookmarking

บทนำ

ปัจจุบันการเครือข่ายสังคมออนไลน์เพื่อการสืบค้นข้อมูลมีความแพร่หลาย และมีแนวโน้มเพิ่มมากขึ้น โดยเฉพาะอย่างยิ่งการเผยแพร่ข้อมูลภาพถ่ายดิจิทัลออนไลน์ ผู้ใช้งานสามารถสืบค้นรูปภาพที่สนใจได้อย่างสะดวกและรวดเร็ว นอกจากนี้ผู้ใช้งานส่วนใหญ่นิยมสืบค้นผลงานภาพที่เกี่ยวข้องผ่านทางเว็บไซต์ที่ให้บริการสืบค้นข้อมูล และผ่านเว็บไซต์โซเชียลบุคมาร์ก (social bookmarking) ซึ่งเป็นการให้บริการบนเว็บที่แบ่งปันข้อมูลที่ผู้ใช้สนใจให้กับผู้อื่นที่สนใจข้อมูลในลักษณะเดียวกัน ทั้งนี้เว็บไซด์ดังกล่าวยังถือเป็นอีกทางเลือกหนึ่งที่จะช่วยให้ผู้ใช้สามารถแลกเปลี่ยนความรู้ และแบ่งปันข้อมูลภาพให้กับกลุ่มผู้ใช้ที่สนใจในเรื่องเดียวกัน ปัจจุบันมีเว็บไซต์โซเชียลบุคมาร์กที่ให้บริการข้อมูลด้านข้อมูลภาพ เช่น Flickr ซึ่งถือเป็นเว็บไซต์ที่ได้รับความนิยมจากผู้ใช้งานมากที่สุด อย่างไรก็ตามการสืบค้นรูปภาพจำเป็นต้องใช้ข้อมูลหรือเนื้อหาที่เกี่ยวข้องกับรูปภาพที่ถูกสร้างโดยผู้ใช้ เพื่อให้การสืบค้นรูปมีความสอดคล้องและตรงกับความต้องการของผู้ใช้มากที่สุด

ทั้งนี้งานวิจัยก่อนหน้านี้ผู้วิจัยได้ทำการศึกษาถึงวิธีการสร้างดัชนี (indexing) ที่เหมาะสมกับการค้นหาภาพถ่าย ซึ่งพบว่าการสร้างดัชนีโดยใช้คำสำคัญของภาพ (tag) ร่วมกับคำอธิบายภาพ หรือเรียกว่า “*DT indexer*” ทำให้ประสิทธิภาพของการค้นหาข้อมูลภาพเพิ่มมากขึ้น งานวิจัยนี้จึงได้มีแนวทางในการปรับปรุงการจัดอันดับผลการค้นหาใหม่จากเนื้อหาที่เกี่ยวข้องเพื่อเพิ่มประสิทธิภาพในการจัดอันดับผลการค้นหา ช่วยให้

ผู้ใช้งานสามารถค้นหาภาพถ่ายได้ตรงกับความต้องการมากยิ่งขึ้น ซึ่งจะส่งผลต่อการเพิ่มประสิทธิภาพการให้บริการของเว็บไซต์เครือข่ายสังคมและโซเชียลบุ๊กมาร์กต่อไป

วัตถุประสงค์ของการวิจัย

จากความสำคัญของปัญหาการวิจัยที่กล่าวมาข้างต้น เพื่อให้มีความชัดเจนในการตอบคำถามงานวิจัย ผู้วิจัยจึงได้กำหนดวัตถุประสงค์ของการวิจัยคือเพื่อเพิ่มประสิทธิภาพในการจัดอันดับผลการค้นหาด้านรูปภาพบนระบบเครือข่ายสังคมและออกแบบ และพัฒนาอัลกอริทึมสำหรับการจัดอันดับผลการค้นหาเชิงรูปภาพ

ขอบเขตการวิจัย

งานวิจัยนี้มุ่งเน้นในการเพิ่มประสิทธิภาพผลการค้นหาบนระบบเครือข่ายสังคมทางด้านข้อมูลภาพเท่านั้น รวมทั้งงานวิจัยนี้ใช้กลุ่มตัวอย่างในการทดลองจากผู้ใช้งานในระบบระบบเครือข่ายสังคมด้านข้อมูลภาพ โดยใช้กลุ่มตัวอย่างซึ่งเป็นนักศึกษา และอาจารย์ในมหาวิทยาลัยราชภัฏสวนสุนันทา

แนวคิดและทฤษฎีที่เกี่ยวข้อง

ขั้นตอนการประมวลผลของการค้นคืนสารสนเทศ

ขั้นตอนการประมวลผลของการค้นคืนสารสนเทศประกอบด้วยขั้นตอนจำนวน 5 ขั้นตอนหลักเป็นดังนี้ (ดร.ศุภชัย ตั้งวงศ์ศานต์, 2551) การทำดัชนี (indexing) เป็นการสร้างตัวแทนเอกสาร การจัดรูปแบบคำสอบถาม (query formulation) เป็นการสร้างตัวแทนคำสอบถาม การเทียบเคียงจับคู่ (matching) ตัวแทนคำสอบถามกับตัวแทนเอกสารการเลือก (selection) รายการผลลัพธ์ที่ตรงประเด็น การปรับเปลี่ยนคำสอบถามใหม่ (query reformulation) ในรอบต่อไป

1. การทำดัชนี

การทำดัชนีเป็นการสร้างตัวแทนเอกสาร วัตถุประสงค์ของการทำดัชนีก็เพื่อทำเป็นตัวแทนเอกสาร โดยจัดเป็นหมวดหมู่อย่างเป็นระบบ อันเป็นการเตรียมการในขั้นต้นของ การค้นคืนสารสนเทศ เพื่อการสืบค้นต่อไปให้ดำเนินการได้อย่างมีประสิทธิภาพและประสิทธิผล ขั้นตอนการทำดัชนีสามารถสร้างด้วยแรงหรือสร้างอย่างอัตโนมัติได้ รูปแสดงขั้นตอนการทำดัชนีจากต้นแหล่งที่เป็นเอกสาร

ขั้นตอนแรกเป็น lexical analysis โดยทำการแปลงสายตัวอักษรในเอกสารหรือสิ่งเป้าหมายกลายเป็นสายของคำศัพท์ เมื่อได้คำแต่ละคำในข้อความแล้ว ให้จัดการคำที่เป็นตัวเลข คำที่เชื่อมด้วยอักขรตัวใหญ่เล็ก หรือคำที่เชื่อมด้วยเครื่องหมายวรรคตอน (punctuation) ตามข้อกำหนด และส่งต่อไปในขั้นต่อไป ลำดับต่อไปเป็นการจำกัดคำโหล (stop-words) อันหมายถึง คำที่ปรากฏบ่อยครั้งมากจนไม่มีผลในการสืบค้นทางปฏิบัติ จึงไม่จำเป็นต้องนำมาสร้างเป็นดัชนี ขั้นต่อมา คือ stemming เป็นการหารากศัพท์ของคำในข้อความเพื่อใช้คำ stem นั้นเป็นตัวแทนในการสร้างดัชนีต่อไป ติดตามด้วยการเลือกคำศัพท์ เช่น กลุ่มของนามเพื่อสร้างดัชนี จากนั้นเป็นการสร้างคำศัพท์สัมพันธ์ (thesaurus construction) เพื่อจัดกลุ่มคำที่มีความหมายเดียวกันให้เชื่อมต่อกันอย่างเป็นระบบ และกำหนดคำศัพท์ควบคุม (controlled vocabulary) เมื่อได้คำศัพท์

