



รายงานวิจัยฉบับสมบูรณ์

โครงการย่อยที่ 1

ระบบเว็บเชิงความหมายเพื่อการจัดการและสนับสนุนการประมวลแบบขนานสำหรับการ
ท่องเที่ยวเชิงสุขภาพ อำเภอหัวหิน

(Semantic Web Systems for Management and Parallel Processing Support for
Health Tourism in Hua Hin)

โดย รศ. ดร. จันทนา จันทราพรชัย และคณะ

31 มกราคม พ.ศ. 2558

เลขที่สัญญา RDG5650033

รายงานวิจัยฉบับสมบูรณ์

โครงการย่อย 1

ระบบเว็บเชิงความหมายเพื่อการจัดการและสนับสนุนการประมวลผลแบบขนานสำหรับการ
ท่องเที่ยวเชิงสุขภาพ อำเภอหัวหิน

(Semantic Web Systems for Management and Parallel Processing Support for
Health Tourism in Hua Hin)

รศ. ดร. จันทนา จันทราพรชัย มหาวิทยาลัยเกษตรศาสตร์

นางสาว ชิดชนก โชคสุชาติ มหาวิทยาลัยศิลปากร

ชุดโครงการ แผนงานวิจัยระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพ

กรณีศึกษาอำเภอหัวหิน

โดยอาศัยออนโทโลยีและเทคโนโลยีการประมวลผลแบบขนานบนกลุ่มเมฆ

สนับสนุนโดยสำนักงานคณะกรรมการวิจัยแห่งชาติ (วช.) และ

สำนักงานกองทุนสนับสนุนการวิจัย (สกว.)

(ความเห็นในรายงานนี้เป็นของคณะผู้วิจัย วช.-สกว. ไม่จำเป็นต้องเห็นด้วยเสมอไป)

กิตติกรรมประกาศ

คณะผู้จัดทำโครงการวิจัย เรื่อง ระบบเว็บเชิงความหมายเพื่อการจัดการและสนับสนุนการประมวลแบบขนานสำหรับการท่องเที่ยวเชิงสุขภาพ อำเภอหัวหิน ขอขอบพระคุณบุคลากรต่างๆ ที่เกี่ยวข้องกับโครงการวิจัยนี้ โดยเฉพาะผศ. ดร. ฤชงค์ อุทโยภาส ที่ปรึกษาแผนงานวิจัยที่ให้คำแนะนำในการวิจัย และขอขอบคุณ ภาควิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเกษตรศาสตร์ และภาควิชาคอมพิวเตอร์ คณะวิทยาศาสตร์ มหาวิทยาลัยศิลปากร ที่ให้การสนับสนุนด้านพื้นที่ และอุปกรณ์ในการดำเนินการ รวมทั้งบุคลากรเจ้าหน้าที่ธุรการที่ช่วยงานในโครงการวิจัยนี้

คณะผู้วิจัยขอขอบพระคุณสำนักงานกองทุนสนับสนุนการวิจัย (สกว.) โครงการวิจัยการท่องเที่ยวและบริการ ขอขอบพระคุณสำนักงานคณะกรรมการวิจัยแห่งชาติ (วช.) ที่ให้ทุนสนับสนุนในการการทำวิจัยครั้งนี้ ขอขอบพระคุณบริษัท กสท. โทรคมนาคมแห่งประเทศไทย จำกัด ที่ให้การสนับสนุนการใช้แพลตฟอร์ม IRIS เพื่อทดสอบระบบ

ขอขอบคุณผู้มีส่วนร่วมในงานวิจัยนี้ ได้แก่ คุณนพพร วุฒิกุล คุณมนตรี ชูภู และเทศบาลเมืองหัวหินที่ให้ความร่วมมือในการสนับสนุนการนำผลงานจากแผนงานวิจัยนี้ไปใช้ประโยชน์ ในการให้คำแนะนำรวมทั้งประเมินส่วนต่างๆ ของโครงการ

สุดท้ายคณะผู้วิจัยขอขอบคุณทีมงานและผู้ช่วยวิจัยทุกท่าน รวมทั้งบุคลากรผู้มีพระคุณที่ให้การอุปถัมภ์ทั้งร่างกายและใจแก่คณะผู้วิจัยทุกคนที่ทำให้งานวิจัยนี้สำเร็จลุล่วงไปด้วยดี

คณะผู้วิจัย

บทสรุปผู้บริหาร

ในแผนพัฒนาเศรษฐกิจและสังคมแห่งชาติ ฉบับที่ 11 (พ.ศ.2554 - พ.ศ.2559) การพัฒนาศักยภาพระดับคุณภาพบริการในอุตสาหกรรมบริการและการท่องเที่ยวโดยเฉพาะจากการท่องเที่ยวเชิงสุขภาพสามารถสร้างมูลค่าเพิ่มให้กับสาขาบริการที่เป็นมิตรกับสิ่งแวดล้อม การพำนักระยะยาวและดูแลผู้สูงอายุในท้องถิ่นด้วยความคิดสร้างสรรค์และนวัตกรรมแบบใหม่ เทคโนโลยีปัจจุบันในสาขาคอมพิวเตอร์ที่จะมีบทบาทในการปรับปรุงการให้บริการนี้ ได้แก่ การประมวลผลแบบขนานบนกลุ่มเมฆ และการบริการเว็บเชิงความหมายด้วยออนโทโลยีนั้น ซึ่งจะมีส่วนช่วยในประเด็นต่างๆ เช่น

- 1) ความรวดเร็วในการสืบค้นข้อมูลด้วยการประมวลผลแบบขนาน
- 2) คำตอบจากการสืบค้นที่ตรงตามความต้องการของผู้ค้นหาข้อมูล ด้วยการค้นหาเชิงความหมายในขอบเขตโดเมนความรู้เฉพาะทาง
- 3) ความสามารถในการขยายฐานข้อมูลผู้ให้บริการการท่องเที่ยวในรูปแบบต่างๆ ที่ไม่ยึดติดพื้นที่ทางกายภาพด้วยเทคโนโลยีกลุ่มเมฆ

ในโครงการวิจัยนี้ ผู้วิจัยจึงขอเสนอ ต้นแบบระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพโดยอาศัยการผสมผสานกับเทคโนโลยีนี้สำหรับกรณีศึกษา อำเภอหัวหิน ในโครงการย่อยนี้จะเน้นการนำเสนอข้อมูลเพื่อให้สอดคล้องกับการประมวลผลแบบขนานและประมวลผลแบบกลุ่มเมฆ พัฒนาวิธีการประมวลผลแบบขนานที่เหมาะสม มีประสิทธิภาพเพื่อการค้นหาข้อมูลที่รวดเร็ว

วัตถุประสงค์

เพื่อพัฒนาระบบการค้นหาข้อมูล สำหรับให้บริการข้อมูลการท่องเที่ยวเชิงสุขภาพสำหรับอำเภอหัวหิน โดยเน้นที่การประมวลผลข้อมูลแบบขนาน และสำหรับกลุ่มเมฆอย่างมีประสิทธิภาพ

ขอบเขตการวิจัย

พัฒนาระบบบริการข้อมูลท่องเที่ยวเชิงสุขภาพ ใช้เทคโนโลยีเว็บเชิงความหมาย และการประมวลผลวิธีแบบขนาน ทำงานได้บนกลุ่มเมฆ

ขั้นตอนการดำเนินงาน

1. สำรวจความต้องการซอฟต์แวร์ในปัจจุบัน
2. การศึกษางานวิจัยที่เกี่ยวข้องด้านการประมวลผลแบบขนาน และกลุ่มเมฆ
3. พัฒนาตัวแบบของระบบการประมวลผลแบบกลุ่มเมฆ
4. พัฒนาอัลกอริธึมการประมวลผลวิธีแบบขนาน และการแทนค่าข้อมูล
5. พัฒนาโปรแกรมการประมวลผลวิธีแบบขนาน
6. การออกแบบหน้าจอ การติดต่อผู้ใช้ ระบบออนโทโลยี
7. ออกแบบฐานข้อมูลระบบ ได้แก่ ข้อมูลผู้ใช้งาน ข้อมูลสถานที่ท่องเที่ยว
8. พัฒนาระบบจัดการการค้นหา พัฒนาโปรแกรม ทดลองใช้

9. เปรียบเทียบผลการค้นหาวิธี ด้านประสิทธิภาพ
10. ประเมินผลการใช้งานกับผู้ใช้งานบทบาทต่างๆ กัน
11. เขียนสรุปรายงานและเอกสารวิชาการ
12. จัดเตรียมการอบรมผู้ใช้ จัดทำคู่มือผู้ใช้งาน

ผลการดำเนินงาน

ผลการวิจัยดำเนินงานแบ่งเป็น 3 ส่วน ดังนี้ ส่วนแรกเกี่ยวกับการสำรวจลักษณะของซอฟต์แวร์ที่ใช้ในปัจจุบันของหน่วยงาน และข้อมูลที่หน่วยงานใช้อยู่ ส่วนที่ 2 เป็นผลการประเมินประสิทธิภาพด้านเวลาของระบบ เว็บบางความหมาย การประมวลผลแบบขนาน ส่วนที่ 3 เป็นผลการประเมินประสิทธิภาพด้วยเวลาสำหรับการประมวลผลการดึงข้อมูล และค้นหาข้อมูลแบบขนาน ส่วนการประเมินผู้ใช้งานนั้นอยู่ในภาคผนวกของรายงาน แผนงานวิจัย จึงไม่ขอกล่าวถึงในที่นี้

1. การสำรวจข้อมูลของเทศบาลด้านลักษณะหน่วยงาน ข้อมูลที่มีและซอฟต์แวร์ต่างๆ

คณะผู้วิจัยได้ สอบถามข้อมูลซอฟต์แวร์ที่ใช้ภายในหน่วยงานเทศบาล และการทำงานของหน่วยงาน แหล่งข้อมูลต่างๆ ของการท่องเที่ยวที่ทางเทศบาลใช้อยู่ ด้วยการสัมภาษณ์และแบบสอบถาม โดยรายละเอียดอยู่ในรายงาน แผนการวิจัย

ข้อมูลเกี่ยวข้องกับหน่วยงานเทศบาลเมืองหัวหิน

ผู้ที่เกี่ยวข้องกับโครงการนี้ในหน่วยงานของเทศบาลเมือง หัวหิน คือ นายกเทศมนตรีมอบหมายงานให้รองนายกเทศมนตรีด้านการท่องเที่ยว คือคุณมนตรี ชูภู เป็นผู้อนุมัติการประเมินระบบนี้ รองนายกเทศมนตรีผู้รับผิดชอบด้านการท่องเที่ยว ได้มอบหมายฝ่ายอำนวยการ คือคุณเบญจวรรณ เมธาวีกุล มีหน้าที่ควบคุมดูแล และรับผิดชอบการปฏิบัติงานในหน้าที่ที่เกี่ยวข้องกับงานส่งเสริมการท่องเที่ยว มีหัวหน้างานคือ คุณอุกฤษฏ์ ป้อทองคำ มีพนักงาน 7 คน สำหรับงานส่งเสริมการท่องเที่ยว มีหน้าที่เกี่ยวกับ

- งานส่งเสริมและเผยแพร่การท่องเที่ยวของเทศบาลภายในประเทศ
- งานควบคุมกิจการบริการท่องเที่ยวภายในประเทศของเทศบาลและหน่วยงานอื่น
- งานต้อนรับ แนะนำ อำนวยความสะดวกและบริการนำเที่ยว รวมทั้งดำเนินธุรกิจการท่องเที่ยว
- งานประชาสัมพันธ์ ให้คำแนะนำเผยแพร่ความรู้เกี่ยวกับแหล่งท่องเที่ยวหรือกิจการ ที่เกี่ยวกับการท่องเที่ยวภายในจังหวัดและเทศบาล
- งานจัดทำรายงานและสถิติการเคลื่อนไหวเกี่ยวกับธุรกิจการท่องเที่ยว
- งานอื่นที่เกี่ยวข้องหรือตามที่ได้รับมอบหมาย

ข้อมูลที่ได้จากการสำรวจความต้องการและข้อมูลที่มีอยู่

รายละเอียดของการสำรวจความต้องการซอฟต์แวร์จากหน่วยงานเทศบาลเมืองหัวหินสำหรับการสำรวจข้อมูลในพื้นที่นั้นได้ทำ 2 ครั้ง นั่นคือ ครั้งที่ 1 เมื่อวันที่ 23 สิงหาคม พ.ศ. 2556 และได้กำหนดวันนัดหมายเพื่อพูดคุยรายละเอียดกันในเดือนกันยายน แต่ได้มีการเลื่อนเป็นครั้งที่ 2 วันที่ 4 ตุลาคม พ.ศ. 2556 ดังรายละเอียดดังต่อไปนี้

ระดับผู้บริหารนั้นมีความต้องการใช้ข้อมูลที่ผ่านการรวบรวมอย่างเป็นระบบ เนื่องจากการนำเสนอข้อมูลจากหน่วยงานต่อสาธารณชนนั้นควรมีรายละเอียดรอบด้านมากขึ้น

การสัมภาษณ์ผู้อำนวยการฝ่ายอำนวยการ และสำรวจจากเทศบาลเมืองหัวหิน หัวหน้างานส่งเสริมการท่องเที่ยว และเจ้าหน้าที่การท่องเที่ยวแห่งประเทศไทยในเขตเมืองหัวหินและพนักงานปฏิบัติการ คุณ เกษมศักดิ์ ปัดออด มีแบบสำรวจความต้องการการใช้งานซอฟต์แวร์ของระบบการจัดการท่องเที่ยวเชิงสุขภาพ สรุปผลการสำรวจได้ดังนี้

- ระดับหน่วยงานส่งเสริมการท่องเที่ยวมีศูนย์บริการนักท่องเที่ยว 2 แห่งคือเทศบาลเมืองหัวหินและหอนาฬิกา
- การรวบรวมข้อมูลเป็นการทำงานแบบระบบมือ
- การค้นหาข้อมูลเกี่ยวกับที่พัก ร้านอาหาร สถานที่แอ่งภัย ส่วนมากใช้ระบบเดินสำรวจหรือค้นหาทางเว็บไซต์เฉพาะที่ ในขณะที่สถานพยาบาล สปา นวดแผนไทย อายุรเวช นั้นไม่เคยสำรวจ
- การติดต่อสื่อสารระหว่างหน่วยงานเทศบาลและที่พักชั่วคราวอาทิโรงแรมในพื้นที่มักใช้โทรศัพท์และการเดินเอกสาร สำหรับการใช้งานโปรแกรมจองโรงแรมหรือเว็บไซต์เกี่ยวกับการท่องเที่ยว เช่น Agoda, TripAdvisor, AtSiam นั้น ไม่เคยใช้ สำหรับโปรแกรมการค้นหาเว็บไซต์นั้นใช้กูเกิล ส่วนเครือข่ายสังคมออนไลน์นั้นใช้เฟสบุ๊กเพื่อกระจายข่าวสารเทศบาลและการท่องเที่ยว
- สำหรับเครื่องมือที่ใช้ค้นหาผ่านอินเทอร์เน็ตคือคอมพิวเตอร์ส่วนบุคคลที่เชื่อมต่ออินเทอร์เน็ตทั้งแลนและระบบไร้สาย
- การทำรายงานนำเสนอผู้บริหารประกอบด้วยเชื้อชาติของนักท่องเที่ยวและจำนวนผู้สอบถามเป็นรายวันรายเดือนและรายปีผ่านโปรแกรมไมโครซอฟต์เอ็กเซล
- สำหรับเอกสารเผยแพร่ข้อมูลจัดทำเป็นแผ่นพับสองภาษาคือภาษาไทยและภาษาอังกฤษและมีการปรับปรุงประมาณปีละครั้ง
- สิ่งที่ต้องการจากระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพคือข้อมูลโรงแรมเบื้องต้นเพิ่มเติมจากการลงพื้นที่เองและจากกูเกิล ประโยชน์ที่จะนำไปใช้นอกจากรายงานสำหรับผู้บริหารคือ เป็นแหล่งข้อมูลสำหรับนักวิจัย นักเรียน นักศึกษา ในกรณีที่มีการร้องขอข้อมูลอย่างละเอียด
- ข้อมูลที่ร้องขอเพิ่มเติมคือหมายเลขโทรศัพท์และอีเมลของสถานทูต ข้อมูลสำนักงานตรวจคนเข้าเมืองเพื่อต่อวีซ่า
- ต้องการใช้ระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพ เพื่อรวบรวมข้อมูลร้อยละ 45 เพื่อตรวจสอบข้อมูลร้อยละ 45 ใช้ประกอบรายงานสำหรับผู้บริหารร้อยละ 10
- หากโปรแกรมเสร็จแล้วมีความต้องการอบรมเพื่อใช้งานโปรแกรม

นอกจากนี้ได้มีการสอบถามข้อมูลการวางแผนท่องเที่ยวจากสำนักงานการท่องเที่ยวแห่งประเทศไทย สำนักงานประจวบคีรีขันธ์วันที่ 23 สิงหาคม พ.ศ. 2556 ในเรื่องการวางแผนการท่องเที่ยวสำหรับ

นักท่องเที่ยวทั้งชาวไทยและชาวต่างประเทศซึ่งมีการเก็บข้อมูลนักท่องเที่ยวที่คล้ายคลึงกับเทศบาลเมืองคือ เชื้อชาติและจำนวนนักท่องเที่ยวที่มาสอบถามรายวัน สำหรับข้อมูลรายได้และนักท่องเที่ยวโดยรวมของภาคกลางจะได้รับมาจากการท่องเที่ยวแห่งประเทศไทยสำนักงานใหญ่การวางแผนท่องเที่ยวในส่วนข้อมูลทั่วไปและกิจกรรมการท่องเที่ยวที่เผยแพร่บนเจ้าหน้าที่เป็นผู้รวบรวมขึ้นเองส่วนเครื่องมือค้นหาเว็บไซต์นั้นใช้งานผ่านกูเกิลสำหรับรายละเอียดของหน้าที่สำนักงาน กิจกรรมการท่องเที่ยว รายละเอียดนักท่องเที่ยวและแผนการท่องเที่ยวมีให้ตามแผ่นพับของ ททท.