ตามกระบวนการข้างต้น ผลลัพธ์ที่ได้นำมาจัดทำดัชนี และจัดเก็บตามการออกแบบของโครงสร้างที่สำคัญก็จะมีรูปแบบของแฟ้มข้อมูลผกผัน (inverted file) เพื่อการสืบค้นต่อไป

2. การจัดรูปแบบคำสอบถาม

การจัดรูปแบบคำสอบถามเป็นการสร้างตัวแทนคำสอบถาม การเขียนคำสอบถามเพื่อการค้นหาสารสนเทศที่ต้องการทำได้ในหลากหลายรูปแบบที่ง่าย และสะดวก คือ การใส่คำสำคัญ (keywords) เป็นหนึ่งคำศัพท์หรือหลาย ๆ คำเรียงต่อกันได้ตามที่ต้องการ เช่น ต้องการค้นหาเรื่อง Quantum Computing เขียนโดย Joseph Stelmach ก็เพียงเขียนคำสอบถามว่า Quantum Computing Joseph Stelmach ป้อนถามโปรแกรมค้นหาก็จะได้ผลลัพธ์ตามที่เทียบเคียงได้ การเขียนรูปแบบข้างต้นสามารถขยายผลโดยเชื่อมต่อกด้วย Boolean Operators อันได้แก่ AND, OR และ NOT เช่น (quantum or computing) AND (not stelmach) การเขียนคำสอบถามด้วยตัวเชื่อมเหล่านี้ จะทำให้การเขียนมีความซับซ้อน แต่ในขณะเดียวกันก็มีความหมายที่เจาะจงมากขึ้น อันทำให้ผลลัพธ์สอดคล้องตามความต้องการมากขึ้น อีกวิธีหนึ่งเป็นการสืบค้นแบบ Pattern โดยการอนุญาตให้ใช้ตัว Wild Card Character แทนตัวอักษรที่เป็น prefix หรือ suffix ของคำเช่น ใช้สัญลักษณ์ * แทน Wild Card ตัวอย่างการเขียนที่เป็น Pattern คือ compu*,*ters, *กุมาร* เป็นต้น

3. การเทียบเคียงจับคู่

การเทียบเคียงจับคู่ตัวแทนคำสอบถามกับตัวแทนเอกสาร เป็นการจับคู่ระหว่างคำสอบถามกับเอกสารโดยใช้ตัวแทนจากที่ได้มาในขั้นตอนก่อนหน้านี้ วิธีการขึ้นอยู่กับโมเดล (model) ของการค้นหาสารสนเทศ หากเป็น Boolean model หรือโมเดลที่เป็นลักษณะเดียวกันนี้จะให้ค่าการเทียบเคียงเป็นเพียง 2 ค่า คือ สอดคล้องหรือไม่สอดคล้อง จะไม่มีค่าระหว่างกลางทำให้ผลของการสืบค้นได้ค่าการเรียกกลับที่ต่ำสำหรับ Vector model หรือโมเดลในทำนองเดียวกัน ค่าการเทียบเคียงจะได้จากผลคูณภายในของเวกเตอร์ของคำสอบถามและของเอกสาร ซึ่งเป็นการวัดความใกล้เคียงสอดคล้องของแต่ละคู่ หากได้ค่าสูงแสดงถึงความเกี่ยวข้องสอดคล้องกันมาก ในทางตรงกันข้ามหากได้ค่าต่ำ หมายถึง ความสอดคล้องน้อยจากค่าที่ได้ยังสามารถนำไปจัดลำดับของเอกสารในผลลัพธ์ตามความสำคัญก่อนหลัง ค่าสูงอยู่ก่อน ค่าต่ำอยู่หลัง อนึ่ง ค่าการจับคู่ยังอาจปรับได้ในทางคำนวณโดยการให้น้ำหนักของเทอมในคำสอบถาม รวมทั้งค่าน้ำหนักใน Term-Document Matrix ที่สอดคล้องกันอย่างเหมาะสม

การจัดลำดับของเอกสารที่นิยมในระบบโปรแกรมค้นหาที่สำคัญ เช่น Google ยังได้จัดลำดับที่เรียกว่า PageRank โดยมีกระบวนการคำนวณจากความสำคัญของเว็บไซต์ นั้นมีการถูกอ้างอิงมากน้อย โดยดูจากจำนวนสายเชื่อมโยง (hyperlink)

4. การเลือกรายการผลลัพธ์ที่ตรงประเด็น

การเลือกรายการผลลัพธ์ที่ตรงประเด็น เป็นการเลือกของผู้ใช้ในผลลัพธ์ที่ปรากฏของเอกสารที่สอดคล้องตรงประเด็น ซึ่งผลลัพธ์ของการสืบค้นอาจจะเรียงลำดับความสำคัญในกลุ่มหัวเรื่องต่าง ๆ หรือกลุ่มประเภทต่าง ๆ อย่างอัตโนมัติ (classification หรือ clustering) นอกจากนี้ เพื่อช่วยผู้ใช้ในการเลือกที่ชัดเจนผลลัพธ์อาจจะแสดงเป็นหัวเรื่อง (title) ส่วนของบทคัดย่อ (abstract) พร้อมทั้งเน้นเป็นแถบสว่าง (highlight) ที่ชัดเจนในคำศัพท์หรือเทอมของคำสอบถามที่ต้องการค้นหา

5. การปรับเปลี่ยนคำสอบถาม

การปรับเปลี่ยนคำสอบถามใหม่ในรอบต่อไป ในระบบการค้นคืนสารสนเทศ ผู้ใช้มักจะมีปัญหาการตั้งคำถาม ปัญหาไม่ใช่ตั้งไม่ได้แต่การตั้งที่ให้ผลการสืบค้นที่มีประสิทธิภาพนั้นทำได้ไม่มากนัก บ่อยครั้งที่ผู้ใช้ต้องเสียเวลาในการเลือกชุดเอกสารที่เป็นผลลัพธ์ที่ได้มามากมายแต่มีขยะปะปนมากี่มาก แนวทางหนึ่งในการแก้ปัญหาเพื่อสร้างความพึงพอใจแก่ผู้ใช้คือ การปรับเปลี่ยนคำสอบถามใหม่ (query reformulation) ด้วยเทคนิคของ query expansion หรือของ relevance feedback หรือรวมกันทั้งสองเทคนิค

Query expansion หมายถึง การเพิ่มทอมในคำสอบถามของระบบ การค้นคืนสารสนเทศ เพื่อให้การสืบค้นในรอบต่อไปมีผลลัพธ์ที่ดีขึ้น ซึ่งทอมที่เพิ่มขึ้นนั้นอาจจะมาจากการ feedback หรือไม่ก็แล้วแต่กรณี รวมทั้งอาจเปลี่ยนค่าน้ำหนักของทอมเพื่อความเหมาะสม ส่วน relevance feedback หมายถึง การบ่อนความเกี่ยวพันย้อนกลับ อันเป็นการใช้ประโยชน์จากข้อมูลย้อนกลับนี้ไปปรับเปลี่ยนคำสอบถามเก่ามาเป็นคำสอบถามใหม่ของการสืบค้นในแต่ละรอบ วิธีการเป็นไปได้ทั้งอัตโนมัติและกึ่งอัตโนมัติโดยใช้แรงคนร่วมปฏิบัติการด้วย

สำหรับแหล่งข้อมูลเพื่อการปรับเปลี่ยนคำสอบถามใหม่เป็นไปได้ใน 3 แหล่งคือ จาก Local Analysis, Global Analysis และ Query Reuse สำหรับ Local Analysis หมายถึง แหล่งข้อมูลที่ทำวิเคราะห์หามาจากชุดเอกสารรวมทั้งหมดและดึงทอมที่เป็นคำสัมพันธ์ (thesaurus) มาพิจารณา ส่วน Query Reuse หมายถึงการใช้ซ้ำของคำสอบถามจากแหล่งฐานข้อมูล query based

การสร้างโมเดล

โมเดล หมายถึง รูปแบบที่แสดงในเชิงตรรกะ (logical view) เพื่อจำลององค์ประกอบในระบบ รวมทั้งจำลองการปฏิบัติการ โดยทั่วไปรูปแบบอาจจะเป็นรูปภาพ หรือเป็นสัญลักษณ์ และมีสายโยงต่อกันไปมา รวมทั้งอาจเขียนเป็นสัญลักษณ์ทางคณิตศาสตร์และสมการคณิตศาสตร์เพื่อแสดงขบวนการปฏิบัติการที่ปรากฏทางกายภาพ

สำหรับโมเดลของระบบการสืบค้น มีรูปแบบเป็นการเฉพาะที่แสดงเชิงตรรกะเพื่อจำลององค์ประกอบต่างๆ ในระบบ เช่นตัวเอกสาร คลังเอกสาร ข้อมูลสารสนเทศที่ผู้ใช้งานต้องการหรือคำสอบถาม รวมทั้งปฏิบัติการเทียบเคียงเพื่อหาผลลัพธ์ ปฏิบัติการการจัดกลุ่ม การจัดหมวดหมู่ของคลังเอกสาร และอื่น ๆ เป็นต้น

โมเดลพื้นฐานหลักของระบบการค้นคืนข้อมูล จำนวน 2 รูปแบบ คือ CBM (Classical Boolean Model) และ VSM (Vector Space Model)

การประเมินประสิทธิภาพแบบ NDCG

Normalized Discount Cumulative Gain (NDCG) เป็นทฤษฎีที่ใช้ในการประเมินผลโปรแกรมค้นหาที่มีการให้ระดับคะแนนความเกี่ยวข้องของเอกสาร รวมทั้งพิจารณาลำดับของผลการค้นหา และ DCG เป็นการวัดความเหมาะสมของเอกสารโดยสนใจตำแหน่งหรือลำดับของเอกสาร ค่าระดับคะแนนที่ได้รับเป็นสะสมจากลำดับบนของรายการผลลัพธ์การค้นหาไปยังลำดับล่างของผลลัพธ์การค้นหา โดยค่าคะแนนจะลดลงเมื่อผลความพึงพอใจอยู่ในลำดับที่ต่ำ (Jarvelin, K., and Kekalainen, J. 2006) โดยสมการ NDCG เป็นดังนี้

$$NDCG_q = M_q \sum_{j=1}^k \frac{(2^{r(j)} - 1)}{\log(1 + j)} \quad (1)$$

เมื่อ k คือ ระดับหรือเกณฑ์ที่ใช้, $r(j)$ คือค่าลำดับความเกี่ยวข้องของเอกสารที่ได้จากประเมินโดยผู้ใช้, M_q คือ ค่าคงที่ที่เกิดจากความสมบูรณ์ในการจัดลำดับโดยมีค่ามากที่สุดคือ 1 ทั้งนี้ NDCG จะให้รางวัลกับเอกสารที่เกี่ยวข้องที่ปรากฏในลำดับของการจัดอันดับผลการค้นหาและลงโทษเอกสารที่ไม่เกี่ยวข้องโดยการลดคะแนน NDCG

งานวิจัยที่เกี่ยวข้อง

งานวิจัยจำนวนมากมุ่งพัฒนาประสิทธิภาพของผลการค้นหาโดยพยายามที่จะเปรียบเทียบ และประเมินผลโดยทำการเทียบกับโปรแกรมค้นหา (search engine) ที่มีการให้บริการอยู่ในปัจจุบัน ดังงานวิจัยต่อไปนี Gordon และ Pathank (1999) ได้ดำเนินการศึกษา และเปรียบเทียบประสิทธิภาพของโปรแกรมค้นหาจำนวน 8 เว็บไซต์ รวมทั้งศึกษาความสัมพันธ์ของแต่ละเว็บ, Thomas (2006) พยายามที่จะประเมินคุณภาพของโปรแกรมค้นหา ซึ่งผลการประเมินแสดงให้เห็นว่าการจัดอันดับตามคุณภาพจะนำไปสู่ผลลัพธ์ที่ดีขึ้น โดยคุณภาพของโมเดลจะเป็นประโยชน์ในการระบุคุณลักษณะที่สำคัญอาจเกิดขึ้น และลักษณะคุณภาพของเว็บ, Long และเพื่อน (2007) ได้ทำการประเมิน และเปรียบเทียบผลของโปรแกรมค้นหา จำนวน 3 เว็บไซต์ที่พัฒนาขึ้นด้วยภาษาจีนโดยทำการทดลองที่ใช้การประเมินจากผู้ใช้งาน นอกจากนี้ วิณาวดี ม่วงอันและ ศทาวุธ แก้วบรรจง (2016) นำเสนอระบบสืบค้นเว็บเซอร์วิสด้วยการพิจารณารายละเอียดเชิงความหมายของเว็บเซอร์วิส โดยทำการสร้างดัชนีจากส่วนประวัติการบริการ, ส่วนแบบจำลองการบริการ และส่วนพื้นฐานการบริการ ของเว็บเซอร์วิส ผลการทดลองเบื้องต้นสำหรับการใช้แบบจำลองการสืบค้นนี้ชี้ให้เห็นว่ามีค่าเฉลี่ยความแม่นยำที่ร้อยละ 63.60 %