ข้อมูลจากหน่วยงานเพื่อตรวจสอบความถูกต้อง

สำหรับสถานบริการสปาหรือนวดแผนโบราณที่เป็นเป้าหมายของงานวิจัย ในกรณีที่เป็นโรงแรมหรือที่พักขนาดใหญ่ที่มีอยู่ใน Agoda หรือ TripAdvisor นั้น จะถูกดึงเข้าสู่ระบบโดยโปรแกรมคอมพิวเตอร์ อย่างไรก็ตามยังมีสถานบริการริมชายหาดและสถานบริการเชิงสุขภาพขนาดเล็ก ซึ่งผู้อำนวยการฝ่ายอำนวยความสะดวกให้ความรู้ในไว้ดังนี้

- กรณีที่เป็นการนวดริมชายหาดของเอกชนนั้นสามารถตรวจสอบได้กับเทศกิจ
- กรณีที่เป็นร้านขนาดเล็ก จะมีการขึ้นทะเบียนกับกองสาธารณสุข สิ่งที่ได้คือชื่อร้านและที่อยู่
- สำหรับพิกัดที่ตั้งนั้นงานแผนที่ภาษีมีบันทึกไว้ในระบบพิกัดกริดแบบ UTM รายชื่อผู้เสียภาษีและที่อยู่
- เมื่อต้องการนำมาตรวจสอบกับระบบที่ผู้วิจัยเก็บไว้เพื่อให้ได้ชื่อร้านและพิกัดละติจูด/ลองจิจูด ต้องนำข้อมูลมาจากทั้ง 3 หน่วยงานนี้มาตรวจสอบร่วมกัน

2. ผลการประเมินประสิทธิภาพ

การวัดประสิทธิภาพด้านเวลา มี 2 แบบ แบบแรกเป็นการวัดประสิทธิภาพด้านเวลาการค้นหาด้วยเว็บเชิงความหมายโดยค้นหาในออนไลน์ และแบบที่สองเป็นการค้นหาสายอักขระด้วยการประมวลผลแบบขนานด้วยหน่วยประมวลผลกราฟิก

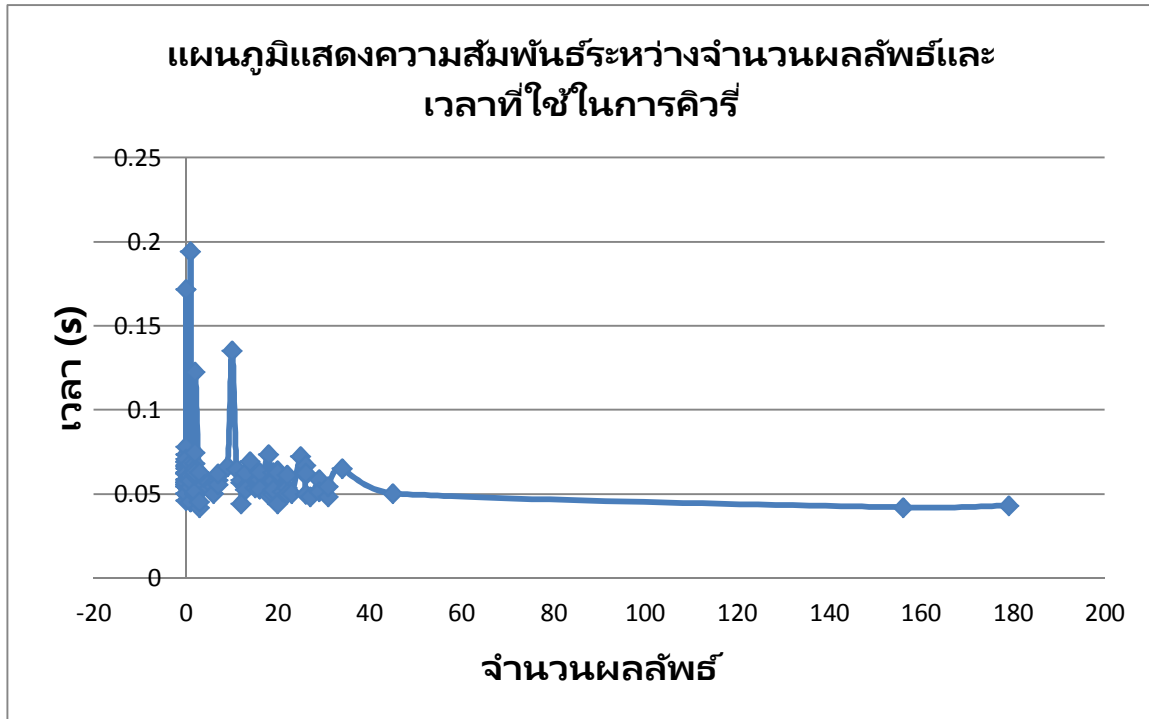
เวลาการค้นหาด้วยเว็บเชิงความหมาย

ในการทดสอบประสิทธิภาพ คณะผู้วิจัยได้ทดสอบวัดความเร็วคิวรีจากจำนวนคิวรี 100 ตัวอย่าง ทั้งที่ไม่มี ความซับซ้อน มีความซับซ้อนน้อยไปจนถึงตัวอย่างที่มีความซับซ้อนมากเพื่อวัดเวลาที่ใช้ในการคิวรีข้อมูล ตารางที่ 1 แสดงรายละเอียดข้อมูลออนไลน์

ตารางที่ 1 ข้อมูลออนโทโลยี

ข้อมูลทั้งหมดในออนโทโลยี	จำนวน
Axiom (ข้อกำหนดความสัมพันธ์)	24663
Logical axiom (ความสัมพันธ์เชิงตรรกะ)	21094
Class (คลาส)	229
Object Property (ความสัมพันธ์)	33
Data Property (คุณสมบัติ)	29
Individual (สมาชิกของคลาส)	3008
ข้อจำกัดความสัมพันธ์ในคลาส	จำนวน
SubClassOf (คลาสย่อย)	227
EquivalentClasses (คลาสเทียบเท่า)	19
DisjointClasses (คลาสที่ไม่เกี่ยวข้องกัน)	236
Hidden GCI (General Concept Inclusions)	13
ข้อจำกัดความสัมพันธ์ใน Object Property	จำนวน
SubObjectPropertyOf (ความสัมพันธ์ย่อย)	2
EquivalentObjectProperties (ความสัมพันธ์เทียบเท่า)	2
InverseObjectProperties (ความสัมพันธ์ที่แบบผกผัน)	16
FunctionalObjectProperty (ความสัมพันธ์ระหว่างกันแบบหนึ่งต่อหนึ่ง)	2
TransitiveObjectProperty (ความสัมพันธ์ที่มีการส่งต่อไปเรื่อยๆ)	4
ObjectPropertyDomain (สมาชิกภายในเซตข้อมูลที่ใช้จับคู่กับอีกเซตหนึ่ง)	33
ObjectPropertyRange (เซตข้อมูลที่ถูกจับคู่โดยความสัมพันธ์)	26
ข้อจำกัดความสัมพันธ์ใน Annotation	จำนวน
AnnotationAssertion (คำอธิบายประกอบคุณสมบัติ ความสัมพันธ์ และข้อมูล)	271
ข้อจำกัดความสัมพันธ์ใน Individual	
ClassAssertion (การกำหนดสมาชิกในคลาส)	3076
ObjectPropertyAssertion (การกำหนดความสัมพันธ์ของข้อมูล)	9943
DataPropertyAssertion (การระบุคุณสมบัติของข้อมูล)	7454

ผลการทดลองวัดเวลาที่ใช้ในการค้นหาพบว่า แต่ละตัวอย่างที่ใช้ทดลองควิรีใช้เวลาในการค้นหาใกล้เคียงกัน และใช้เวลาไม่มาก โดยเฉลี่ยใช้เวลาน้อยกว่า 1 วินาที (เวลาโดยเฉลี่ย 0.061618 วินาที) ทั้งนี้ความซับซ้อนของคำค้นและจำนวนผลลัพธ์ของการสืบค้นไม่ส่งผลให้ใช้เวลาในการค้นหาแตกต่างกันมาก



ภาพที่ 2 แผนภูมิแสดงความสัมพันธ์ระหว่างคำค้นและเวลาที่ใช้ในการคิวรี แสดงให้เห็นว่าความซับซ้อนของคำค้นไม่ส่งผลให้เวลาที่ใช้ในการคิวรีแตกต่างกันมาก แต่ตัวอย่างที่ใช้ทดลองในการคิวรีมีเวลาใกล้เคียงกัน และใช้เวลาโดยเฉลี่ยไม่ถึง 1 วินาที (เวลาเฉลี่ย 0.061618 วินาที) ทั้งนี้จากแผนภูมิจะเห็นได้ว่าความสูงของกราฟที่ได้มีลักษณะเป็นเส้นตรง และมีบางข้อมูลที่มีค่าเวลาที่ใช้ในการคิวรีมากกว่าข้อมูลอื่นๆ เนื่องจากมีปัจจัยภายนอกที่ส่งผลให้เวลาที่ใช้ในการคิวรีมากขึ้น เช่น ความเร็วของอินเทอร์เน็ต เป็นต้น

การประเมินเวลาของการนำเข้าข้อมูล

ส่วนเวลาการค้นหาค้นด้วย Java Concurrency โดยการแตกเธรดพบว่าเมื่อมีจำนวนเธรดมากขึ้น จะช่วยให้เวลาที่ใช้ในการคิวรีลดลง โดยในเบื้องต้น ทดสอบกับข้อมูลจาก DBpedia BTC2012¹ โดยเพิ่มข้อมูลมีขนาด 4GB จำนวน 21,398,608 Triples

คณะผู้วิจัยทดสอบเทียบเวลาการทำงานของการทำงานเมื่อใช้เครื่องคอมพิวเตอร์เพียง 1 โหนด กับเวลาเมื่อใช้เครื่องคอมพิวเตอร์จำนวน 4 โหนด โดยงาน Map มีจำนวนตั้งแต่ 4 ถึง 61 และขนาดบล็อกข้อมูลมีขนาดตั้งแต่ 64, 128, 256, 512 และ 1024 MB และทดสอบกับงาน Reduce 1 หรือ 2 งาน ผลการทดลองพบว่าเวลาที่ใช้ลดลงเมื่อจำนวนงาน Map ลดลง และเมื่อขนาดบล็อกใหญ่ขึ้น สำหรับคลัสเตอร์ขนาดเล็กนี้ จากงานวิจัยในอดีตพบว่าจำนวนงาน Map จะต้องสอดคล้องกับจำนวนคอร์ทั้งหมดที่มี และจำนวนงาน Reduce มีเพียง 1 หรือ 2 งาน ก็เพียงพอ แต่ควรกำหนดขนาดบล็อก ให้ใหญ่ให้เหมาะสมเพียงพอเพื่อให้เวลาการทำ I/O ไม่มากเกินไปกว่าเวลาคำนวณในบล็อกนั้นจริงๆ

ผลการทดสอบพบว่าประสิทธิภาพด้านเวลาของการทำงานบน 1 โหนด กับการทำงานบน 4 โหนด พบว่าการทำงานบน 4 โหนดทำงานได้เร็วกว่า เมื่อให้จำนวนงาน Map เป็น 8 ถ้ามีจำนวนโหนดมากขึ้น ก็จะได้ผลด้านเวลาดีขึ้น และเมื่อกำหนดขนาดของแฟ้มให้ใหญ่ในระดับ GB การทำงานของ MapReduce เป็นการทำงานที่เน้นข้อมูลมากๆ และมีการ Overlap การคำนวณระหว่างงานเมื่องาน Map ทำงานเสร็จ Reduce ก็ทำงานต่อได้ และงาน Map อื่นๆ ก็ทำงานต่อไป จึงจะทำให้การทำงานนั้นต่อเนื่องงาน Map จึงต้องประมวลผลกับข้อมูลจำนวนมาก นอกจากนี้เวลาในการเริ่มต้นต่างๆ เช่นการเริ่มต้นของ Java Virtual Machine (JVM) ก็ไม่ต่ำกว่า 0.3-0.5 วินาที^{2,3} ดังนั้น ขนาดข้อมูลในบล็อกควรใหญ่พอเพียงด้วยให้งาน Map ทั้งหลายทำงานพร้อมกันด้วย แต่ก็ไม่ทำให้เกิดจำนวนงาน Map มากเกินไปเพราะจะทำให้ Overhead มากเนื่องจากมีงานในระบบเป็นจำนวนงานเกินไปเทียบกับจำนวนแกนจริงของเครื่องคอมพิวเตอร์

เมื่อเทียบกับเวลาการรวมแฟ้มข้อมูลทั่วไปด้วยโปรแกรม Java แบบลำดับจะใช้เวลามากกว่า 10 นาที สำหรับข้อมูลขนาด 4GB หรือบางครั้งอาจจะรันไม่ได้เลย การปรับพารามิเตอร์ต่างๆ เช่นจำนวนงาน Map จำนวนงาน Reduce ขนาดของบล็อก ขนาดของแฟ้มข้อมูลทั้งหมด จำนวน VM และอื่นๆ ของ Hadoop ให้เหมาะกับคลัสเตอร์ที่มีคอยข้างขึ้นกับลักษณะของงานและคลัสเตอร์อย่างมาก

การประเมินเวลาของการค้นหาค้นด้วยการแทนค่าด้วยต้นไม้ไบนารี

เมื่อใช้การแทนค่าแบบไบนารีหรือต้นไม้ 0-1 จะเห็นว่าจะมีขั้นตอนเพิ่มกว่าการใช้รูปแบบ RDF ทั่วไป เพราะจำต้องมีการแปลงจาก RDF เป็น HDT และจาก HDT เป็นไบนารี จากนั้นจึงนำข้อมูลที่ผ่านการแปลงเรียบร้อยแล้วเข้าหน่วยประมวลผลกราฟิกส์ เพื่อทำการค้นหา การแปลงเป็นไบนารีย่อมได้เปรียบในการย่อขนาดข้อมูลซึ่งจะทำให้สามารถนำข้อมูลเข้าไปในหน่วยประมวลผลกราฟิกส์ได้มากกว่ารูปแบบ RDF

¹ <http://km.aifb.kit.edu/projects/btc-2012/>

² <http://stackoverflow.com/questions/10862096/hadoop-map-reduce-how-to-speed-up-job-launch-setup>.

³ <http://blog.griddynamics.com/2010/07/war-story-optimizing-one-hadoop-job.html>.

คณะผู้วิจัยได้ทำการเปรียบเทียบจำนวนทอมในหน่วยหลักล้านของแฟ้มข้อมูลในรูปแบบ RDF รูปแบบ HDT และรูปแบบไบนารีที่นำเสนอ จะเห็นว่าในงานวิจัยนี้นำเสนอรูปแบบข้อมูลที่มีขนาดเล็กกว่าแบบดั้งเดิมและแบบ HDT ในทุกกรณี

จากนั้นทำการทดสอบวัดเวลาในการโอนย้ายข้อมูลจากหน่วยประมวลผลกลางไปหน่วยประมวลผลกราฟิกส์ เมื่อข้อมูลเป็นแบบไบนารีจะทำให้สามารถโอนย้ายข้อมูลได้เร็วกว่า

จากการวัดเวลาที่ใช้ในการค้นหาข้อมูลในหน่วยประมวลผลกลางและหน่วยประมวลผลกราฟิกส์ ตามลำดับ เมื่อข้อมูลอยู่ในรูปแบบ HDT ที่ได้กำหนดการสุ่ม ID ของ Subject, Predicate, Object จำนวน 1,024 ตัว และวัดเวลาเฉลี่ยของการค้นหาจำนวน 25 ครั้ง เมื่อข้อมูลอยู่ในรูปแบบ SPO โดยใช้ข้อมูลต่อไปนี้ในการทดสอบ swdf, wn3.0, Freebase และ DBPedia จะเห็นว่าวิธีการแบบไบนารี สามารถนำเข้าข้อมูลได้มากกว่า และค้นหาข้อมูลได้พร้อมกันหลายตำแหน่ง ทำให้ใช้เวลาน้อยกว่า อย่างไรก็ตามวิธีการใช้วิธีนี้จำเป็นต้องมีเวลาที่ต้องใช้ในการแปลงข้อมูลมาพิจารณาด้วย

3. การประเมินความพึงพอใจของผู้ใช้บริการเว็บเชิงความหมาย

ในการวิจัยได้ทำการประเมินความพึงพอใจของผู้ใช้งานระบบเว็บเชิงความหมายกำหนดให้ผู้ใช้ระบบจำนวน 100 คนทำแบบสอบถามเพื่อประเมินประสิทธิภาพการทำงานของระบบใน 3 ด้านคือ

- ความพึงพอใจด้านคุณภาพของข้อมูล
- ความพึงพอใจด้านการออกแบบและการจัดรูปแบบ
- ความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับ

ผลการประเมินความพึงพอใจด้านคุณภาพของข้อมูลจะเห็นได้ว่า ค่าเฉลี่ยของความพึงพอใจด้านคุณภาพของข้อมูลเท่ากับ 3.96 และมีค่าเฉลี่ยของส่วนเบี่ยงเบนมาตรฐานเท่ากับ 0.57 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพในด้านคุณภาพของข้อมูลอยู่ในระดับดี

ผลการประเมินความพึงพอใจด้านการออกแบบและการจัดรูปแบบ การประเมินความพึงพอใจด้านการออกแบบและการจัดรูปแบบซึ่งได้ผลลัพธ์ว่าความพึงพอใจด้านนี้มีค่าเฉลี่ยเท่ากับ 4.08 และมีค่าส่วนเบี่ยงเบนมาตรฐานโดยเฉลี่ยเท่ากับ 0.54 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพในการออกแบบและการจัดรูปแบบอยู่ในระดับดี

ผลการประเมินความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับ การประเมินความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับซึ่งได้ผลลัพธ์ว่าความพึงพอใจด้านนี้มีค่าเฉลี่ยเท่ากับ 3.90 และมีค่าส่วนเบี่ยงเบนมาตรฐานโดยเฉลี่ยเท่ากับ 0.54 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพในด้านประโยชน์ของข้อมูลที่ได้รับอยู่ในระดับดี

ผลการประเมินความพึงพอใจโดยรวม ซึ่งได้ผลลัพธ์ว่าความพึงพอใจโดยรวมมีค่าเฉลี่ยเท่ากับ 3.96 และมีค่าส่วนเบี่ยงเบนมาตรฐานโดยเฉลี่ยเท่ากับ 0.53 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพโดยรวมอยู่ในระดับดี

บทสรุปและข้อเสนอแนะ

ในโครงการย่อยนี้ ได้ทำการสำรวจความต้องการของซอฟต์แวร์ปัจจุบัน รวมทั้งลักษณะโครงสร้างการทำงานของหน่วยงานเทศบาลเมือง พบว่าการค้นหาข้อมูลส่วนใหญ่เป็นการค้นหาจากเว็บไซต์ทั่วไปเช่น Agoda, TripAdvisor ผ่าน Google Search Engine เป็นต้น นอกจากนี้ข้อมูลอื่นๆ มาจากการเดินสำรวจ ซอฟต์แวร์ในการจัดเก็บข้อมูล และเก็บสถิติได้แก่ Excel และ Word เพื่อทำรายงานสรุปต่างๆ แหล่งข้อมูลต่างๆ ในการ

รวบรวมข้อมูลสถานประกอบการได้แก่ จากเว็บไซต์ Agoda, TripAdvisor, AtSiam, HeyHuaHin และจาก Official Website ของสถานบริการ รวมทั้งการเดินสำรวจในพื้นที่หัวหินและชะอำ ข้อมูลเหล่านี้ถูกนำเข้าออนโทโลยี เชื่อมต่อกับหน้าจอผู้ใช้งานที่ได้ออกแบบไว้ เพื่อสามารถค้นหาได้ สำหรับการประมวลผลการค้นหา เมื่อผู้ใช้ระบุเงื่อนไขการค้นหาจะถูกแปลงเป็นคิวรีส่งไปยังมิดเดิลแวร์เพื่อติดต่อกับออนโทโลยี และค้นหาข้อมูลที่ต้องการขึ้นมาแสดงผล

คณะผู้วิจัยการนำข้อมูลเข้าแบบขนานด้วยเทคโนโลยี Java แบบพร้อมกัน (Concurrent) และ MapReduce โดยพบว่าการนำเข้าข้อมูลโดยการค้นหาเว็บแบบพร้อมกัน ทำให้ทำงานได้เร็วขึ้น ประมาณ 3-4 เท่าสำหรับกรณีใช้ 4 เธรด และการทำงานแบบ MapReduce จะช่วยทำการรวมข้อมูลจากแฟ้มข้อมูลขนาดใหญ่ เช่น 4 GB ได้ดีภายในเวลาไม่เกิน 7 วินาที เทียบกับวิธีปกติโดยโปรแกรมภาษา Java แบบ 1 เธรดที่ใช้เวลาร่วมสิบกว่านาที นอกจากนี้คณะผู้วิจัยได้ทำการวัดประสิทธิภาพการค้นหาจากออนโทโลยี โดยจะพบว่าเวลาเฉลี่ยในการค้นหาไม่เกิน 30 วินาที และคณะผู้วิจัยได้ทำการนำเสนอโมเดลการค้นหาด้วยวิธีการแบบขนานในโมเดลต้นไม้แบบไบนารีกับหน่วยประมวลผลกราฟิกส์เพื่อประมวลผลแบบขนาน โดยนำข้อมูลออนโทโลยีในรูปแบบ Triple เข้ามาในหน่วยประมวลผลกราฟิกส์เพื่อค้นหาแบบพร้อมกัน ซึ่งสามารถนำเข้าข้อมูลได้มากขึ้นและนำมาค้นหาได้เร็วขึ้นนั่นเอง

คณะผู้วิจัยได้ประเมินความพึงพอใจจากผู้ใช้จำนวน 200 คนเพื่อวัดความพึงพอใจของข้อมูลและการออกแบบหน้าเว็บอีกด้วย การออกแบบระบบด้วยออนโทโลยีให้ข้อดี ในการค้นหาข้อมูลต่างๆ ได้ตรงตามความต้องการของผู้ใช้มากขึ้น เนื่องจากเป็นการค้นหาตามความหมายของคำค้นในบริบทที่ผู้ใช้ต้องการจริง อย่างไรก็ตามเครื่องมือที่สนับสนุนการออกแบบออนโทโลียังมีขีดจำกัดในด้านต่างๆ เช่น กลไกการอนุมาน ลอจิกที่รองรับความเสถียรของระบบเชื่อมต่อ และการทำงานที่ช้ากว่าระบบฐานข้อมูลเชิงสัมพันธ์ที่มีอยู่ ดังนั้นการพัฒนาอย่างต่อเนื่องของเทคโนโลยีเหล่านี้จะทำให้ระบบค้นหาเชิงความหมายสมบูรณ์ขึ้น

นอกจากนี้การรองรับการประมวลผลแบบขนานของเครื่องมือซอฟต์แวร์ที่มีอยู่ยังมีขีดจำกัดทำให้เมื่อต้องการใช้กลไกการค้นหาแบบขนานมาช่วย ต้องทำการออกแบบการเชื่อมต่อ รวมทั้งเวลาการประมวลผลขึ้นก่อนการค้นหาหรือ Preprocess ต่างๆ ค่อนข้างมากซึ่งเป็น Overhead ในการทำงานโดยรวมอยู่

ในงานวิจัยในอนาคตการพัฒนาโครงสร้างข้อมูลการแทนค่าออนโทโลยีพื้นฐานนี้ให้เหมาะกับการประมวลผลแบบขนานไม่ว่าจะด้วยลักษณะคลัสเตอร์ หรือหน่วยประมวลผลกราฟิกส์ก็ตาม ไลบรารีต่างๆ ที่รองรับจะมีส่วนสำคัญในการเชื่อมต่อการประมวลผลออนโทโลยีกับเว็บเชิงความหมายให้เป็นอัตโนมัติและด้วยประสิทธิภาพที่ดีต่อไป

บทคัดย่อ

ในงานวิจัยนี้ คณะผู้วิจัยได้พัฒนาระบบค้นหาสำหรับการบริการข้อมูลเชิงสุขภาพในอำเภอหัวหิน ระบบค้นหาได้อาศัยเทคโนโลยีการประมวลแบบขนานและกลุ่มเมฆ เชื่อมต่อกับข้อมูลท่องเที่ยวเชิงสุขภาพที่จัดเก็บในรูปแบบออนโทโลยี สนับสนุนการเทคโนโลยีเชิงความหมาย คณะผู้วิจัยได้นำระบบขึ้นไปวางบนกลุ่มเมฆในลักษณะงานประยุกต์โดยการสนับสนุนของ กสท. (ตะวันตก) ภายในแพลตฟอร์มของ IRIS

การประมวลแบบขนานที่ใช้ในงานวิจัยนี้มี 2 ด้านได้แก่ ด้านเสาะแสวงหา ดึงข้อมูล และเก็บข้อมูล โดยใช้เทคโนโลยี Java Concurrent (Java 1.7) ร่วมกับการพัฒนาโปรแกรมสำหรับการจัดทำกลไกการดึงรวบรวมข้อมูลหน้าเว็บ นอกจากนี้ได้ใช้เฟรมเวิร์ก MapReduce ในการพัฒนาส่วนของการรวมข้อมูล อีกด้านได้แก่การออกแบบการแทนค่าข้อมูลสำหรับการประมวลคิวรีภายในสถาปัตยกรรม GPU ซึ่งเป็นแบบใช้หน่วยความจำร่วมกันแบบหลายแกน ในรายงานได้นำเสนอกลไกทั้งหมดตั้งแต่การรับข้อมูลจากผู้ใช้งานแปลงไปเป็นคิวรีสำหรับการค้นหา การนำเข้าข้อมูลใน GPU สำหรับค้นหา และการนำผลลัพธ์การค้นหาจากการประมวลผลการค้นหาแบบขนานมาแสดง

การวัดผลด้านประสิทธิภาพมีในการค้นหาข้อมูลเชิงความหมายพบว่าใช้เวลาประมาณ 0.06 วินาทีในการค้นหาแต่ละคิวรี และ Speedup ที่ได้จากการใช้เครื่องที่มีจำนวน 4 แกนในการดึงข้อมูลแบบขนานประมาณ 3.6 เท่า และการใช้ MapReduce ในการรวมไฟล์ข้อมูลขนาด 4 GB ประมาณ 6-7 วินาที สำหรับการค้นหาใน GPU ด้วยวิธีการแทนค่าข้อมูลที่น่าเสนอนั้น ทำให้สามารถนำเข้าข้อมูลได้มากกว่า และสามารถนำเข้าข้อมูลได้ถึง 18 ล้านเทอม ใช้เวลาค้นหา 34 มิลลิวินาที สำหรับด้านความพึงพอใจในการค้นหาเว็บเชิงความหมายจากผู้ใช้ 200 คน ได้คะแนนเฉลี่ยด้านความถูกต้องของข้อมูล และความพอใจในหน้าจอการใช้งาน 3.96/4.00

Abstract

In this project, Researchers develop the search engine for health information services in Hua Hin district. The search engine employs the parallel processing and cloud computing. The search engine connects with the health tourism data which is stored in a prototype ontology format. It supports the semantic web technology. Researchers particularly develop the search engine as a cloud application on IRIS sponsored by CAT Telecom Co. Ltd. (Thailand).

The parallel processing used in this research is two folds. First, it is for the purpose of data extraction and collection using Java Concurrent Package (1.7) together with the development of the parallel programs to implement the web crawler engine. Also, the MapReduce framework is used to combine the variety of data extraction. Secondly, the data representation for parallel query processing is designed. The parallel query processing is based on the shared-memory architecture which is on the GPU platform. The research reports the whole process from the user queries until the search results are derived and given back to the user using the representation and the parallel query.

Researchers measured the performance of the semantic query and find that it takes about 0.06 seconds to query the semantic web. Researchers measured the performance by using the parallel processing using the speedup, the speedup obtained by the parallel web crawler is about 3.6 for using 4-core computer and using MapReduce to merge the data about 4 GB takes only 6-7 seconds. The prototype of the searching in GPU can perform the search with larger data than using the search directly with the traditional RDF format. It can search 18 million terms using 34 milliseconds for searching. The user satisfaction of the semantic web development is on average 3.96/4.0 from 200 users.

สารบัญ

บทที่ 1 บทนำ.....	1
1.1 ที่มาความสำคัญ.....	1
1.2 วัตถุประสงค์หลักของโครงการ.....	1
1.3 ขอบเขตการวิจัย.....	1
1.4 วิธีการดำเนินการวิจัย.....	2
1.5 ประโยชน์ที่คาดว่าจะได้รับ	2
บทที่ 2 การทบทวนเอกสารและวรรณกรรมที่เกี่ยวข้อง	4
2.1 งานวิจัยเกี่ยวกับการบริการบนกลุ่มเมฆ	4
2.2 แพลตฟอร์มสำหรับการท่องเที่ยวทางอิเล็กทรอนิกส์.....	9
2.3 การคำนวณแบบขนานด้วย MapReduce.....	15
2.4 งานวิจัยที่เกี่ยวข้องด้านการประมวลผลกับข้อมูลขนาดใหญ่.....	19
2.5 งานวิจัยที่เกี่ยวข้องด้านการแทนค่าข้อมูล.....	22
2.6 การประมวลผลบนกลุ่มเมฆกับการบริการเว็บเชิงความหมาย	23
2.7 เทคโนโลยีของ Hadoop.....	25
2.8 Java Concurrency [32].....	31
บทที่ 3 การดำเนินการวิจัย.....	36
3.1 สถาปัตยกรรมโครงสร้างสำหรับการนำข้อมูลเข้าแบบขนาน.....	37
3.2 การพัฒนาอัลกอริธึมการประมวลผลคิวรีแบบขนาน และการแทนค่าข้อมูล	40
3.2.1 การแปลง RDF เป็นชุดตัวเลขไบนารี.....	42
3.2.2 การค้นหาแบบขนานด้วยวิธีบรูซฟอร์ซของหน่วยประมวลผลกราฟิกส์.....	46

3.3 Use Case Diagram และหน้าจอผู้ใช้งาน	50
3.4 การเชื่อมต่อข้อมูลกับระบบอื่นๆ	54
3.5 โครงสร้างต้นแบบระบบบนกลุ่มเมฆ	55
3.6 การประเมินความพึงพอใจ	56
บทที่ 4 ผลการวิจัย	58
4.1 การสำรวจข้อมูลของเทศบาลด้านลักษณะหน่วยงาน ข้อมูลที่มีและซอฟต์แวร์ต่างๆ.....	58
4.1.1 ข้อมูลเกี่ยวข้องกับหน่วยงานเทศบาลเมืองหัวหิน.....	58
4.1.2 สรุปข้อมูลที่ได้จากการสำรวจความต้องการและข้อมูลที่มีอยู่	59
4.1.3 ข้อมูลจากหน่วยงานเพื่อตรวจสอบความถูกต้อง	60
4.2 ผลการประเมินประสิทธิภาพ	61
4.2.1 เวลาการค้นหาด้วยเว็บเชิงความหมาย.....	61
4.2.2 การประเมินเวลาของการนำเข้าข้อมูล	64
4.2.3 การประเมินเวลาของการค้นหาด้วยต้นไม้ไบนารี.....	67
4.3 การประเมินความพึงพอใจของผู้ใช้บริการเว็บเชิงความหมาย.....	69
บทที่ 5 สรุปผลการวิจัย อภิปรายผล และข้อเสนอแนะ.....	72
5.1 ข้อเสนอแนะ.....	74
เอกสารอ้างอิง.....	75

ภาคผนวก ก. การติดตั้ง Hadoop และการทดสอบประสิทธิภาพ MapReduce

ภาคผนวก ข. วิธีที่ใช้ทดสอบประสิทธิภาพด้านเวลา

สารบัญภาพ

ภาพที่ 2.1 การบริการของกลุ่มเมฆ 5 ระดับและตัวอย่างการบริการ Storage	7
ภาพที่ 2.2 โครงงานของการท่องเที่ยวอิเล็กทรอนิกส์ [3].....	10
ภาพที่ 2.3 กระบวนการให้สิทธิ์ใน Role Ontology [3].....	12
ภาพที่ 2.4 การไหลของข้อมูลสำหรับการทำดัชนีด้วยวิธีการแบบขนาน [10].....	16
ภาพที่ 2.5 การทำงานของ LLGrid MapReduce	22
ภาพที่ 2.6 แผนผังการทำงานของ MapReduce [21]	27
ภาพที่ 2.7 โมเดลการเขียนโปรแกรม [27,28]	28
ภาพที่ 2.8 สถาปัตยกรรมระดับ Logical ของ Hadoop [29].....	29
ภาพที่ 2.9 การวาง Data node และ Task Tracker บนเครื่องๆ เดียว [30]	29
ภาพที่ 2.10 ลำดับการทำงานของงาน MapReduce ใน Hadoop [30].....	31
ภาพที่ 2.11 แสดงภาพรวมของการทำงานของ Fork/Join.....	34
ภาพที่ 2.12 กระบวนการ Work Stealing	34
ภาพที่ 2.13 การทำงานของหน่วยประมวลผลกลางระหว่างเทรดเดียว (Single Threaded Model) กับ เฟรมเวิร์ก Fork/Join	35
ภาพที่ 3.1 สถาปัตยกรรมโครงสร้างการนำข้อมูลเข้าแบบขนาน	37
ภาพที่ 3.2 สถาปัตยกรรมซอฟต์แวร์.....	38
ภาพที่ 3.3 ขั้นตอน MapReduce	39
ภาพที่ 3.4 ภาพรวมการแปลงข้อมูลออนโทโลยีเป็นข้อมูลเข้าประมวลผลคิวรีแบบขนาน	42
ภาพที่ 3.5 ข้อมูลที่ได้จากการแปลงแฟมออนโทโลยี RDF เป็น HDT	43
ภาพที่ 3.6 ชุดข้อมูลที่ได้หลังจากแปลงตัวเลข SPO ของ HDT	44