ทั้งนี้ในปัจจุบันเครือข่ายแท็ก (social tag) หรือการสร้างแท็ก(tag) บนระบบเครือข่ายสังคมเป็นกระบวนการที่สามารถทำได้ง่ายและสะดวก รวมถึงการจัดการ tag ของผู้ใช้แต่ละบุคคล เนื้อหาที่ใช้ร่วมกันสำหรับการเรียกใช้ในอนาคต (Golder & Huberman, 2006) โดยเว็บ Flickr (<http://www.flickr.com>) คือเว็บไซต์เครือข่ายสังคมด้านการแบ่งปันภาพถ่ายที่เป็นที่นิยม ซึ่งมีบริการบริการด้าน อัปโหลดภาพถ่าย แท็ก และวิดีโอ ผู้ใช้สามารถเชิญผู้อื่น (เช่นเพื่อนหรือครอบครัว) เพื่อดูภาพและแท็กเช่นกัน Guy & Tonkin (2006) ได้ศึกษาถึงการใช้แท็กที่ผู้ใช้สร้างหรืออธิบายภาพและเอกสารอื่น ๆ เพื่อวัตถุประสงค์ในการจัดทำดัชนีได้รับการวิพากษ์วิจารณ์เนื่องจากมีความไม่ถูกต้องในการใช้ภาษา ซึ่งได้รับการยอมรับว่าจุดอ่อนนี้จะค่อนข้างหลีกเลี่ยงไม่ได้เนื่องจากมีการบริการที่ฟรี (ไม่สามารถควบคุม) และเป็นธรรมชาติที่เปิดของการติดแท็กทางสังคม (Hammond, Hannay, Lund, & Scott, 2005)

นอกจากนี้ยังมีผู้วิจัยท่านอื่นศึกษาถึงการสร้างดัชนีที่เกี่ยวข้องกับภาพถ่าย เช่น Pauline Rafferty และ Rob Hilderley (2007) ได้นำเสนออัลกอริทึมสำหรับสร้างดัชนี จำนวน 3 รูปแบบคือ 1) expert-led indexing, 2) author-generated indexing และ 3) user-orientated indexing โดยการทดลองเว็บ Flickr และดัชนีแบบ Democratic ซึ่งผลการทดลองพบว่าผู้ใช้สร้างเมตาตาต้าที่ความสอดคล้องกันไปในระดับของ

รายละเอียดความถูกต้อง รวมถึงการใช้งานของผู้ใช้สามารถนำไปสู่การจัดทำดัชนีโดยมีผู้ใช้เป็นศูนย์กลางได้ Jiyi Li, Qiang Ma, Yasuhito Asano, and Masatoshi Yoshikawa (2011) ได้ทำการศึกษาถึงการจัดอันดับผลการค้นหารูปภาพบนระบบเครือข่ายสังคม งานวิจัยนี้คาดหวังว่าการนำเครือข่ายแท็ก มาใช้ในการจัดอันดับรูปภาพสามารถเพิ่มประสิทธิภาพผลการค้นหาได้ ดังนั้นจึงทำการทดลองโดยเปรียบเทียบการจัดอันดับผลการค้นหาโดยใช้ข้อมูลจากเครือข่ายแท็ก และไม่ใช่ ซึ่งผลการทดลองพบว่าการจัดอันดับโดยการนำเครือข่ายแท็ก เข้ามาใช้สามารถเพิ่มประสิทธิภาพผลการค้นหารูปภาพได้ดี

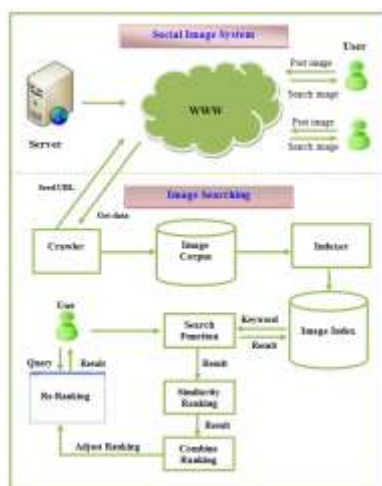
Schmidt และ Stock (2009) ได้ทำการศึกษาการบ่งบอกอารมณ์ของผู้ใช้จากรูปภาพ โดยทำการจัดเก็บดัชนีทางด้านอารมณ์โดยใช้รูปภาพในรูปแบบของอารมณ์พื้นฐาน เช่น โกรธ กลัว รังเกียจ มีความสุข เศร้า จากผู้ใช้ 763 คน เพื่อช่วยในการสืบค้นข้อมูลภาพถ่าย

มีผู้วิจัยได้ศึกษาถึงการนำปัจจัยที่แวดล้อมข้อมูลมาสร้างดัชนีโดยทำการศึกษากับระบบเครือข่ายสังคมด้านงานวิจัย ผู้วิจัยได้สร้างดัชนีจำนวน 3 ดัชนี คือ 1) “tag”(T), “title, 2) abstract”(TA) , 3) “tag, title and abstract”(TTA) และทำการเปรียบเทียบกับ CiteULike ผลการทดลองพบว่า ดัชนีที่สร้างโดยใช้ “tag, title and abstract”(TTA) มีประสิทธิภาพดีที่สุด (พิจิตรา จอมศรี 2009a, 2009b) นอกจากนี้ผู้วิจัยยังได้ทำการศึกษาถึงการสร้างดัชนีที่เหมาะสมสำหรับระบบเครือข่ายสังคมทางด้านรูปภาพ โดยพบว่าการนำคำสำคัญและคำบรรยายที่เกี่ยวข้องกับรูปภาพหรือ (DT) มาสร้างดัชนีสามารถเพิ่มประสิทธิภาพการค้นคืนเนื้อหาภาพได้

จากการทบทวนวรรณกรรมข้างต้น ทำให้พบว่างานวิจัยครั้งนี้มีความแตกต่างจากงานวิจัยอื่นโดยการนำปัจจัยทางด้านจำนวนการเข้าชมรูปภาพ และ จำนวนผู้ใช้ที่ชื่นชอบรูปภาพ มาช่วยในการจัดอันดับผลการค้นหารูปภาพโดยมุ่งเน้นความหลากหลายและความน่าเชื่อถือ ซึ่งนักวิจัยประยุกต์ใช้เทคนิคในการสร้างดัชนีแบบ “DT indexer” มาทำการเพิ่มประสิทธิภาพการจัดอันดับ

วิธีการดำเนินการวิจัย

วิธีการดำเนินการวิจัย แบบออกเป็น 4 กระบวนการหลัก คือ 1) กระบวนการในการพัฒนาระบบ 2) กระบวนการวัดระดับความเกี่ยวข้องของเอกสารโดยผู้ใช้งาน 3) กระบวนการประเมินผลตามทฤษฎี และ 4) ทดสอบสมมติฐานด้วยสถิติ