ภาพที่ 3.7 รายละเอียดขั้นตอนการแปลงและการค้นหา.....	45
ภาพที่ 3.8 เพิ่มข้อมูลชุดตัวเลขของ Subject (Predicate(Object))	45
ภาพที่ 3.9 เพิ่มข้อมูล 1- 0 ที่ได้รับการแปลงจากเพิ่มชุดตัวเลข.....	46
ภาพที่ 3.10 การคืนค่าของ Subject ID ตามจำแนงที่เจอ “11” ที่สามารถตรวจสอบได้จากตำแหน่งแรก	47
ภาพที่ 3.11 การคืนค่าของ Predicate ID “10”ที่สามารถตรวจสอบได้จากตำแหน่งแรกที่เป็นลำดับที่สอง	48
ภาพที่ 3.12 การคืนค่าของ Object ID ตาม “01” ที่สามารถตรวจสอบได้จากตำแหน่งแรกที่เป็นลำดับที่สาม.....	48
ภาพที่ 3.13 การคืนค่าของ Object ID ตาม “00” ที่สามารถตรวจสอบได้จากตำแหน่งสุดท้าย และสามารถคำนวณค่าสุดท้ายของ Object หลังสุดได้.....	49
ภาพที่ 3.14 หน้าที่ของผู้ดูแลระบบ	50
ภาพที่ 3.15 หน้าแรกของเว็บไซต์.....	51
ภาพที่ 3.16 หน้า Wellness Tourism	51
ภาพที่ 3.17 หน้าเว็บเมื่อกดเลือกที่ Map.....	52
ภาพที่ 3.18 หน้าเว็บเมื่อกดเลือก Advanced Search	53
ภาพที่ 3.19 การเชื่อมต่อข้อมูลกับระบบข้อมูลอื่นๆ	54
ภาพที่ 3.20 ระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพต้นแบบ	55
ภาพที่ 4.1 แผนผังหน่วยงานที่เกี่ยวข้องกับงานวิจัย.....	58
ภาพที่ 4.2 แผนภูมิแสดงความสัมพันธ์ระหว่างจำนวนผลลัพธ์และเวลาที่ใช้ในการควรี.....	63
ภาพที่ 4.3 แผนภูมิแสดงความสัมพันธ์ระหว่างคำค้นและเวลาที่ใช้ในการควรี.....	64
ภาพที่ 4.4 เวลาในการทำงานหลายงานของ MapReduce เมื่อทำการรวมเพิ่มข้อมูล.....	65
ภาพที่ 4.5 เวลาในการสร้าง Index โดยใช้ 1 โหนดและ 4 โหนด.....	66

ภาพที่ 4.6 เปรียบเทียบจำนวนเทอม (Term) สำหรับตัวอย่างข้อมูลต่างๆ กันในแกน สำหรับวิธีแทนค่าข้อมูล แบบดั้งเดิม (RDF) แบบ HDT และแบบไบนารีในงานวิจัยนี้.....	67
ภาพที่ 4.7 เปรียบเทียบเวลาในการขนย้ายข้อมูลจาก CPU ไปยัง GPU.....	68
ภาพที่ 4.8 เปรียบเทียบขนาดของข้อมูลและตำแหน่ง SPO ที่พบในแต่ละข้อมูล	69

สารบัญตาราง

ตารางที่ 1.1	ผู้มีส่วนร่วมและบทบาท	2
ตารางที่ 2.1	ค่าใช้จ่ายกับความคุ้มค่าในการลงทุนต่อประมาณการใช้งาน	6
ตารางที่ 2.2	กรอบการทำงานของการทำงานที่เกี่ยวข้องเชิงอิเล็กทรอนิกส์โดยรวม.....	11
ตารางที่ 2.3	วิเคราะห์ค่า bound ของการทำงานที่เกี่ยวข้องเชิงสุขภาพด้านการแพทย์ที่ประเทศสิงคโปร์.....	14
ตารางที่ 2.4	การประมวลผลแบบขนานบนกริดและบนกลุ่มเมฆ	20
ตารางที่ 2.5	โครงสร้างข้อมูลและอัลกอริธึมแบบขนานบนฮาร์ดแวร์ต่างๆ.....	21
ตารางที่ 2.6	ทฤษฎีที่เกี่ยวข้องตามระดับขั้นของการประมวลผลบนกลุ่มเมฆ.....	23
ตารางที่ 2.7	ความแตกต่างและความได้เปรียบของระบบปัจจุบันกับระบบเว็บเชิงความหมายที่ใช้อยู่	24
ตารางที่ 3.1	เกณฑ์การให้คะแนนของแบบประเมิน.....	57
ตารางที่ 3.2	เกณฑ์การแปลความหมายข้อมูลและพิจารณาจากค่าตัวกลางเลขคณิต (ค่าเฉลี่ย).....	57
ตารางที่ 4.1	ข้อมูลออนไลน์.....	62
ตารางที่ 4.2	เวลาการค้นหาด้วย Java Concurrency	65
ตารางที่ 4.3	เวลาในการค้นหาข้อมูลด้วย HDT ในหน่วยประมวลผลกลาง	69
ตารางที่ 4.4	เวลาในการค้นหาข้อมูลในรูปแบบ HDT ในหน่วยประมวลผลกราฟิกส์.....	69
ตารางที่ 4.5	เวลาการค้นหา SPO ในหน่วยประมวลผลกราฟิกส์ (GTX560Ti)	69
ตารางที่ 4.6	ผลการประเมินความพึงพอใจด้านคุณภาพของข้อมูล	70
ตารางที่ 4.7	ผลการประเมินความพึงพอใจด้านการออกแบบและการจัดรูปแบบ.....	70
ตารางที่ 4.8	ผลการประเมินความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับ.....	71
ตารางที่ 4.9	ผลการประเมินความพึงพอใจโดยรวม	71

ตารางที่ 5.1 ตารางเปรียบเทียบวัตถุประสงค์กิจกรรมที่วางแผนไว้และกิจกรรมที่ดำเนินการมา และผลที่ได้รับ ของ โครงการย่อย 1	72
--	----

บทที่ 1 บทนำ

1.1 ที่มาความสำคัญ

ในแผนพัฒนาเศรษฐกิจและสังคมแห่งชาติ ฉบับที่ 11 (พ.ศ. 2554 - พ.ศ.2559) การพัฒนา ศักยภาพระดับคุณภาพบริการในอุตสาหกรรมบริการและการท่องเที่ยวโดยเฉพาะจากการท่องเที่ยวเชิงสุขภาพ สามารถสร้างมูลค่าเพิ่มให้กับสาขาบริการที่เป็นมิตรกับสิ่งแวดล้อม การพำนักระยะยาวและดูแลผู้สูงอายุใน ท้องถิ่นด้วยความคิดสร้างสรรค์และนวัตกรรมแบบใหม่ เทคโนโลยีปัจจุบันในสาขาคอมพิวเตอร์ที่จะมีบทบาทใน การปรับปรุงการให้บริการนี้ ได้แก่ การประมวลผลแบบขนานบนกลุ่มเมฆ และการบริการเว็บเชิงความหมายด้วย ออนโทโลยีนั้น ซึ่งจะมีส่วนช่วยในประเด็นต่างๆ เช่น

- 1) ความรวดเร็วในการสืบค้นข้อมูลด้วยการประมวลผลแบบขนาน
- 2) คำตอบจากการสืบค้นที่ตรงตามความต้องการของผู้ค้นหาข้อมูล ด้วยการค้นหาเชิงความหมาย ในขอบเขตโดเมนความรู้เฉพาะทาง
- 3) ความสามารถในการขยายฐานข้อมูลผู้ให้บริการการท่องเที่ยวในรูปแบบต่างๆ ที่ไม่ยึดติดพื้นที่ ทางกายภาพด้วยเทคโนโลยีบนกลุ่มเมฆ

ในโครงการวิจัยนี้ ผู้วิจัยขอเสนอ ต้นแบบระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพโดยอาศัย การผสมผสานกับเทคโนโลยีนี้สำหรับกรณีศึกษา อำเภอหัวหิน ในโครงการย่อยนี้จะเน้นการนำเสนอข้อมูลเพื่อให้ สอดคล้องกับการประมวลผลแบบขนานและประมวลผลบนกลุ่มเมฆ พัฒนาวิธีการประมวลผลแบบขนานที่ เหมาะสม มีประสิทธิภาพเพื่อการค้นหาข้อมูลที่รวดเร็ว

1.2 วัตถุประสงค์หลักของโครงการ

เพื่อพัฒนาระบบการค้นหาข้อมูล สำหรับให้บริการข้อมูลการท่องเที่ยวเชิงสุขภาพสำหรับอำเภอหัว หิน โดยเน้นที่การประมวลผลข้อมูลแบบขนาน และสำหรับบนกลุ่มเมฆอย่างมีประสิทธิภาพ

1.3 ขอบเขตการวิจัย

ออกแบบและพัฒนาระบบบริการข้อมูลและจัดการข้อมูลท่องเที่ยวเชิงสุขภาพสำหรับอำเภอหัวหิน โดยอาศัยเทคโนโลยีออนโทโลยี โดยระบบประกอบด้วยระบบที่รองรับการทำงานของผู้ใช้ในระดับต่างๆ ที่สามารถ รวบรวมได้ พัฒนาวิธีการประมวลผลแบบขนานมาช่วยในการค้นหาข้อมูลและประมวลผลข้อมูล ประเมิน ประสิทธิภาพที่ประมวลผลด้านความเร็ว พัฒนาการสืบค้นเว็บเชิงความหมาย เพื่อนำรูปแบบออนโทโลยีแสดงผล การค้นหาผ่านเว็บ นำระบบบริการข้อมูลดังกล่าวไปติดตั้งทดสอบกับให้กับเทศบาลอำเภอหัวหิน

ขอบเขตพื้นที่ที่ศึกษาได้แก่ พื้นที่ทางภูมิศาสตร์ เทศบาลเมืองหัวหิน ส่วนกรณีอำเภอชะอำหรือ ไกลเคียง สามารถดึงข้อมูลที่มีอยู่จากสื่ออิเล็กทรอนิกส์ได้ และตาราง 1.1 แสดงผู้ใช้งานในระบบและบทบาทของ ผู้ใช้

ตารางที่ 1.1 ผู้มีส่วนร่วมและบทบาท

ผู้ใช้	ระบบที่ใช้
ผู้บริหาร หัวหน้าส่วน	แสดงผลของระดับผู้บริหาร รายงานสถิติผู้เข้าใช้ รายงานสรุปความคิดเห็น เป็นต้น
ผู้ดูแลระบบ เจ้าหน้าที่ท้องถิ่น	แสดงผลของระดับผู้ดูแลระบบ การปรับปรุงข้อมูลระบบในส่วนของเหตุการณ์ (Event) ข่าวสารของเทศบาล
นักท่องเที่ยว	หน้าจอรับและแสดงผลลัพธ์จากการค้นหา เช่นข้อมูลสปา ข้อมูลสิ่งอำนวยความสะดวก บริการ ข้อมูลสถานพยาบาล ข้อมูลการอีเวนต์ ข้อมูลส่งเสริมการขาย ข้อมูลสภาพอากาศ ข้อมูลอัตราแลกเปลี่ยนประจำวัน ข้อมูลสถานที่พักแรม และอื่นๆ
ผู้ประกอบการสถานบริการ	แสดงผลข้อมูลเฉพาะสถานประกอบการ

1.4 วิธีการดำเนินการวิจัย

- 1.4.1 สำรวจความต้องการซอฟต์แวร์ในปัจจุบัน
- 1.4.2 การศึกษางานวิจัยที่เกี่ยวข้องด้านการประมวลผลแบบขนาน และกลุ่มเมฆ
- 1.4.3 พัฒนาตัวแบบของระบบการประมวลผลแบบกลุ่มเมฆ
- 1.4.4 พัฒนาอัลกอริธึมการประมวลผลคิวรีแบบขนาน และการแทนค่าข้อมูล
- 1.4.5 พัฒนาโปรแกรมการประมวลผลคิวรีแบบขนาน
- 1.4.6 การออกแบบหน้าจอ การติดต่อผู้ใช้ ระบบออนไลน์
- 1.4.7 ออกแบบฐานข้อมูลระบบ ได้แก่ ข้อมูลผู้ใช้ระบบ ข้อมูลสถานที่ท่องเที่ยว
- 1.4.8 พัฒนาระบบจัดการการค้นหา พัฒนาโปรแกรม ทดลองใช้
- 1.4.9 เปรียบเทียบผลการค้นหาคิวรี ด้านประสิทธิภาพ
- 1.4.10 ประเมินผลการใช้งานกับผู้ใช้บทบาทต่างๆ กัน
- 1.4.11 เขียนสรุปรายงานและเอกสารวิชาการ
- 1.4.12 จัดเตรียมการอบรมผู้ใช้ จัดทำคู่มือผู้ใช้งาน

1.5 ประโยชน์ที่คาดว่าจะได้รับ

- 1.5.1 ระบบจัดการและให้บริการข้อมูล เพื่อสนับสนุนการท่องเที่ยวเชิงสุขภาพ กรณีศึกษาการท่องเที่ยวเชิงสุขภาพอำเภอหัวหิน

1.5.2 อัลกอริธึม โปรแกรมการประมวลผลวิธีแบบขนาน และตัวแบบการประมวลผลบนกลุ่มเมฆ เพื่อสนับสนุนการท่องเที่ยวเชิงสุขภาพ กรณีศึกษาการท่องเที่ยวเชิงสุขภาพอำเภอหัวหิน

1.5.3 ผลการประเมินด้านประสิทธิภาพ ผลการประเมินด้านความถูกต้องในการค้นหา

1.5.4 คู่มือการใช้งานระบบ สรุปรายงาน เอกสารวิชาการ ผลงานตีพิมพ์

1.5.5 นักวิจัย บัณฑิตและมหาบัณฑิตในสาขาคอมพิวเตอร์ที่มีทักษะในการแก้ปัญหา

บทที่ 2 การทบทวนเอกสารและวรรณกรรมที่เกี่ยวข้อง

งานวิจัยนี้เป็นการศึกษาระบบเว็บเชิงความหมายเพื่อการจัดการและสนับสนุนการประมวลแบบขนานสำหรับการท่องเที่ยวเชิงสุขภาพ อำเภอหัวหิน โดยใช้ทฤษฎีการบริการเว็บเชิงความหมายสร้างระบบต้นแบบการเรียกใช้บริการด้านการท่องเที่ยว และการประมวลผลแบบขนานเพื่อเพิ่มประสิทธิภาพของการรับส่งข้อมูลออนไลน์ การประมวลผลฐานความรู้บนกลุ่มเมฆ การคิวรีข้อมูล Triple และการค้นหาฐานความรู้ออนไลน์ในสถานที่ท่องเที่ยวจริง เพื่อให้สามารถประเมินประสิทธิภาพของการใช้ทรัพยากรทั้งหมด ในการค้นหาฐานความรู้ด้านการท่องเที่ยวเชิงสุขภาพที่มีการเชื่อมต่อของข้อมูลทั่วโลกที่มีการปรับปรุงสม่ำเสมอจากลิงค์เดต้า (Linked Data Cloud) นอกจากนี้ยังมีระบบการบริหารจัดการท่องเที่ยวสำหรับกลุ่มผู้ใช้ที่เป็นนักท่องเที่ยว ผู้ดูแลระบบ ผู้บริหารและนักเทคนิค

จึงขอเสนอวรรณกรรมที่เกี่ยวข้องกับการสร้างระบบโดยรวม นอกเหนือจากที่กล่าวถึงในแผนงานวิจัย ยังมีงานวิจัยที่เกี่ยวข้องกับการบริการกลุ่มเมฆ แพลตฟอร์มเกี่ยวกับงานบริการเกี่ยวกับท่องเที่ยว และการประมวลแบบขนาน รวมทั้งพื้นฐานด้านการประมวลผลบนกลุ่มเมฆ (Cloud Computing) การบริการเว็บเชิงความหมาย (Semantic Web Service) และการประมวลผลแบบขนาน (Parallel Processing)

2.1 งานวิจัยเกี่ยวกับการบริการบนกลุ่มเมฆ

บริการจากเว็บเซอร์วิสจากอเมซอน (Amazon) เป็นมาตรฐานที่ได้รับการยอมรับในกลุ่ม Infrastructure as a service (IaaS) ที่ผู้ใช้มีการจ่ายค่าบริการตามการใช้งานผ่าน VM Elastic Compute Cloud (EC2) ¹ ซึ่งถือว่าอเมซอน เป็นผู้นำตลาดในด้าน Xen-based virtualized infrastructure ² ในด้านการตลาด อเมซอนมักจะจัดการส่งเสริมการขายด้วยการลดราคา นอกจากนี้อเมซอนยังสร้างพันธมิตรในกลุ่มเมฆสาธารณะในเรื่องการรักษาสิ่งแวดล้อมที่ได้รับการรับรองเรื่องความปลอดภัยและนำเสนอที่เก็บข้อมูลส่วนกลางที่ครอบคลุมทั่วโลก

ในการประมวลผลกลุ่มเมฆที่มีลักษณะเป็นศูนย์ข้อมูล (Data Center) เช่นนี้ ผู้ให้บริการต้องคุณลักษณะสำคัญ [1] ดังนี้

- 1) มีเจ้าของเป็นหน่วยงานต่อไปนี เช่น องค์กรหรือบริษัท ผู้ให้บริการ (Service Provider) หรือรัฐบาล
- 2) สามารถสร้างกลุ่มของทรัพยากรที่ใช้ในการประมวลผลทั้งหมด และบริการต่างๆ ที่ต้องใช้ร่วมกัน และเปิดให้ผู้ใช้สมัครใช้งานได้

¹ <http://aws.amazon.com/ec2/>

² <http://www.networkworld.com/supp/2012/enterprise2/040912-ecs-iaas-companies-257611.html>

3) มีการเก็บค่าบริการทั้งทรัพยากรและบริการ ตามการใช้งานจริงและ/หรือตามขนาดของแพคเกจบริการพื้นฐาน

4) มีการให้บริการกระจายตามภูมิภาคของโลก ซึ่งในกรณีนี้เป็นการสร้างความโปร่งใสต่อกลุ่มผู้ใช้งาน

5) มีการเพิ่มสิ่งที่เตรียมการและการตั้งค่า และถอนออกได้โดยอัตโนมัติ ทั้งทรัพยากรและบริการที่อยู่บนการบริการพื้นฐานแบบช่วยตัวเองของระบบ โดยปกติจะเปิดให้ผู้ใช้ใช้ได้หลังการสมัครอยู่แล้ว ซึ่งเป็นการตั้งค่าของโปรแกรมที่ไม่ต้องมีมนุษย์เป็นผู้ควบคุมและส่งออกมาเป็นลำดับในเวลาไม่กี่วินาที

6) ทั้งทรัพยากรและการบริการถูกส่งออกมาในลักษณะเสมือนจริง ทั้งเซิร์ฟเวอร์ ดิสก์ และอุปกรณ์เครือข่ายซึ่งล้วนแต่เป็นการใช้ทรัพยากรเสมือนจริง โดยดึงจากโครงสร้างพื้นฐานทางกายภาพที่ผู้ใช้ไม่จำเป็นต้องมองเห็นได้

7) โครงสร้างทางกายภาพควรมีการเปลี่ยนแปลงน้อยมาก แต่ทรัพยากรและการบริการเสมือนนั้นสามารถเปลี่ยนได้ทันที

8) ทรัพยากรและบริการนั้นเป็นลักษณะกายภาพเสมือน ทั้งเซิร์ฟเวอร์ ดิสก์ และอุปกรณ์เครือข่ายหรือในส่วนนามธรรม ทั้งฟังก์ชันการจัดเก็บวัตถุ (blob storage) ฟังก์ชันคิวข้อความ ฟังก์ชันอีเมล ฟังก์ชันมัลติคาสต์ ทั้งหมดนี้สามารถเข้าถึงได้โดยการรันโค้ดหรือสคริปต์เพื่อเซตของ API สำหรับบริการเหล่านี้ หรืออาจใช้ทั้งโค้ดและสคริปต์ผสมกัน

ตัวอย่างสำหรับการบริการ Amazon Elastic Compute Cloud (EC2) มีรูปแบบต่างๆ ดังนี้

1) 750 ชั่วโมงสำหรับ Amazon EC2 Linux (ยกเว้น SuSe Enterprise Server และ RHEL) Micro Instance usage หน่วยความจำ ขนาด 613 MB และรองรับ 32 บิต และ 64 บิต โดยมีจำนวนชั่วโมงรันอย่างต่อเนื่องในแต่ละเดือน *

2) 750 ชั่วโมงสำหรับ Amazon EC2 Microsoft Windows Server† Micro Instance usage หน่วยความจำ ขนาด 613 MB และรองรับ 32 บิต และ 64 บิต โดยมีจำนวนชั่วโมงรันอย่างต่อเนื่องในแต่ละเดือน *

3) 750 ชั่วโมงสำหรับ Amazon Elastic Load Balancer และการประมวลผลข้อมูลขนาด 15 GB *

4) 30 GB ชั่วโมงสำหรับ Amazon Elastic Block Storage และ จำนวน I/O ขนาด 2 ล้านครั้ง และมีพื้นที่ 1 GB snapshot storage*

ข้อยกเว้นมีดังนี้

* เฉพาะลูกค้าใหม่ที่เปิดใช้บริการ เมื่อครบ 12 เดือนแล้วจะจ่ายตามการใช้งานจริง

สำหรับการจ่ายตามจริงมีนโยบายดังต่อไปนี้ กรณีที่เลือกใช้แหล่งทรัพยากรที่ใกล้ที่สุดคือสิงคโปร์และเลือกชนิดการใช้งาน (on demand) ซึ่งเหมาะกับโครงการระยะเวลาไม่นาน (เช่นไม่เกิน 1 ปี) มีข้อดีคือไม่ต้องเสียค่าธรรมเนียมเบื้องต้น สามารถวางแผนการ การจัดซื้อและบำรุงรักษาฮาร์ดแวร์ที่มีความยืดหยุ่นสูงได้ รวมถึงการย้ายระบบด้วยต้นทุนที่น้อยกว่า ราคาในตาราง 2.1 นั้นสามารถใช้ได้ทั้งกิจการสาธารณะและส่วนบุคคลบน

ระบบปฏิบัติการเฉพาะ ซึ่งในกรณีวินโดวส์ คือ Windows Server® 2003 R2, 2008 และ 2008 R2 นอกจากนี้ยังใช้วินโดวส์กับ SQL Server, Amazon EC2 บน SUSE Linux Enterprise Server, Amazon Red Hat Enterprise Linux และ Amazon EC2 บน IBM ซึ่งการคิดราคาจะแตกต่างกันไป สามารถดูเพิ่มเติมได้ที่ <http://aws.amazon.com/ec2/#pricing>

ทั้งนี้ผลิตภัณฑ์ของอเมซอน หรือใน Google App Engine³ ล้วนมีคุณสมบัติเหล่านี้ทั้งสิ้น

ตัวอย่างของการบริการของอเมซอนที่มีการส่งเสริมการขายสำหรับลูกค้าใหม่เป็นเวลา 1 ปีดังนี้ AWS Free Usage Tier (ต่อเดือน)⁴

ตารางที่ 2.1 ค่าใช้จ่ายกับความคุ้มค่าในการลงทุนต่อปริมาณการใช้งาน⁵

	การใช้งานระบบ Linux/UNIX	การใช้งานในระบบ Windows
แพคเกจมาตรฐาน (Standard On-Demand Instances)		
Small (Default)	\$0.085 ต่อชั่วโมง	\$0.115 ต่อชั่วโมง
Medium	\$0.170 ต่อชั่วโมง	\$0.230 ต่อชั่วโมง
Large	\$0.340 ต่อชั่วโมง	\$0.460 ต่อชั่วโมง
Extra Large	\$0.680 ต่อชั่วโมง	\$0.920 ต่อชั่วโมง
แพคเกจตามความต้องการขนาดเล็ก (Micro On-Demand Instances)		
Micro	\$0.020 ต่อชั่วโมง	\$0.020 ต่อชั่วโมง
แพคเกจตามความต้องการขนาดใหญ่ (High-Memory On-Demand Instances)		
Extra Large	\$0.506 ต่อชั่วโมง	\$0.570 ต่อชั่วโมง
Double Extra Large	\$1.012 ต่อชั่วโมง	\$1.140 ต่อชั่วโมง
Quadruple Extra Large	\$2.024 ต่อชั่วโมง	\$2.280 ต่อชั่วโมง
แพคเกจตามความต้องการใช้ซีพียู (High-CPU On-Demand Instances)		

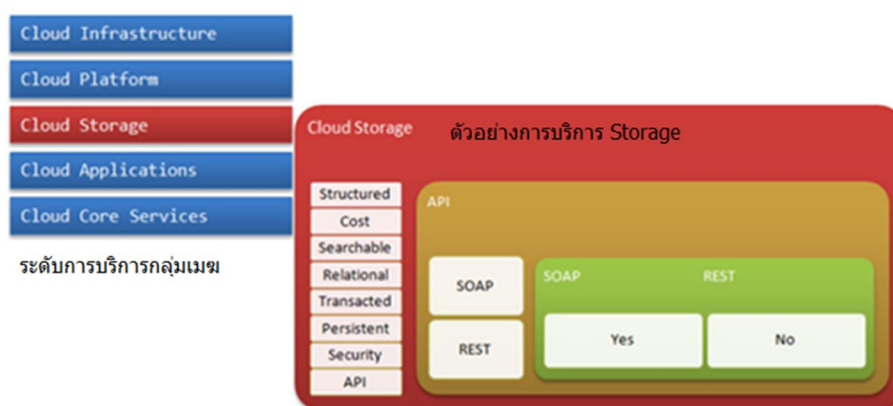
³ <http://code.google.com/appengine/>, [Online] Access: 17 Aug 2012.

⁴ <http://aws.amazon.com/free/>

⁵ <http://aws.amazon.com/ec2/#pricing>

Medium	\$0.186 ต่อชั่วโมง	\$0.285 ต่อชั่วโมง
Extra Large	\$0.744 ต่อชั่วโมง	\$1.140 ต่อชั่วโมง
แพ็คเกจตามความต้องการใช้คลัสเตอร์ (Cluster Compute Instances)		
Quadruple Extra Large	N/A*	N/A*
แพ็คเกจตามความต้องการใช้จีพียู (Cluster GPU Instances)		
Quadruple Extra Large	N/A*	N/A*
* Cluster and High I/O Instances are not available in all regions		

สถาปัตยกรรมของแพลตฟอร์มนั้นเป็นส่วนหนึ่งที่สำคัญ ในการย้ายแอปพลิเคชันไปบนระบบประมวลผลกลุ่มเมฆ⁶ มี 2 วิธีในการพิจารณาเลือกแพลตฟอร์ม วิธีแรกคือ การเลือกคุณสมบัติของแอปพลิเคชันให้เหมาะกับกลุ่มเมฆ เพื่อตรวจสอบดูว่าบริการของกลุ่มเมฆเหมาะสมกับแอปพลิเคชันนั้นๆ หรือไม่ จากนั้นจึงกำหนดประเภทของบริการที่ใช้ วิธีที่ 2 คือการประเมินบริการจากผู้ให้บริการ โดยเป็นการประเมินความเหมาะสมของบริการในการเป็นโฮสต์ของแอปพลิเคชัน เพื่อจะได้เลือกบริการที่เข้ากันได้ และอาจมีการพัฒนาเพิ่มเติมต่อไป



ภาพที่ 2.1 การบริการของกลุ่มเมฆ 5 ระดับและตัวอย่างการบริการ Storage

(ที่มา: <http://i.msdn.microsoft.com/dynimg/IC263304.jpg>, [Online], Access: Aug 17, 2012.)

ภาพที่ 2.1 เป็นตัวอย่างการให้บริการเก็บข้อมูลจากบริการกลุ่มเมฆของบริษัทไมโครซอฟต์ ซึ่งผู้ใช้ต้องพิจารณาความเหมาะสมของบริการที่เลือกใช้ ระดับการบริการมีหลายระดับตั้งแต่ระดับ Infrastructure ไปจนถึง Core Service ในรูปด้านขวาเป็นการเลือกบริการระดับ Storage ซึ่งจะมีการสนับสนุนการบริการเว็บใน

⁶ <http://msdn.microsoft.com/en-us/library/dd430340>, [Online] Access: 17 Aug 2012.

รูปแบบ SOAP⁷ แต่ไม่รองรับ REST⁸ และมีการรองรับด้านความปลอดภัย การทำงานแบบ Transaction ฐานข้อมูลสัมพันธ์ เป็นต้น

Martens [2] ทำการวัดบริการกลุ่มเมฆแบ่งโดยเป็น โบรกเกอร์ สถานที่ ประเภทบริการและ ทรัพยากร Service-level Agreement (SLA) และระบบรักษาความปลอดภัย สำหรับการวัดด้วย KPI ประกอบด้วย ชนิดและการทำงาน การวัดความเสี่ยงนั้นพิจารณาที่ความคิดเห็นต่อความเสี่ยง ผลกระทบใน ภาพรวม และประเภทความเสี่ยง และมุมมองด้านความร่วมมือที่ประกอบด้วย ระดับความร่วมมือ การ ตรวจสอบความร่วมมือ ประเภทของการประมวลผลข้อมูลและประเภทข้อบังคับด้านความร่วมมือ

⁷ <http://en.wikipedia.org/wiki/SOAP>

⁸ http://en.wikipedia.org/wiki/Representational_state_transfer

2.2 แพลตฟอร์มสำหรับการท่องเที่ยวทางอิเล็กทรอนิกส์

งานบริการด้านการท่องเที่ยวอิเล็กทรอนิกส์ (e-Tourism) สามารถนำไปทำงานบนกลุ่มเมฆได้ ในงานวิจัย [3] ได้เสนอแบบงานบริการที่เกี่ยวข้องกับการท่องเที่ยวเป็นงานย่อย ประกอบด้วย แพ็กเกจแบบพลวัต (dynamic packaging) การวางแผนการท่องเที่ยว การเปรียบเทียบราคา การออกแบบเส้นทางการเดินทาง และการโฆษณาทางการตลาดและการส่งเสริมการขาย ซึ่งสามารถได้แบ่งตามสถานการณ์การทำงาน ตามภาพที่ 2.2 ที่ได้แสดงโครงงานการท่องเที่ยวอิเล็กทรอนิกส์ ที่แบ่งขั้นตอนการทำงานที่สำคัญดังนี้

ขั้นตอนที่ 1: นักท่องเที่ยวส่งคำร้องไปยังตัวแทนการท่องเที่ยวเพื่อให้ตัวแทนการท่องเที่ยวทำการวางแผนการท่องเที่ยวโดยใช้องค์กรเสมือน (Virtual Organization) ทำให้สามารถใช้งานข้ามองค์กรได้ เป็นการเปิดสิทธิ์การบริการแก่การท่องเที่ยวแบบหนึ่ง

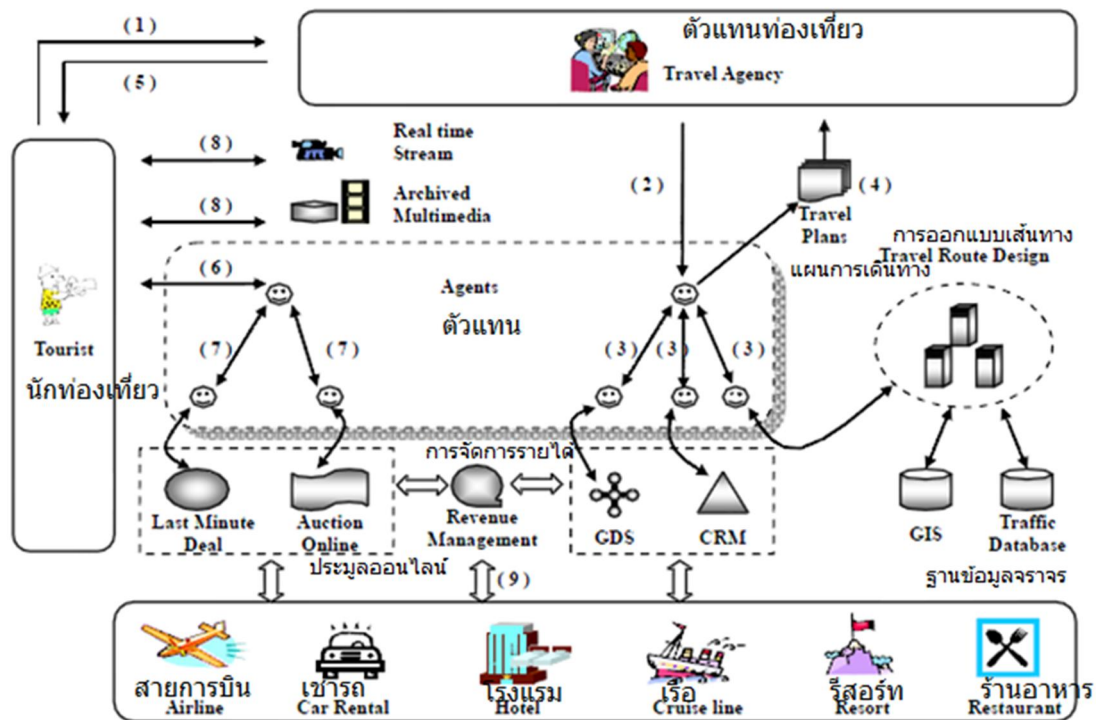
ขั้นตอนที่ 2: ตัวแทนการท่องเที่ยวส่งแผนการท่องเที่ยวให้แก่เอเจนต์

ขั้นตอนที่ 3: ตัวแทนติดต่อตัวแทนอื่นๆ และกำหนดงานแก่สาขาย่อย งานของสาขาแรก คือ ออกแบบเส้นทางการท่องเที่ยว ซึ่งใช้การคำนวณที่ซับซ้อนตามฐานข้อมูลสารสนเทศทางภูมิศาสตร์ (GIS) งานของอีกสองสาขาคือการคิวรี่ไปที่กรมอุตุนิยมวิทยาเพื่อขอผลการพยากรณ์อากาศ และอัตราค่าที่พักจากหน่วยงาน Global Distribution Framework (GDS) และ Central Reservation Management (CRM)

ขั้นตอนที่ 4: ตัวแทนแรกประมวลผลจากแต่ละพันธมิตรทางธุรกิจและส่งแผนการท่องเที่ยวมาให้ตัวแทนการท่องเที่ยว

ขั้นตอนที่ 5: ตัวแทนการท่องเที่ยวส่งแผนการท่องเที่ยวให้นักท่องเที่ยว

จากขั้นตอนการทำงานของการท่องเที่ยวอิเล็กทรอนิกส์ข้างต้นนี้ สามารถนำไปประยุกต์ใช้กับโครงสร้างในอนาคตได้ ซึ่งประกอบด้วย เอเจนต์ส่วนบุคคลและเอเจนต์อัจฉริยะ การบริการข้อมูลเชิงความหมาย การให้ลำดับการเลือกโดยสนับสนุนการตัดสินใจของมนุษย์ การเชื่อมต่อหลายระบบให้ราบรื่น การประมวลผลระหว่างระบบแบบกระจาย การเก็บข้อมูลและความร่วมมือระหว่างพันธมิตร ซึ่งสามารถนำมาเรียบเรียงตามระดับขั้นการส่งข้อมูลได้ดังตารางที่ 2.2



ภาพที่ 2.2 โครงงานของการท่องเที่ยวอิเล็กทรอนิกส์ [3]