ภาพที่ 1 กระบวนการในการพัฒนาอัลกอริทึม สำหรับสร้างดัชนีด้านการค้นหารูปภาพบนระบบเครือข่ายสังคม

กระบวนการในการพัฒนาระบบ

กระบวนการในการพัฒนาระบบประกอบด้วย 5 กระบวนการดังภาพที่ 1 โดยมีรายละเอียดต่อไปนี้

1. *การดึงข้อมูล (Crawler Data)*: คือ การพัฒนาโปรแกรมค้นหาที่มีหน้าที่ในการดึงข้อมูลจากเว็บไซต์เครือข่ายสังคม และ โซเชียลบุ๊กมาร์ก มาจัดเก็บในฐานข้อมูล โดยตัว Crawler จะเก็บข้อมูลโดยผ่านทางเส้นทางที่เชื่อมโยงกันไปมาของเว็บต่าง ๆ ที่ต้องการดึงข้อมูล ซึ่งโปรแกรมรวบรวมรูปภาพมีหน้าที่ในการดึงข้อมูลรูปภาพจากเครื่องแม่ข่ายของระบบโซเชียลบุ๊กมาร์ก เช่น Flickr โดยมีการดึงข้อมูลต่อไปนี้ ค่าสำคัญของภาพ คำบรรยายได้ภาพ ชื่อภาพ จำนวนผู้เข้าชม รหัสรูปภาพ การเชื่อมโยงยังภาพจริง จำนวนผู้ชื่นชอบ การแสดงความคิดเห็น และวันที่โพสต์ ซึ่งข้อมูลเหล่านี้มีประโยชน์ต่อระบบในการตรวจสอบความสนใจของผู้ใช้และยังช่วยให้ระบบการดัชนีสำหรับแต่ละรูปภาพ

2. *ฐานข้อมูลรูปภาพ (Image Corpus)*: คือ ฐานข้อมูลที่ใช้ในการจัดเก็บรูปภาพที่ตัว Crawler ดึงมาจากระบบโซเชียลบุ๊กมาร์กโดยข้อมูลที่ใช้ในการทดลองมีรายละเอียดดังนี้

ตารางที่ 1 รายการข้อมูลที่ใช้ในการทดลอง

รายการ	รายละเอียด
แหล่งข้อมูล	Flickr
จำนวนงานวิจัย	63,345 ภาพ
ข้อมูลกลุ่มผู้ใช้	2,500 ผู้ใช้งาน

3. *ดัชนี (Indexer)*: กระบวนการในการสร้างดัชนีจะใช้เทคนิค TF-IDF ในการสร้างดัชนี TF-IDF คือ เป็นการพิจารณาความสามารถในการ แบ่งแยกของคำ (t_j) ในเอกสาร (d_i) จากความถี่ของคำในเอกสาร และ อัตราส่วนกลับของจำนวนเอกสารทั้งหมดกับจำนวนเอกสารที่มีคำ t_j ความสำคัญจะเพิ่มขึ้นตามสัดส่วนของจำนวนครั้งคำปรากฏอยู่ในเอกสาร แต่สามารถชดเชยตามความถี่ของคำที่มีอยู่ในฐานข้อมูล โดยทำการสร้างดัชนีโดยการผสมระหว่างสองปัจจัยคือคำอธิบายรูปภาพ และค่าสำคัญของรูปภาพ หรือ “DT indexer” ที่เป็น *Similarity ranking* จากนั้นจึงจัดเก็บลงในฐานข้อมูลดัชนี หรือ *DTIndex*

4. *การจัดอันดับตามความเหมือน (Similarity Ranking)*: คือ การเปรียบเทียบคำค้นหากับดัชนี DT ที่สร้างขึ้น คะแนนที่ได้จากการเปรียบเทียบถูกคำนวณจากสมการที่ (1) ซึ่งเป็นฟังก์ชันของ Apache Lucene พัฒนาโดย Hatcher และ Gospodnetic จากสมการ q หมายถึงคำค้นหา และ d หมายถึง เอกสารรูปภาพ

กำหนดให้

$$score(q, d) = \sum_{t \in q} (tf(t \in d) \times idf(t)^2 \times B_q \times B_d \times L) \times C \quad (2)$$

$tf(t \text{ in } d)$ คือ Term frequency หรือ ความถี่ของเทอม (t) ที่ปรากฏในเอกสาร (d)

$idf(t)$ คือ Inverse document frequency หรือ เป็นค่าที่ใช้บ่งบอกความสำคัญของเทอมที่เทียบกับทุกๆ เอกสารที่เก็บอยู่ในฐานข้อมูล

$boost(t.field \text{ in } d)$ คือ ค่า boost ที่ถูกสร้างขึ้นจากการสร้างดัชนี

$lengthNorm(t.field \text{ in } d)$ คือ กระบวนการ normalization value of a field,

$coord(q, d)$ คือ กระบวนการ coordination factor

$querynorm(q)$ คือ กระบวนการ normalization value for a query

5. การจัดอันดับแบบรวม (Combine Ranking): คือ ขั้นตอนในการนำ Static Ranking ซึ่งเป็นข้อมูลที่ได้จากผู้ใช้โดยตรง งานวิจัยนี้จึงทำการศึกษาการจัดอันดับแบบรวม โดยการนำดัชนีที่เหมาะสมที่สุด มาพิจารณาร่วมกับจำนวนผู้เข้าชมรูปภาพ (View: V) กับ จำนวนผู้ใช้ที่ชื่นชอบรูปภาพ (Favorite: F) โดยมีรายละเอียดดังนี้

1. จำนวนผู้เข้าชมรูปภาพ

ปัจจัยด้านจำนวนผู้เข้าชมรูปภาพ หรือ $VRank$ บ่งบอกถึงจำนวนผู้ใช้ที่เข้ามาสนใจรูปภาพนั้น โดยภาพที่มีจำนวนผู้เข้าชมมากกว่าจะมีค่าอันดับสูงกว่าภาพที่มีจำนวนผู้ชมน้อย ดังสมการที่ 3 กำหนดให้ V คือ คะแนนการจัดอันดับด้วยจำนวนผู้เข้าชมรูปภาพ x คือ จำนวนผู้เข้าชมของแต่ละภาพ เมื่อ $x_i = X$

$$V = \frac{x_i}{\max_j (x_j)} \quad (3)$$

โดยที่การแบ่งระดับคะแนนของจำนวนผู้เข้าชม

$$V = \begin{cases} 5; & \text{if } x = \geq 2001; \text{ range level is more over 2000} \\ 4; & \text{if } x = 2000 - 1501; \text{ range level is 1501 to 2000} \\ 3; & \text{if } x = 1500 - 1001; \text{ range level is 1001 to 1500} \\ 2; & \text{if } x = 1000 - 501; \text{ range level is 501 to 1000} \\ 1; & \text{if } x = 500 - 101; \text{ range level is 101 to 500} \\ 0; & \text{if } x = \leq 100; \text{ range level is least then 101} \end{cases}$$