ตารางที่ 2.2 กรอบการทำงานของการทำงานท่องเที่ยวเชิงอิเล็กทรอนิกส์โดยรวม

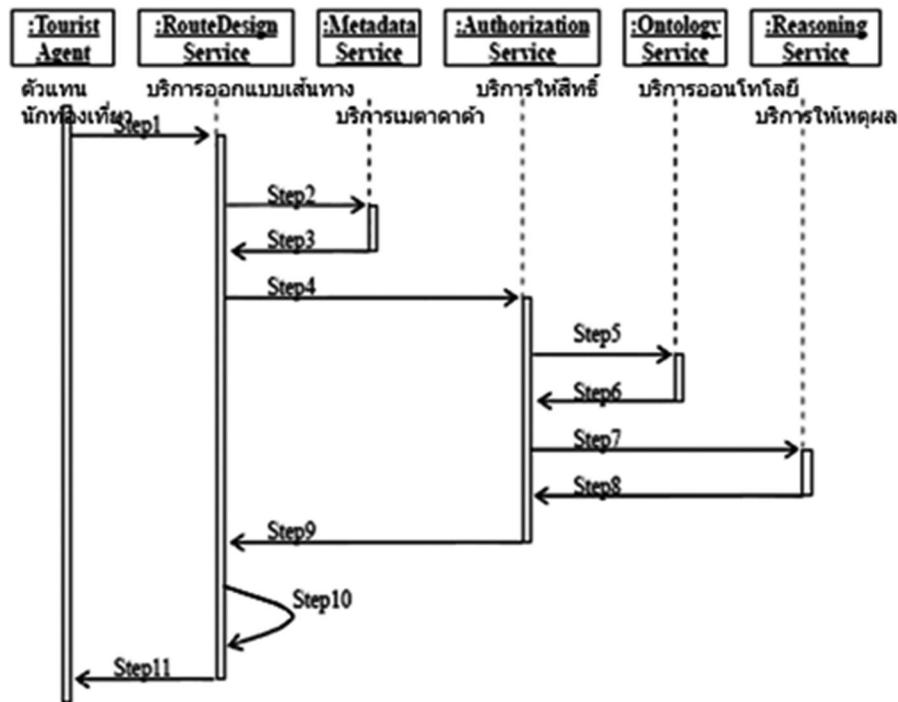
ระดับชั้น	การบริการ	หน้าที่
มุมมองการบริการ	- ผู้บริโภค - ผู้ให้บริการ - องค์กรเสมือน	นักท่องเที่ยวและตัวแทนการท่องเที่ยวใช้บริการ ตัวแทนการท่องเที่ยวและตัวแทนธุรกิจ โรงแรม สนามบิน เรือสำราญ และอื่นๆ ใช้บริการ สนับสนุนการบริการส่วนผู้บริโภค
มุมมองเนื้อหา	- ความรู้	ช่วยจัดการความรู้ระหว่างชั้นเทคนิคและงานบริการผู้ใช้ส่วนผู้ให้บริการ โดยประมวลความรู้มาจากข้อมูลและสารสนเทศที่ดึงมาจากแหล่งเว็บ แบบกระจายในชั้นเทคนิค
มุมมองเทคโนโลยี	- โครงสร้างระบบ บนกลุ่มเมฆ - ตัวแทนอัจฉริยะ - เว็บเชิง ความหมายและ บริการ	โครงสร้างระบบบนกลุ่มเมฆ ซึ่งดัดแปลงมาจากระบบกริด [1] ส่วนนี้ช่วย สนับสนุนการทำงานระหว่างองค์กรแบบเปิดและมีมาตรฐานมากขึ้น สนับสนุนการประมวลผลระหว่างระบบแบบกระจาย การเก็บข้อมูล และจัดการการไหลของงาน ระหว่างตัวแทนอัจฉริยะ และเว็บเชิง ความหมายและการบริการ ในชั้นเทคนิคมีการเชื่อมต่อกันและเชื่อมกับบริการความรู้ (Knowledge service) ในชั้นเนื้อหา มีหน้าที่หลักคือการเชื่อมต่อระบบอย่างแนบเนียน และทำการอนุมานความรู้ในระดับโมเดลและอินสแตนซ์

ในส่วนโครงสร้างออนโทโลยีประกอบด้วยส่วนแนวคิด (Concepts) บริการ (Services) ทรัพยากร (Resources) และบทบาท (Roles) โดยมีหน้าที่โดยละเอียดดังนี้

- **แนวคิด (Concepts ontology)** ทำการกำหนดคำมาตรฐานตามที่ WTO (World Tourism Organization) กำหนดไว้ซึ่งมีความเป็นสากล [5] และเป็นศัพท์ตามมาตรฐานแล้ว
- **บริการ (Services ontology)** กำหนดหน้าจอบริการ และฟังก์ชันตามมาตรฐาน OTA⁹ ตัวอย่างระดับชั้นของการบริการการท่องเที่ยว อาทิ ที่แบ่งเป็นประเภทบริการพาหนะและคมนาคม และประเภทของบริการทัวร์ที่แบ่งโดยละเอียดเป็นทัวร์ที่นำเสนอโดยตรง และบริการการค้นหาคำสนใจคือบริการการค้นหาด้านการท่องเที่ยว อาทิ จุดมุ่งหมายของการค้นหาแบ่งเป็นขนาดของข้อมูล (ขนาดมากที่สุดหรือน้อยที่สุด) และประเภทของตำแหน่งที่ตั้ง
- **ทรัพยากร (Resources ontology)** กล่าวถึงความสามารถของฮาร์ดแวร์ ซอฟต์แวร์ และการเชื่อมต่อที่สนับสนุนบริการต่างๆ ที่เปิดให้เรียกใช้ ตัวอย่างของการบริการฮาร์ดแวร์คือความสามารถของหน่วยประมวลผลกลางและความเร็วของเครือข่าย

⁹ <http://www.opentravel.org/specifications/default.aspx>

- บทบาท (Roles ontology) ประกอบด้วยการอธิบายบทบาทของผู้ใช้ คือนักท่องเที่ยว ตัวแทนการท่องเที่ยวและผู้ให้บริการต่างๆ กระบวนการให้สิทธิ์ตามหน้าที่เป็นดังภาพที่ 2.3 กระบวนการให้สิทธิ์ใน Role Ontology



ภาพที่ 2.3 กระบวนการให้สิทธิ์ใน Role Ontology [3]

ภาพที่ 2.3 แสดงกระบวนการให้สิทธิ์ใน Role ontology ซึ่งประกอบด้วย 10 ขั้นตอน คือ ขั้นตอน ที่ 1 นักท่องเที่ยวเข้าสู่องค์กรเสมือนเพื่อขอ**บริการออกแบบเส้นทาง**การเดินทางผ่านตัวแทนส่วนตัว **บริการออกแบบเส้นทาง** ทำหน้าที่เก็บคุณสมบัติของนักท่องเที่ยวตามบริการเมทาตาทาในขั้นตอนที่ 2 และ 3 ถัดมาขั้นตอนที่ 4 **บริการออกแบบเส้นทาง**สร้างการร้องขอสิทธิที่มีคุณสมบัติตามฐานความรู้ RDF ตามจุดให้บริการลูกค้า ในขั้นตอนที่ 5 และ 6 **บริการให้สิทธิ์**ได้รับออนโทโลยี VO ที่ประกอบด้วยนิยามของ Role จาก **ออนโทโลยีบริการ** และในขั้นที่ 7 และ 8 **บริการให้สิทธิ์**จะไปนำ**บริการทางเหตุผล**มาอนุมานหน้าที่ของนักท่องเที่ยวผ่าน**ออนโทโลยี องค์กรเสมือน**และพารามิเตอร์คุณสมบัติของนักท่องเที่ยว จากนั้น**บริการให้สิทธิ์**เปรียบเทียบหน้าที่ของนักท่องเที่ยวที่อนุมานได้กับหน้าที่การเข้าถึงพื้นฐานเพื่อประเมินว่านักท่องเที่ยวใช้บริการได้บ้าง ในขั้นตอนที่ 9 ถ้าอนุญาตให้เข้าถึงได้ไปขั้นตอนที่ 10 และส่งค่าเส้นทางท่องเที่ยวแก่ตัวแทนของนักท่องเที่ยว ถ้ามีปฏิเสธการเข้าถึง จะไม่คำนวณเส้นทางและตัวแทนของนักท่องเที่ยวจะได้รับข้อความปฏิเสธในการเข้าถึง

วิธีการติดตามความสนใจของนักท่องเที่ยว การวิเคราะห์ความชอบและแสดงผลอัตโนมัติ เป็นการกำหนดค่าที่สนใจไว้ในอะเรย์ IV_{ij} ($i=1...n$ เป็นจำนวนสิ่งที่สนใจ, $j=0, 1$) ถูกใช้โดยตัวแทนเพื่อแสดงระดับความสนใจของนักท่องเที่ยว

เมื่อ $j=0$, $IV_{ij}=\{\text{สิ่งที่สนใจ}\};$

เมื่อ $j=1$, $IV_{ij}=\{1,2,3,4,5\}.$

ระดับความสนใจ: กำหนดเป็นค่าดังนี้

- 1: ถ้าจำนวนบริการในเดือนน้อยกว่า 5
- 2: ถ้าจำนวนบริการในเดือนมากกว่า 5 แต่น้อยกว่า
- 3: ถ้าจำนวนบริการในเดือนมากกว่า 10 แต่น้อยกว่า 15
- 4: ถ้าจำนวนบริการในเดือนมากกว่า 15 แต่น้อยกว่า 20
- 5: ถ้าจำนวนบริการในเดือนมากกว่า 20

จากค่าที่กำหนดข้างต้นสิ่งที่สนใจคือเซตที่ได้มาจากการร้องขอจากข้อมูลของนักท่องเที่ยวในเดือนที่แล้ว ระดับความสนใจ คือค่าความถี่ของการร้องขอมาจากข้อมูลของนักท่องเที่ยวในเดือนที่แล้ว ถ้านักท่องเที่ยวร้องขอสิ่งนั้นมาก ค่าระดับความสนใจนี้จะมีมาก เช่นถ้ามีการร้องขอ ดังนี้ “Request (op: airFare, source: Beijing, destination: New York)” เป็นการร้องขอข้อมูลราคาตั๋วเครื่องบินจากปักกิ่งไปยังนิวยอร์ก ถ้าความถี่ในการร้องขอนี้ มีค่ามากกว่าระดับที่กำหนด การร้องขอนี้จะเข้าไปอยู่ในรายชื่อเส้นทางโปรด ในการเข้าใช้ระบบครั้งหน้า ระบบก็จะแสดงเส้นทางบินจากปักกิ่งไปนิวยอร์กของนักท่องเที่ยว

นักวิจัย Lee [6] ได้ทำการวิเคราะห์ค่าขอบเขต (bound) ของการท่องเที่ยวเชิงสุขภาพด้านการแพทย์ที่มีประเทศสิงคโปร์เป็นประเทศผู้นำอยู่ โดยทำการวิเคราะห์ค่าตัวแปรที่สำคัญ 2 ตัว คือ TOUR ซึ่งเป็นความต้องการของนักท่องเที่ยวมาจาก Singapore Tourism Board (STB) และ HEALTH เป็นจำนวนผู้ให้บริการทางการแพทย์ต่อประชากร 10,000 คนที่ได้มาจากระบบสาธารณสุข (Ministry of Health) ของประเทศสิงคโปร์ ประโยชน์ที่ได้จะช่วยให้การวิเคราะห์การลงทุนด้านบุคลากรของรัฐบาลในปัจจุบันให้เพียงพอต่อนักท่องเที่ยวที่เข้ามาในระยะยาว การทดสอบค่า bound ต่างๆ เป็นดังตาราง 2.3

ตารางที่ 2.3 วิเคราะห์ค่า bound ของการท่องเที่ยวเชิงสุขภาพด้านการแพทย์ที่ประเทศสิงคโปร์

การนำไปใช้	สูตร
พิจารณาจากความสัมพันธ์ระยะยาวระหว่าง 2 ตัวแปร คือ TOUR และ HEALTH แสดงด้วยโมเดล unrestricted error correction	$\Delta \text{TOUR}_t = a_0 + \sum_{i=1}^p a_{Ti} \Delta \text{TOUR}_{t-i} + \sum_{i=0}^p a_{Hi} \Delta \text{HEALTH}_{t-i} + a_1 \text{TOUR}_{t-1} + a_2 \text{HEALTH}_{t-1} + \varepsilon_{1t}$ $\Delta \text{HEALTH}_t = b_0 + \sum_{i=1}^p b_{Hi} \Delta \text{HEALTH}_{t-i} + \sum_{i=0}^p b_{Ti} \Delta \text{TOUR}_{t-i} + b_1 \text{HEALTH}_{t-1} + b_2 \text{TOUR}_{t-1} + \varepsilon_{2t}$
ค่าฐานของ bound testing มาจาก ARDL (Autoregressive distributed lag) ซึ่งสามารถสร้างโมเดล ARDL ตามสัมประสิทธิ์โดยประมาณระยะยาว	$(1 - c_1 L - \dots - c_e L^e) \text{TOUR}_t = d_0 + (1 - d_1 L - \dots - d_f L^f) \text{HEALTH}_t + \varepsilon_{3t}$ โดยค่า TOUR สามารถนำมาปรับพารามิเตอร์ตามค่า residual ในภายหลังได้
Causality test ตามงานของ Granger ที่ Lee [4] กล่าวอ้าง ผลลัพธ์คือค่า short-run (F-test) และ long-run (t-test)	$\Delta \text{TOUR}_t = g_0 + \sum_{i=1}^p g_{Ti} \Delta \text{TOUR}_{t-i} + \sum_{i=1}^p g_{Hi} \Delta \text{HEALTH}_{t-i} + \gamma \text{ECT}_{t-1} + \varepsilon_{4t}$ $\Delta \text{HEALTH}_t = h_0 + \sum_{i=1}^p h_{Hi} \Delta \text{HEALTH}_{t-i} + \sum_{i=1}^p h_{Ti} \Delta \text{TOUR}_{t-i} + \varepsilon_{5t}$

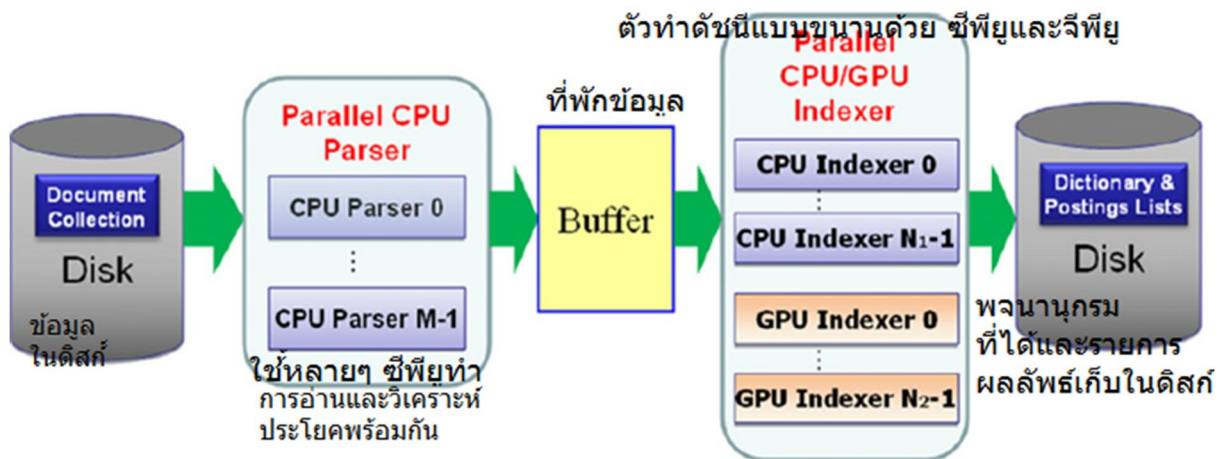
2.3 การคำนวณแบบขนานด้วย MapReduce

การทำงานในรูปแบบ MapReduce เป็นการสร้างโมเดลสำหรับการคำนวณแบบขนานสำหรับระบบงานแบบกระจายสำหรับข้อมูลปริมาณมาก ประกอบด้วยขั้นตอนการแมป (Map) โดยจะทำการแมปงานย่อยไปยัง worker ต่างๆ เพื่อแก้ปัญหาแล้วผลลัพธ์ชั่วคราวที่ได้จะอยู่ในรูป $\langle \text{key}, \text{value} \rangle$ และนำออกมาส่งต่อให้ worker เพื่อทำการรีดิวซ์ (Reduce) ซึ่งเป็นการรวมผลลัพธ์เข้าด้วยกัน โดยโหนดมาสเตอร์จะนำ $\langle \text{key}, \text{value} \rangle$ มาจับกลุ่มอัตโนมัติ โดยข้อมูลที่มี key เดียวกันจะเข้ามาอยู่รวมกลุ่มกัน

มีงานวิจัยเป็นจำนวนมากที่ทำงานแบบขนานในขั้นตอนการค้นหาข้อมูลนี้ เพื่อส่งคำตอบกลับไปยังผู้ใช้งาน และยังมีการประมวลผลดัชนีการค้นหาด้วยวิธี MapReduce โดยตรง ในการทำดัชนีจะใช้เทอม (term) ที่กำหนดเป็น key และ ID ของเอกสารเป็น value ซึ่ง worker ที่ทำการ Reduce จะรับรายชื่อ ID ที่ยังไม่ได้เรียงลำดับมาโดยตรง และได้มีงานวิจัยที่ใช้ MapReduce แบบที่ทำในรอบเดียว (Single-Pass) มาปรับปรุง โดย McCreadie และคณะวิจัยใน [6] ได้นำเสนอให้ Map worker ส่ง $\langle \text{term}, \text{partial posting list} \rangle$ แทน เพื่อลดจำนวนการส่งและขนาดของการส่งทั้งหมดระหว่างขั้นตอนการ Map และการ Reduce เนื่องจากการรวมเทอมที่ซ้ำกัน จะลดความถี่ในการส่งลง กลยุทธ์นี้ทำให้ความเร็วของการประมวลผลเพิ่มขึ้นเมื่อจำนวนของคอร์ของเครื่องมีมากขึ้น

Lin และคณะ [7] ได้พัฒนาวิธี MapReduce แบบ scalable MapReduce Indexing โดยการเปลี่ยน $\langle \text{term}, \text{posting}\{\text{document ID}, \text{term frequency}\} \rangle$ เป็น $\langle \text{tuple}\{\text{term}, \text{document ID}\}, \text{term frequency} \rangle$ การทำเช่นนี้ส่งผลส่วนให้ key นั้นไม่ซ้ำ และรับรองได้ว่า MapReduce ส่งค่าไปยังขั้น Reduce ตามลำดับ ส่งผลให้ผลลัพธ์ที่ส่งไปเรียงกันโดยไม่เหลือผลลัพธ์ให้รอไปทำยังรอบต่อไป ซึ่งการทำงานแบบนี้ Lin และ Dyer [7] ได้เผยแพร่เป็นหนังสืออัลกอริธึมสำหรับ MapReduce ที่รู้จักกันดีว่ามีอัตรา throughput ที่เหมาะสำหรับการทำดัชนีของข้อความแบบเต็ม (full text)

ยังมีงานวิจัยเกี่ยวกับการประมวลผลแบบขนานด้วยอัลกอริธึม MapReduce ในขั้นตอนการทำดัชนี ได้ใช้หน่วยประมวลผลกลางและหน่วยประมวลผลกราฟิกส์แบบผสม [8] อัลกอริธึมของการทำดัชนีแบบนี้ยังเป็นการลดรายการข้อมูลที่เก็บลงชั่วคราว ค่า document ID ของเอกสารถูกเก็บแบบเรียงลำดับอยู่ในแต่ละรายการที่ต้องใช้ทำงานต่อไป ดังนั้นแนวคิดอย่างง่ายคือจะทำการเข้ารหัสที่ช่องว่างระหว่าง document ID ของเอกสารที่ติดกัน 2 รายการแทนการรวมค่าตามการเข้ารหัสแบบไบนารี (byte encoding) และแบบ γ และการบีบอัดแบบ Golomb



ภาพที่ 2.4 การไหลของข้อมูลสำหรับการทำดัชนีด้วยวิธีการแบบขนาน [10]

ภาพที่ 2.4 เป็นการนำหน่วยประมวลผลกราฟิกส์หรือจีพียู (GPU) มาช่วยหน่วยประมวลผลกลางหรือซีพียู (CPU) ในการประมวลผล โดยได้มีการทดสอบแล้วว่า การนำหน่วยประมวลผลกราฟิกส์เข้ามาช่วยประมวลผล ทำให้สามารถลดต้นทุนและเวลาได้ วิธีการก็คือทำการแบ่งข้อมูลและส่งไปวิเคราะห์ประโยค (parse) พร้อมกันเป็นลักษณะ Trie-Collection นำเข้าไปทำงานในหน่วยประมวลผลกลางและในหน่วยประมวลผลกราฟิก จากการทดลองพบว่าการนำหน่วยประมวลผลกราฟิกส์เข้ามาช่วยส่งผลให้ประสิทธิภาพการทำงานดีกว่า [9,10]

นอกจากการสร้างดัชนีเพื่อจัดการแสดงผลแก่ผู้ใช้แล้วยังมีการทำเหมืองเว็บ (Web mining) ด้วยเว็บเชิงความหมายเพื่อแสดงคำตอบแก่ผู้ใช้อย่างเหมาะสม รวมทั้งการแปลงดัชนีแบบอื่น [9] นอกจากนี้ยังม้งานวิจัยที่การศึกษาขั้นตอนการทำงานแบบขนานของวิธีการ MapReduce ในกระบวนการอื่นที่เกี่ยวข้อง [10] ทำการพัฒนาวิธีการประมวลผลแบบขนานที่รองรับการค้นหาข้อมูลดังกล่าวที่อาศัยหลักการแบ่งและส่งข้อมูลขนาดใหญ่ผ่านเครือข่ายประยุกต์ใช้บิตเมตริกซ์ของ G.T. Williams, Medha Atre และคณะ [12] และใช้หลักการของ MapReduce แบ่งการจัดการเป็นการประมวลผลแบบขนานนั้น อ้างอิงหลักการการคิวรีแบบขนานของRDF/OWL บน MapReduce [13] การประมวลผลด้วยภาษา SPARQL สำหรับการทำงานบนกลุ่มเมฆ [14] และการเก็บข้อมูลแบบ Triple สำหรับ MapReduce [15] การปรับปรุงมิตเดลแวร์เพื่อควบคุมการทำงานของหน่วยประมวลผลกราฟิก [16] อัลกอริธึมการทำงานแบบขนานในส่วนของกฎและข้อมูลสำหรับออนโทโลยี [17] ที่กล่าวถึงในแผนงานวิจัยแล้วและเทคนิคการกระจายข้อมูลสำหรับแอปพลิเคชันแบบขนานในลักษณะโฮตส์ และไคลเอนท์ [18]

ในงานวิจัยของ Hung-chih Yang, Ali Dasdan, Ruey-Lung Hsiao และ D. Stott Parker [19] เป็นการเพิ่มการทำงานในส่วนของการรวมกันในการทำ MapReduce จากเดิมที่ไม่สนับสนุนการประมวลผลสำหรับข้อมูลที่ต่างกัน จึงได้ทำการพัฒนาตัวแบบขึ้นมาใหม่ ให้มีประสิทธิภาพในการรวมข้อมูลที่ถูกแบ่งเป็นส่วนๆ หรือถูกแบ่งประเภทไว้ ซึ่งนำมาใช้กับการค้นหากับข้อมูลจำนวนมหาศาล

ในที่นี้รูปแบบ MapReduce นั้นสืบทอดมาจาก Google MapReduce ยกเว้นแต่คุณสมบัติย่อยๆ เท่านั้นที่เปลี่ยนแปลง ที่เพิ่มเข้าไปได้แก่ ฟังก์ชันการรวม (merge) ฟังก์ชัน processor ฟังก์ชันการแบ่ง (partition selector) และการทำซ้ำแบบปรับได้ (configurable iterator) ในส่วนการรวม ผู้ใช้อาจต้องการเหตุผลในการประมวลผลข้อมูลต่างๆ ไป ทั้งนี้ขึ้นอยู่กับข้อมูลด้วย หลังจากงานของ Map และ Reduce ทำเสร็จแล้ว จะเรียกการทำการรวมไปยังคลัสเตอร์แต่ละโหนดเช่นเดียวกันกับ MapReduce ที่ต้องพิจารณาเหตุผลในการประมวลผลข้อมูลตามข้อมูลที่เข้ามา หน้าที่ฟังก์ชันต่างๆ ได้แก่

- **Partition selector:** ในการรวม ให้ผู้ใช้กำหนดฟังก์ชัน partition selector เพื่อมากำหนดข้อมูลที่แบ่งแล้ว ซึ่งผลิตโดย up-stream reducers จะถูกเรียกไปยังการรวมหรือ merge ต่อ
- **Processor function:** ตัว processor ถูกวางอยู่ในที่ที่ผู้ใช้สามารถกำหนดตรรกะของการประมวลผลข้อมูลจากแต่ละแหล่งที่มา
- **Merge function:** ผู้ใช้สามารถกำหนดตรรกะของการประมวลผลข้อมูล ในการรวมข้อมูลจาก 2 แหล่งที่มา ที่ซึ่งข้อมูลเหล่านี้เหมาะสมกับเงื่อนไขการรวมได้
- **Configurable iterator:** เพื่อจัดการความสัมพันธ์ของตัว merger ของ logical iterator จำนวน 2 ตัว ผู้ใช้สามารถระบุอัลกอริธึมการรวม ต่างๆ เองได้

Djoerd Hiemstra และ Claudia Hauff [20] กล่าวถึงการค้นหาข้อมูลจากเอกสารหน้าเว็บที่เก็บรวบรวมไว้ ได้แก่ ClueWeb09 ด้วยการใช้ MapReduce บน Hadoop ช่วยในการดึงข้อมูลขนาดใหญ่นี้ โดยระบบจะสแกนเอกสารทั้งหมดและอ่านแต่ละหน้าเว็บเพียงครั้งเดียว เพื่อสกัดข้อความออกมา ทั้งยังมีการเก็บสถิติคำศัพท์ที่มีอยู่ในเอกสารเหล่านั้น แล้วทำการค้นหาแบบเชิงเส้น จากนั้นจะลบหน้าเว็บที่ไม่เกี่ยวข้องออกไป การทดลองแบ่งเป็น 4 ประเภท คือ การค้นหาแบบธรรมดา การค้นหาแบบบอกหัวเรื่อง การค้นหาแบบลบหน้าเว็บที่ไม่เกี่ยวข้องออกไป สุดท้ายการค้นหาแบบบอกหัวเรื่องและลบหน้าเว็บที่ไม่เกี่ยวข้องออกไป ผลการทดลองพบว่าแต่ละประเภทการค้นหาหามีค่าความแม่นยำ (Precision) เพิ่มขึ้นตามลำดับ แสดงให้เห็นว่ามีจำนวนเอกสารที่เกี่ยวข้องกับความต้องการและถูกสกัดออกมาได้มากกว่าจำนวนเอกสารที่มีอยู่ทั้งหมด เนื่องจากระบบสามารถกำจัดเอกสารที่ไม่เกี่ยวข้องออกไปได้เป็นจำนวนมาก แต่ถ้ามีจำนวนเอกสารเพิ่มมากขึ้นโดยแบ่งประเภทการค้นหาตามรูปแบบเดิม พบว่ามีค่าความแม่นยำลดลงตามลำดับ เนื่องจากมีจำนวนเอกสารที่ไม่ได้เกี่ยวข้องมากนั่นเอง

Jeffrey Dean และ Sanjay Ghemawat [21] นำเสนอการประยุกต์การประมวลผลแบบ MapReduce บนกลุ่มคอมพิวเตอร์ขนาดใหญ่ อธิบายการทำงานและยกตัวอย่างโปรแกรม โดยสรุปได้ดังนี้ MapReduce เป็นรูปแบบการเขียนโปรแกรมและการดำเนินการต่างๆ ที่เกี่ยวข้องกับการประมวลผลและสร้างชุดข้อมูลขนาดใหญ่ ซึ่งทำงานแบบขนานได้อย่างอัตโนมัติ สามารถคำนวณข้อมูลได้ปริมาณมาก โดยแบ่งข้อมูลเป็นกลุ่มเล็กๆ ทั้งยังจัดเวลาการทำงานเมื่อมีการกระทำซ้ำเครื่อง ดูแลการสื่อสารระหว่างคอมพิวเตอร์ และจัดการข้อผิดพลาดที่เกิดขึ้น วิธีการนี้ช่วยให้โปรแกรมเมอร์ที่ไม่มีประสบการณ์กับระบบแบบขนานและแบบกระจายเขียนโปรแกรมได้ง่ายขึ้น เพราะมีการซ่อนรายละเอียดของการทำงานภายในต่างๆ ไว้ วิธีการนี้ได้รับแรงบันดาลใจจากวิธี Map และ Reduce ของการทำงานในหลายๆ ภาษาโปรแกรม

การทำงานของ MapReduce แบ่งเป็น 2 ส่วน คือ Map และ Reduce โดยการทำงานของ Map จะแบ่งข้อมูลที่ได้รับมาเป็นขนาดเล็กๆ ข้อมูลที่ถูกแบ่งได้รับการประมวลผลจากคอมพิวเตอร์เครื่องต่างๆ ในเครือข่าย โดยมีโหนดหลัก (master) เป็นคอมพิวเตอร์เครื่องหนึ่งทำหน้าที่กำหนดและควบคุมการทำงานทั้งหมดของเครื่องคอมพิวเตอร์ที่เหลือซึ่งเรียกว่า โหนดทำงาน (worker หรือ slave) เมื่อประมวลผลเสร็จแล้ว จึงเก็บข้อมูลลงบัฟเฟอร์ในดิสก์หรือหน่วยความจำของเครื่องตัวเอง จากนั้นการทำงานของ Reduce ก็จะมาอ่านข้อมูลในบัฟเฟอร์แล้วจัดกลุ่มข้อมูล โดยโหนดหลักจะเก็บสถานะของการทำงาน Map และ Reduce คือ ว่าง (idle) กำลังทำงาน (in-progress) และทำงานเสร็จสมบูรณ์ (complete) โหนดหลักมีการติดต่อกับโหนดทำงานเป็นระยะๆ ถ้าหากไม่มีการตอบสนองกลับมาก็จะจดจำว่าโหนดทำงานนั้นล้มเหลว และเปลี่ยนสถานะโหนดนั้นให้เป็นว่าง

การที่มีข้อจำกัดด้านทรัพยากรเครือข่าย MapReduce จึงเก็บข้อมูลนำเข้าไว้ในหน่วยความจำของเครื่องตัวเอง ทำให้ลดปริมาณข้อมูลที่ต้องส่งผ่านเครือข่าย และเพื่อความสะดวกในการเข้าถึงและค้นหาข้อมูลอีกด้วย และด้วยภาระงานของ Map และ Reduce ที่แตกต่างกัน จึงต้องมีการจัดสมดุลของภาระงานแบบพลวัต (Dynamic load balance) เพื่อให้การทำงานของแต่ละเครื่องมีประสิทธิภาพสูงสุด และสามารถกู้คืนได้เร็วหากเกิดข้อผิดพลาด

การทำงานแบบ MapReduce สนับสนุนการอ่านและเขียนข้อมูลได้หลายรูปแบบ โดยเพิ่มข้อมูลแต่ละประเภทจะรู้ด้วยตัวเองว่าควรแบ่งเพิ่มข้อมูลอย่างไรเพื่อให้สื่อความหมายในการประมวลผล ขณะที่โหนดหลักทำงานจะมีการรายงานสถานะของข้อมูล คือ ความก้าวหน้าของการคำนวณ เช่น จำนวนงานที่เสร็จแล้ว จำนวนงานที่กำลังประมวลผล หรือขนาดของข้อมูล เพื่อให้ผู้ใช้ทำนายเวลาที่ใช้ในการทำงานได้ รวมทั้งตัดสินใจได้ว่าควรเพิ่มทรัพยากรอะไรให้กับระบบเพื่อการคำนวณบ้าง เช่น รายงานการใช้ซีพียู การใช้ดิสก์ รายงานแบนด์วิดท์ของเครือข่าย

การทำงานที่สามารถใช้การคำนวณแบบ MapReduce ได้เช่นโปรแกรม Grep ที่มีการสแกนและค้นหาข้อความจากเอกสารขนาดใหญ่ โดยแบ่งข้อมูลออกเป็นกลุ่ม ผลการคำนวณพบว่าอัตราที่ข้อมูลถูกสแกนสูงขึ้นเมื่อจำนวนเครื่องคอมพิวเตอร์มากขึ้น งานที่สำคัญอีกชิ้นหนึ่งคือ สร้างระบบดัชนี ซึ่งผลิตโครงสร้างข้อมูลสำหรับการบริการค้นหาหน้าเว็บของ Google เป็นต้น

นอกจากนี้ยังมีการงานวิจัยใน [22] เน้นการจัดสรรทรัพยากรแบบทำซ้ำ (Iterative resource allocation) สำหรับหน่วยความจำ สำหรับอัลกอริธึมการค้นหาแบบขนานบนกลุ่มเมฆ กริด และคลัสเตอร์ที่ใช้งานร่วมกัน อัลกอริธึมแบบ iterative allocation นี้ทำงานช้าจนกว่าจะสามารถแก้ปัญหาได้เสร็จสิ้น จะเห็นได้ว่าอัลกอริธึมนี้สามารถทำงานได้ง่าย แต่มีส่วนสำคัญคือ ฟังก์ชัน *Increase()* ซึ่งรับค่าของจำนวนของ HAUus เพื่อจัดสรรพื้นที่หน่วยความจำในรอบต่อไป ดังอัลกอริธึมต่อไปนี้

Algorithm 1 Generic, Iterative Allocation (IA)

```
numHAUs  $\leftarrow$  1
while true do
    result  $\leftarrow$  Run algorithm a on problem p
    if result = solved then
        return success
    else
        numHAUs  $\leftarrow$  Increase(numHAUs)
```

2.4 งานวิจัยที่เกี่ยวข้องด้านการประมวลผลกับข้อมูลขนาดใหญ่

การประมวลผลแบบขนานกับข้อมูลปริมาณมากในปัจจุบัน ส่วนใหญ่จะใช้อัลกอริธึม MapReduce [21] เป็นพื้นฐาน โดยมีการแบ่งข้อมูลออกเป็นแฟ้มข้อมูลก่อน จากนั้นจึงจับคู่แฟ้มข้อมูลที่มีปัญหาแบบเดียวกันเข้าด้วยกัน เมื่อแก้ปัญหาสเสร็จแล้วจึงนำคำตอบที่ได้มารวมกัน จากตารางที่ 2.4 แสดงถึงงานวิจัยที่มีการประมวลผลแบบขนานบนกลุ่มเมฆ โดย Rohloff และ Schantz [23] ใช้กรอบการทำงานแบบ MapReduce แล้วสามารถทำการประมวลผลข้อมูลขนาดใหญ่บนกลุ่มเมฆได้อย่างรวดเร็วร่วมกับการวางแผนการประมวลผล ขณะที่การประมวลผลข้อมูลของ Zhang [3] มีลักษณะเป็นการประยุกต์ใช้กริดที่เน้นในขอบเขตของการท่องเที่ยว ดังตารางที่ 2.4