2. จำนวนผู้ใช้ที่ชื่นชอบรูปภาพ

ปัจจัยด้านจำนวนผู้ชื่นชอบรูปภาพ หรือ $FRank$ บ่งบอกถึงจำนวนผู้ใช้ที่กดชื่นชอบรูปภาพนั้น โดยภาพที่มีจำนวนผู้ชื่นชอบมากกว่าจะมีค่าการจัดอันดับสูงกว่าภาพที่มีจำนวนผู้ชื่นชอบน้อย ดังสมการที่ 4 กำหนดให้ F คือ คะแนนการจัดอันดับด้วยจำนวนผู้เข้าชมรูปภาพ y คือ จำนวนผู้เข้าชมของแต่ละภาพ เมื่อ $y_i = Y$

$$F = \frac{Y_i}{\max_j (Y_j)} \quad (4)$$

โดยที่การแบ่งระดับคะแนนของผู้ใช้ที่ชื่นชอบรูปภาพ

$$Y = \begin{cases} 5; & \text{if } y \geq 101; \text{ range level is more over 101} \\ 4; & \text{if } y = 100 - 76; \text{ range level is 100 to 76} \\ 3; & \text{if } y = 75 - 51; \text{ range level is 75 to 51} \\ 2; & \text{if } y = 50 - 26; \text{ range level is 50 to 26} \\ 1; & \text{if } y = 25 - 1; \text{ range level is 25 to 1} \\ 0; & \text{if } y = 0; \text{ range level is not have Favorite} \end{cases}$$

นอกจากนี้ยังได้นำทฤษฎีถ่วงน้ำหนักมาทำการศึกษาค้นคว้าถึงสัดส่วนที่เหมาะสมที่สุด สำหรับการจัดอันดับรูปภาพ โดยให้ Let ω_c เป็นค่าการถ่วงน้ำหนักโดย $\{\omega_c = 0.9, 0.8, \text{ และ } 0.5\}$

การจัดอันดับตามความเหมือนโดยใช้ดัชนีที่สร้างขึ้นร่วมกับจำนวนผู้เข้าชมรูปภาพ (DTVRank) เป็นการรวมคะแนนที่ได้จากการจัดอันดับ SimRank ร่วมกับการจัดอันดับแบบ VRank ดังสมการที่ 5

$$DTVRank = (SimRank \times \omega_c) + (VRank \times (1 - \omega_c)) \quad (5)$$

การจัดอันดับตามความเหมือนโดยใช้ดัชนีที่สร้างขึ้นร่วมกับจำนวนผู้ที่ชื่นชอบรูปภาพ (DTFRank) เป็นการรวมคะแนนที่ได้จากการจัดอันดับ SimRank ร่วมกับการจัดอันดับแบบ Frank ดังสมการที่ 6

$$DTFRank = (SimRank \times \omega_c) + (FRank \times (1 - \omega_c)) \quad (6)$$

$$DTFRank = (SimRank \times 0.5) + (FRank \times (1 - 0.5)) \quad (7)$$

การจัดอันดับตามความเหมือนโดยใช้ดัชนีที่สร้างขึ้นร่วมกับจำนวนผู้เข้าชมรูปภาพและจำนวนผู้ที่ชื่นชอบรูปภาพ (DTVFRank) เป็นการรวมคะแนนที่ได้จากการจัดอันดับ SimRank ร่วมกับการจัดอันดับแบบ VRank ทั้งนี้โดย VFRank คือผลรวมของค่าคะแนนทั้ง 2 ปัจจัยรวมกันเท่ากับ 1 ดังสมการที่ 9

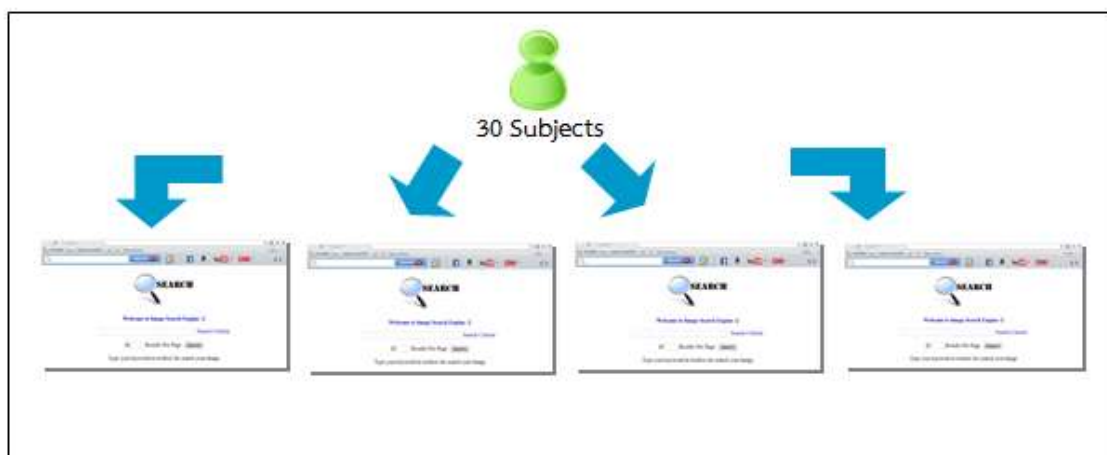
$$VFRank = (VRank \times 0.5) + (FRank \times (1 - 0.5)) \quad (8)$$

$$DTVFRank = (SimRank \times \omega_c) + (VFRank \times (1 - \omega_c)) \quad (9)$$

3).ผลลัพธ์การค้นหา (Search Result) คือ ขั้นตอนการแสดงผลการค้นหาได้มีการสร้างส่วนต่อประสานใหม่เพื่อป้องกันความอคติของผู้ใช้หรือผู้ประเมิน โดยผู้ประเมินระบบสามารถค้นหาภาพผ่านส่วนต่อประสานที่สร้างขึ้นจากอัลกอริทึมในการจัดอันดับแต่ละวิธี ผู้ประเมินสามารถดูผลลัพธ์การค้นหาได้ ซึ่งผลลัพธ์แต่ละลำดับประกอบด้วย คำบรรยายใต้ภาพ และคำสำคัญของภาพ พร้อมการเชื่อมโยงไปยังรูปภาพจริง

กระบวนการวัดระดับความเกี่ยวข้องของเอกสารโดยผู้ทดสอบระบบ

กระบวนการวัดระดับความเกี่ยวข้องของเอกสารโดยผู้ใช้งานเป็นขั้นตอนต่อจากการพัฒนาระบบเสร็จสิ้นโดยผู้วิจัยจะต้องประชาสัมพันธ์เพื่อรับสมัครผู้ใช้งาน 30 คน โดยกำหนดคำค้นหาที่จะใช้ในการค้นคืนสารสนเทศให้กับทดสอบระบบ จากนั้นผู้ทดสอบระบบจะต้องทำการประเมินผลจากผลลัพธ์ที่ได้ในแต่ละ คำค้นหาตามระดับคะแนนที่นักวิจัยกำหนด



ภาพที่ 2 กระบวนการในวัดระดับความพึงพอใจของผู้ใช้งานระบบ

ทั้งนี้คะแนนประเมินผลการค้นหาถูกกำหนดไว้ 5 ระดับ ดังนี้ ระดับ 5 คือ ผลลัพธ์เกี่ยวข้องกับคำค้นหามาก ระดับ 4 คือ ผลลัพธ์เกี่ยวข้องกับคำค้นหามาก ระดับ 3 คือ ผลลัพธ์เกี่ยวข้องกับคำค้นหาน้อยระดับ 2 คือ ผลลัพธ์เกี่ยวข้องกับคำค้นหาน้อยที่สุด ระดับ 1 คือ ผลลัพธ์ไม่เกี่ยวข้องกับคำค้นหา

กระบวนการประเมินผล

งานวิจัยนี้เลือกใช้ทฤษฎี NDCG ในการประเมินผลการทดลอง หลังจากที่ได้ผลการประเมินความพึงพอใจจากผู้ใช้งานโดยงานวิจัยครั้งนี้ได้ทำการประเมินผลค่า NDCG จากการจัดอันดับที่แตกต่างกันจำนวน 4 วิธี และนำค่าเฉลี่ยของ NDCG มาสมมติฐานด้วยทฤษฎีทางสถิติ

จากสมมติฐานหลักที่กำหนดขึ้นคือ ประสิทธิภาพของอัลกอริทึม สำหรับการจัดอันดับผลการค้นหา งานวิจัยแต่ละวิธีไม่แตกต่างกัน ดังนี้

1) สมมติฐานหลัก

H_0 : ค่าเฉลี่ย NDCG ของ ทั้ง 4 การจัดอันดับ คือ *SimRank* , *DTVRank* , *DTFRank* และ *DTVFRank* ไม่แตกต่างกัน

$$(\mu_{SimRank} = \mu_{DTVRank} = \mu_{DTFRank} = \mu_{DTVFRank})$$

2) สมมติฐานรอง

H_1 : ค่าเฉลี่ย NDCG ของ ทั้ง 4 การจัดอันดับ คือ *SimRank* , *DTVRank* , *DTFRank* และ *DTVFRank* แตกต่างกัน

$$(\mu_{SimRank} \neq \mu_{DTVRank} \neq \mu_{DTFRank} \neq \mu_{DTVFRank})$$

งานวิจัยนี้จึงทำการทดสอบสมมติฐานโดยใช้ทฤษฎีทางสถิติเข้ามาช่วยในการทดสอบ คือ ANOVA โดยการใช้ค่า NDCG ที่ได้จากระบบประเมินผล

ผลการวิจัย

ผลจากการคำนวณค่าเฉลี่ยการประเมินผลจากการใช้งานระบบซึ่งจะเห็นได้ว่าค่า NDCG เฉลี่ยของการการจัดอันดับตามความเหมือนโดยใช้ดัชนีที่สร้างขึ้นร่วมกับจำนวนผู้เข้าชมรูปภาพและจำนวนผู้ที่ชื่นชอบรูปภาพ *DTVFRank* ที่อัตราส่วน 50:50 มีค่ามากที่สุด แสดงดังตารางที่ 2

ตารางที่ 2 แสดงค่า NDCG เฉลี่ย 20 อันดับของอัลกอริทึม 4 วิธี

Rank no.	Average of NDCG Scores			
	<i>SimRank</i>	<i>DTVRank (50:50)</i>	<i>DTFRank (50:50)</i>	<i>DTVFRank (50:50)</i>
1	0.683851	0.722066	0.713724	0.735612
2	0.684061	0.706951	0.708336	0.731022
3	0.689697	0.701659	0.726659	0.730025
4	0.710210	0.711181	0.729395	0.729769
5	0.696722	0.705407	0.705407	0.727394
6	0.683427	0.683573	0.683573	0.724309
7	0.685951	0.686546	0.685649	0.722905
8	0.695729	0.696743	0.695428	0.721454
9	0.694902	0.696305	0.695388	0.721156
10	0.692595	0.694358	0.694333	0.721028
11	0.696932	0.700026	0.702123	0.716773
12	0.693958	0.696946	0.697953	0.713572
13	0.694594	0.697845	0.698836	0.710278
14	0.695137	0.698637	0.698737	0.709898
15	0.689295	0.692882	0.695996	0.704899
16	0.684158	0.687673	0.687344	0.703778
17	0.685951	0.68926	0.688201	0.693881
18	0.689026	0.690219	0.689567	0.696966
19	0.68715	0.691541	0.690658	0.692434
20	0.692717	0.680723	0.680112	0.691703

งานวิจัยนี้ทำการทดสอบสมมติฐานโดย ANOVA ซึ่งเป็นทฤษฎีทางสถิติเข้ามาช่วยในการทดสอบสมมติฐานที่กำหนดไว้ โดยจากการทดลองเบื้องต้นพบว่าความแปรปรวนในแต่ละอัลกอริทึม

ตารางที่ 3 แสดงการวิเคราะห์แบบ ANOVA

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	.008	2	.002	.027	.001
Within Groups	152.529	2347	.067		
Total	153.532	2349			

จากการวิเคราะห์ค่าตาราง Test of Homogeneity of Variance แสดงผลการทดสอบสมมติฐานตามข้อตกลงเบื้องต้น เรื่อง การเท่ากันของความแปรปรวนที่ระดับนัยสำคัญ กำหนดให้ $\alpha = .05$ ผลการวิเคราะห์ ค่าสถิติ Levene มีค่า Sig. ซึ่งมากกว่า $\alpha = .05$ จึงไม่ปฏิเสธ H_0 แสดงว่าความแปรปรวนของประชากรทั้ง 4 กลุ่มเท่ากัน และจากตาราง ANOVA ดังตารางที่ 3 เป็นการวิเคราะห์ความแตกต่างของค่าเฉลี่ยของ 4 วิธีของการจัดอันดับผลการค้นหาพบว่าค่า Sig. น้อยกว่าระดับนัยสำคัญที่กำหนด แสดงว่าค่าเฉลี่ยอย่างน้อย 1 กลุ่มแตกต่างกัน และเมื่อตรวจสอบด้วยตาราง Multiple Comparisons โดยใช้วิธี Dunnett T3 พบว่าการจัดอันดับโดยใช้ *DTVFRank* มีความแตกต่างจากวิธีการจัดอันดับแบบอื่นและให้ผลลัพธ์ที่มีประสิทธิภาพสูงสุด

สรุปผลการวิจัย

งานวิจัยฉบับนี้ได้นำเสนอแนวทางในการจัดเรียงลำดับผลการค้นหาใหม่ ที่เหมาะสมสำหรับโซเชียลบุคมาร์กด้านรูปภาพ โดยการนำปัจจัยที่ผู้ใช้เป็นผู้กำหนดมาสนับสนุนและเพิ่มประสิทธิภาพในการจัดอันดับผลการค้นหา จึงได้ทำการพัฒนาเทคนิคการจัดอันดับใหม่ จำนวน 4 เทคนิค คือ การจัดอันดับตามความเหมือน (*SimRank*) การจัดอันดับตามความเหมือน ร่วมกับจำนวนผู้เข้าชมรูปภาพ (*DTVRank*) การจัดอันดับตามความเหมือนร่วมกับจำนวนผู้ที่ชื่นชอบรูปภาพ (*DTFRank*) การจัดอันดับตามความเหมือนร่วมกับจำนวนผู้เข้าชมรูปภาพและจำนวนผู้ที่ชื่นชอบรูปภาพ (*DTVFRank*) เพื่อทำการศึกษาถึงเทคนิคที่มีประสิทธิภาพมากที่สุด

โดยงานวิจัยนี้ได้ทำการทดสอบระบบจากผู้ทดสอบระบบจำนวน 30 คน ผู้ทดสอบแต่ละคนจะทำการค้นหาตามคำค้นหา จำนวน คนละ 3 คำค้นหา โดยมีคะแนนการประเมินผลการค้นหาจำนวน 5 ระดับ

ผู้วิจัยได้ดำเนินการประเมินผลการทดลองโดยใช้เทคนิค NDCG ในการประเมินผล โดยผลการทดลองพบว่า การจัดอันดับแบบ *DTVFRank(50:50)* สามารถที่จะเพิ่มประสิทธิภาพการจัดอันดับผลการค้นหาได้อย่างมีประสิทธิภาพ ทั้งนี้งานวิจัยที่จะทำการพัฒนาต่อไปในอนาคตเป็นการเพิ่มประสิทธิภาพโดยการสร้างประวัติผู้ใช้เพื่อทำการเพิ่มประสิทธิภาพผลการค้นหาสำหรับงานวิจัยบนระบบเครือข่ายสังคมและโซเชียลบุคมาร์กต่อไป

อภิปรายผล

จากผลการทดลองพบว่า *DTVFRank(50:50)* สามารถเพิ่มประสิทธิภาพผลการค้นหาบนระบบโซเซียลบุคมาร์กด้านรูปภาพได้ โดยการจะทำให้ผลการจัดอันดับของรูปภาพที่ตรงกับคำค้นหามากที่สุดอยู่ในอันดับต้นของผลการค้นหา ซึ่งอาจสืบเนื่องมาจากการวิจัยนี้เป็นการนำข้อมูลที่ได้มาจากพฤติกรรมการใช้งานจริงของผู้ใช้งาน โดยที่ภาพที่มีผู้เข้าชมจำนวนมาก และภาพที่มีผู้ใช้ชื่นชอบจำนวนมาก ย่อมเป็นรูปที่ตรงกับคำค้นหา มากกว่ารูปที่มีผู้เข้าชมจำนวนน้อย ซึ่งผลการทดลองพบว่าการจัดอันดับโดยใช้อัตราส่วนดังกล่าวทำให้ประสิทธิภาพผลการค้นหาเอกสารงานวิจัยมีประสิทธิภาพดีขึ้น

นอกจากนี้เนื้อหาที่เกี่ยวข้องกับรูปภาพ เช่น *DTVF* ซึ่งเป็นเทคนิคการจัดอันดับแบบรวม ถือเป็นสิ่งสำคัญสำหรับเป็นตัวแทนเนื้อหาของเอกสาร รวมทั้งการนำปัจจัยที่เกี่ยวข้องกับผู้เข้ามาช่วยในการจัดอันดับผลการค้นหาสามารถเพิ่มประสิทธิภาพผลการค้นหารูปภาพได้

เอกสารอ้างอิง

ภาษาไทย

ศุภชัย ตั้งวงศ์สานต์. 2551. ระบบการจัดเก็บและการสืบค้นสารสนเทศด้วยคอมพิวเตอร์. กรุงเทพฯ: โรงพิมพ์พิทักษ์การพิมพ์, 2551.

วิณาวดี ม่วงอัน, คหาฐ แก้วบรรจง. 2559. ระบบการสืบค้นเว็บเซอร์วิสโดยการใช้เวกเตอร์สเปซโมเดล:

วารสารวิชาการ Veridian E-Journal Science and Technology Silpakorn University
สาขาวิทยาศาสตร์และเทคโนโลยี, 3 , 5 (กันยายน-ตุลาคม) 2559.

ภาษาต่างประเทศ

Krystyna K. Matusiak. (2006). "Towards user-centered indexing in digital image collections."

Journal of OCLC Systems & Services: International digital library perspectives
:283-293.

Schmidt, Stefanie, and Wolfgang G. Stock. (2009). "Collective indexing of emotions in images.

A study in emotional information retrieval." *Journal of the American Society for Information Science and Technology*, 60, 5(May):863–876.

Guy, M., and Tonkin, E. (2006). "Folksonomies: Tidying up tags?." *D-Lib Magazine*, 12(1).

Accessed January 29, 2017. Available from

<http://www.dlib.org/dlib/january06/guy/01guy.html>

Scott A. Golder (2006). "Usage patterns of collaborative tagging systems." *Journal of Information Science*, 32, 2 :198–208.

- Jomsri, Pijitra , Sanguansintukul, Siripun, and Choochaiwattana, Worasit.(2009). “Improve Research paper Searching with social tagging-A Preliminary Investigation.” In **the eight international symposium on natural language processing (SNLP2009)**, October 20-21, 2009. Bangkok.
- Jomsri, Pijitra , Sanguansintukul, Siripun, and Choochaiwattana, Worasit.(2009). “A Comparison of Search Engine Using “Tag Title and Abstract” with CiteULike – An Initial Evaluation.” In **the 4th international conference for internet technology and secured transactions (ICITST-2009)**, November 9-12 ,2009. London ,UK.
- Long, Haiquan, Lv, Benfu, and Zhao, Tianbo. (2007). “Evaluate and Compare Chinese Internet Search Engines Based on Users'Experience.” **Wireless Communications, Networking and Mobile Computing**, September 21-25, 2007.
- Gordon, Michael, and Pathak, Praveen. (1999). “Finding information on the World Wide Web: the retrieval effectiveness of search engines.” **Information Processing and Management: an International Journal**, 35, 2 : 141–180.
- Mandl, Thomas. (2006). “Implementation and Evaluation of a Quality-Based Search Engine.” In **17th ACM conference on hypertext and hypermedia (HT '06)**, Odense, Denmark, 2006 August 22-25.
- Li, Jiyi, Ma, Qiang, Asano, Yasuhito, and Yoshikawa, Masatoshi. (2011). “Ranking Content-Based Social Images Search Results with Social Tags.” In **the seventh asia information retrieval societies conference (AIRS 2011)**, :147-156, 2011.