ตารางที่ 2.4 การประมวลผลแบบขนานบนกริดและบนกลุ่มเมฆ

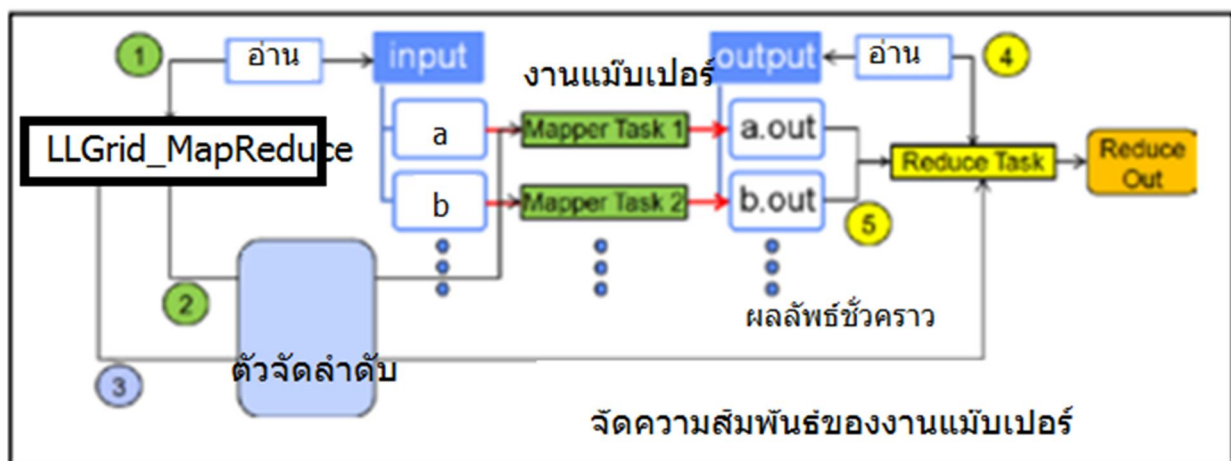
งานวิจัย	ตัวอย่างที่ใช้	สิ่งแวดล้อม	วิธีการ	สรุปผล
Rohloff และ Schantz [15,23]	14 คิวรี (LUBM) 6000 มหาวิทยาลัย (ประมาณ 800 Triple)	ใช้ Hadoop บน Cloudera บริการบนกลุ่มเมฆ ของ Amazon EC2 ปริมาณโหนด คำนวณ 20 XL	ใช้กรอบการทำงาน แบบ MapReduce บนระบบคลาวด์ สำหรับระดับล่าง โดยใช้อินพุต RDF แบบแฟ้มข้อมูล N3	ระบบกระจายการประมวลผลขนาดใหญ่ที่มีประสิทธิภาพสูงโดยใช้กรอบการทำงานแบบ MapReduce เป็นซอฟต์แวร์ใช้การเก็บข้อมูลแบบ Triple ชื่อ SHARD ในระยะสั้นและเพื่อการปรับปรุงในระยะยาว ใช้การแคชของข้อมูลในหน่วยความจำและสำรองข้อมูลบนดิสก์ในระดับซอฟต์แวร์
Zhang[3]	เตรียมแต่ละชั้นของบริการกริดเพื่อโดเมนการท่องเที่ยวอิเล็กทรอนิกส์	บริการกริดสองชั้นคือชั้นของส่วนประกอบพื้นฐานและชั้นแอปพลิเคชันแพลตฟอร์มของกริดคือ Globus 4 และมิดเดิลแวร์คือ Jena	ชั้นส่วนประกอบพื้นฐานเป็นไปตาม S-OGSA ใช้กรอบการทำงานของการทำางานของ OntoGrid ชั้นแอปพลิเคชันใช้บริการการท่องเที่ยวอิเล็กทรอนิกส์เช่น RouteDesign	นำเสนอกริดเชิงความหมายเพื่อโดเมนการท่องเที่ยวอิเล็กทรอนิกส์สำหรับเว็บเซอร์วิสใช้ Apache Axis และ WSRF เพื่อประกาศและประมวลการร้องขอ SOAP การออกแบบออนโทโลยีเป็นการท่องเที่ยวอิเล็กทรอนิกส์แบบ OTA-Compliant ของ TourSearchRQService

สำหรับงานวิจัยที่มีโครงสร้างข้อมูลและอัลกอริธึมแบบขนานนั้น ขึ้นอยู่กับระบบสถาปัตยกรรมระบบและฮาร์ดแวร์ด้วยเช่นกัน จากตารางที่ 2.5 นำเสนองานวิจัยที่มีการประมวลผลข้อมูลแบบข้อความเช่นเดียวกันแต่ Lin และคณะ [24] ได้ใช้ขั้นตอนวิธีแบบ MapReduce บนหน่วยประมวลผลกลาง แต่ Zheng และ Jala [8] นำข้อมูลชุดเดียวกันมาจัดรูปแบบใหม่และนำหน่วยประมวลผลกราฟิกมาช่วยในการประมวลผลร่วมปรากฏว่าสามารถทำงานได้มากกว่าในเวลาน้อยกว่า และได้ทดลองวิธีเดียวกันกับฐานข้อมูลอื่นด้วย

ตารางที่ 2.5 โครงสร้างข้อมูลและอัลกอริธึมแบบขนานบนฮาร์ดแวร์ต่างๆ

ผู้วิจัย	ตัวอย่าง	สิ่งแวดล้อม	วิธีการ	สรุปผล
Lin และคณะ [24]	เซตข้อมูลเป็น ClueWeb09	ใช้ MapReduce แบบ Ivory จำนวน 99 โหนด แต่ละโหนดมีสองแกนรวม 198 แกนบนคลัสเตอร์หน่วยประมวลผลกลาง	การประมวลผลข้อความด้วยการทำดัชนีข้อความทั้งหมดและขยายขนาดการใช้ของอัลกอริธึม MapReduce [8,9 in project II] มี Throughput 91.96 MB ต่อวินาที	สร้างรายการโพสต์ตำแหน่งโดยอัลกอริธึม MapReduce แบบ Ivory ซึ่งยังมีการคิดค่าใช้จ่ายบางส่วนอยู่
Zheng และ JaJa [8]	เซตข้อมูลเป็น ClueWeb09 และ Wikipedia01-07	ใช้ CPU และ GPU ใน 1 โหนดประกอบด้วยหน่วยประมวลผลกลาง 8 แกนและหน่วยประมวลผลกราฟิกส์จำนวน 2 ตัว) ผลลัพธ์ที่ดีที่สุดเกิดจากใช้หน่วยประมวลผลกลาง 6 แกนวิเคราะห์คำและ 2 แกนทำดัชนีร่วมกับหน่วยประมวลผลกราฟิกส์อีกสองตัว	สร้างชุด Index แบบไทร (trie) จาก 0 – 17612 มี throughput 262.76 MB/s ให้ผลดีกว่า Lin [24]	ให้เอกสารจากดิสก์หมายเลข 0 เป็นเอกสารนำเข้าชั่วคราวจากการวิเคราะห์ ถือเป็นขั้นตอนเตรียมการเพื่อทำดัชนีแบบขนานระหว่างซีพียู และหน่วยประมวลผลกราฟิกส์เป็นขั้นเตรียมการของการทำรายการในดิสก์หมายเลข 1 เช่นกัน

สำหรับวิธีการประมวลผลและภาษาที่คล้ายคลึงกับงานวิจัยนี้ได้แก่งานเรื่องการประมวลผลข้อมูลขนาดใหญ่ด้วยคลัสเตอร์ MPI (Message Passing Interface¹⁰) Byun และคณะ [25] ได้มีการนำเสนออัลกอริทึม MapReduce กับงานใดๆ โดยการทำงานเริ่มจากการสแกนแฟ้มข้อมูลอินพุตจากไดเรกทอรีหรืออ่านจากรายการในแฟ้มอินพุตขั้นที่ 1 แล้วจึงเข้าถึงกระบวนการตั้งตารางงานในขั้นที่ 2 โดยการสร้างตารางงานจากหลายๆ งานที่เกิดขึ้น และเรียกตารางนี้ว่าอะเรย์ของงาน แล้วตั้งชื่อให้เป็น งานแมปเพอร์ (Mapper task) ที่ 1 งานแมปเพอร์ที่ 2 ไปเรื่อยๆ ซึ่งในการนี้อาจใช้โปรแกรมช่วยให้มีการทำงานที่พร้อมกัน (concurrent) เมื่ออะเรย์ของงานถูกสร้างและส่งไปประมวลผลแล้ว แต่ละแฟ้มข้อมูลที่เป็นอินพุตจะถูกประมวลผลโดยงานหนึ่งๆ ที่กำหนดไว้ตามคำสั่งเรียกว่าแมปเพอร์ในภาพที่ 2.5 สำหรับในส่วนนี้นั้นแอปพลิเคชันที่สร้างขึ้นสามารถประมวลในรูปแบบใดก็ได้เช่นทำงานบนเชลสคริปต์ (Shellscript) ทำงานด้วยภาษาจาวาหรือทำงานด้วยโปรแกรมใดๆ ที่เขียนด้วยภาษาต่างๆ



ภาพที่ 2.5 การทำงานของ LLGrid MapReduce

เมื่ออ่านแฟ้มอินพุตทั้งหมดแล้ว จะแจกงานไปยังงานแมปเปอเรอร์ ตัวจัดลำดับมาควบคุมลำดับการทำงานของงานแมปเปอเรอร์อีกที เมื่อแต่ละแมปเปอเรอร์ได้ผลลัพธ์ชั่วคราวแล้ว จะส่งไปงาน Reduce เพื่อรวบรวมในขั้นตอนสุดท้าย

พารามิเตอร์ที่เกี่ยวข้องในงานนี้จึงประกอบด้วย จำนวนของงาน วิธีการ Map วิธีการ Reduce ไดเรกทอรีของอินพุตและเอาต์พุต งานวิจัยนี้นำวิธีการนี้มาประยุกต์ใช้ด้วยการนำข้อมูลจากเว็บไซต์มาประมวลผลแบบขนานก่อนในการดึงข้อมูลขึ้นมาเพื่อใส่ลงในออนโทโลยีต่อไป

2.5 งานวิจัยที่เกี่ยวข้องด้านการแทนค่าข้อมูล

ปัจจุบันมีระบบการเก็บข้อมูลแบบ Triple เป็นจำนวนมาก และมีการใช้โครงสร้างข้อมูลแบบกราฟหรือต้นไม้ ช่วยให้สามารถใส่ข้อมูล Triple ได้มากขึ้น ตัวอย่างของระบบ Triple ที่เป็นที่รู้จัก เช่น¹¹ AllegroGraph

¹⁰ http://en.wikipedia.org/wiki/Message_Passing_Interface

¹¹ <http://www.w3.org/wiki/LargeTripleStores>

(ความจุ 1+Trillion) OpenLink Virtuoso v6.1 (ความจุ 15.4Billion+ explicit; uncounted virtual/inferred) BigOWLIM (12B explicit, 20B total) โดยคิดค่าบริการ 100,000 คิวรี่ต่อ \$1 Garlik 4store (15 Billion) Bigdata® (12.7Billion) YARS2 (7 Billion) Jena TDB (1.7Billion) Jena SDB (650 Million) Mulgara (500 Million) RDF gateway (262 Million) Jena with PostgreSQL (200 Million) Kowari (160 Million) 3store ด้วย MySQL 3 (100 Million) และ Sesame (70Million) เป็นต้น ในรายชื่อเหล่านี้มี OWLIM ที่สามารถใช้กับ Amazon EC2 ได้โดยคิดค่าบริการ 1 US\$ ต่อการคิวรี่ SPARQL จำนวน 100,000 คิวรี่

สำหรับกรณีซอฟต์แวร์ Open source มีระบบเก็บฐานความรู้แบบ Triple ที่ทำงานบน Hadoop โดยใช้โครงสร้างข้อมูลกราฟได้แก่ SHARD [26] โดยที่การออกแบบ SHARD ประกอบด้วยฐานข้อมูลแบบ Triple ที่อยู่บน Hadoop File System (HDFS) ที่มีเมธอดเรียกจากไคลเอนท์ ประมวลผลบนกลุ่มเมฆด้วยงานจาก MapReduce ขยายขนาดได้ รองรับการทำงานของคิวรี่ SPARQL และมีการอนุมานขั้นระดับพื้นฐาน และในงานวิจัย [23] ได้นำเสนออัลกอริธึม Clause-Iteration ทำงานบนกราฟโดยใช้วิธีการแบบ MapReduce เช่นกัน

2.6 การประมวลผลบนกลุ่มเมฆกับการบริการเว็บเชิงความหมาย

ตารางที่ 2.6 นำเสนอทฤษฎีการประมวลผลบนกลุ่มเมฆ ตามระดับขั้นของการบริการ จะเห็นว่าในการประมวลผลบนกลุ่มเมฆมีระดับขั้น 3 ระดับได้แก่ Software as a Service, Platform as a Service และ Infrastructure as a Service ในแต่ละขั้นสอดคล้องโครงสร้างบริการของเว็บเชิงความหมายดังนี้

ตารางที่ 2.6 ทฤษฎีที่เกี่ยวข้องตามระดับขั้นของการประมวลผลบนกลุ่มเมฆ

การประมวลผลบนกลุ่มเมฆ	การบริการแบบเว็บเชิงความหมาย
Software as a Service (SaaS)	เว็บแอปพลิเคชัน HTML5 บริการเว็บเซอร์วิส ของผู้ให้บริการต่างๆ เช่นสายการบิน ตัวแทนท่องเที่ยว API เพื่อการส่งออกไปยังเครื่องมือต่างๆ
Platform as a Service (PaaS)	ชั้นตัวกลางของ Semantic Web (เช่น Reasoner, Inference, Query) ตัวจัดการอโนโลยี (Ontology Manager) (XML syntax: RDF, OWL) Web Annotation แอปพลิเคชัน Java Runtime
Infrastructure as a Service (IaaS)	เครื่องเซิร์ฟเวอร์ เครือข่าย ที่จัดเก็บอโนโทโลยี

โดยระบบเว็บเชิงความหมายจะมีความแตกต่างกับระบบที่เป็นอยู่ในปัจจุบัน สรุปได้ดังตาราง 2.7

ตารางที่ 2.7 ความแตกต่างและความได้เปรียบของระบบปัจจุบันกับระบบเว็บเชิงความหมายที่ใช้อยู่

ประเด็นปัญหา	การออกแบบเว็บในปัจจุบันที่มีอยู่ ตัวอย่างคือเว็บของเทศบาล เว็บการท่องเที่ยวแห่งประเทศไทย และเว็บแนะนำที่พักแรม	เว็บเชิงความหมายจากงานวิจัย มีลักษณะคือเว็บที่ตอบผลการค้นหาจากผู้ใช้ตามโครงสร้างออนโทโลยี
การ ค้นหา และนำเสนอแบบหลายมิติ	ค้นหาโดยใช้ตัวหนังสือและมีลำดับของการนำเสนอที่ไม่สัมพันธ์กันจำนวนมาก เช่น ผลการค้นหาคำว่า Cicada ใน Google Search Engine สามารถแสดงผลทางภควิทยาได้ด้วยแต่ความตั้งใจของผู้ค้นหาคือสถานที่ในหัวหิน	ค้นหาตามโครงสร้างออนโทโลยีที่มีความสัมพันธ์ตามโดเมนที่สร้างหลายมิติ เช่น การค้นหาคำว่า Cicada เมื่ออยู่ในโดเมนอำเภอหัวหินแล้วเป็นชื้อตลาด อยู่ที่ตำบลหนองแก
ลักษณะการนำเสนอของเว็บไซต์	นำเสนอแบบแสดงข้อความที่ผู้สร้างต้องการแสดง ซึ่งมีความหมายเพื่อมนุษย์อ่านและเข้าใจ เบื้องหลังมีฐานข้อมูลที่เป็นข้อความเหมือนกันเมื่อเปิดอ่าน มนุษย์ก็สามารถเข้าใจได้ แต่เครื่องจักรไม่เข้าใจสิ่งเหล่านี้	นำเสนอแบบแสดงข้อความที่ผู้สร้างต้องการแสดง ซึ่งมีความหมายเพื่อมนุษย์อ่านและเข้าใจ แต่เบื้องหลังระดับเครื่องสามารถทำให้เครื่องเข้าใจกันได้โดยใช้โครงสร้างออนโทโลยีแบบ XML และเครื่องต้น-ปลายทางสามารถแลกเปลี่ยนบริการกันได้ ตัวอย่างเช่น การรับข้อมูลที่มีการปรับอัตโนมัติ ข่าวสาร RDF จากเหตุการณ์ต่างๆ
วิวัฒนาการของการขยายตัวของโครงสร้างข่าวสาร	ข่าวสารถูกเก็บและบันทึกในที่เดียวกัน การออกแบบสร้างเพื่อดึงข้อมูลจากฐานข้อมูลเดียว แบบบนลงล่างและบำรุงรักษาจากส่วนกลาง	ข่าวสารถูกออกแบบเพื่อเชื่อมต่อกับแหล่งอื่น จึงเป็นแบบขยายกิ่งโครงสร้าง โดยพัฒนาจากล่างขึ้นบน และมีการอัปเดตแบบกระจายได้
การขยายโครงสร้างของชุมชนการสื่อสาร	ชุมชนเพิ่มข่าวสารและแสดงผลตามโครงสร้างของเว็บ เช่นเว็บบอร์ด	ชุมชนสามารถเพิ่มขึ้นและโครงสร้างใหม่และเพิ่มข่าวสารตามโครงสร้างนั้นได้ เช่นออนโทโลยีที่ออกแบบเหมือนกันสามารถนำโครงสร้างมาใช้ใหม่ หรือนำออนโทโลยีโดเมนการท่องเที่ยวในชุมชนอื่นมาเชื่อมกันได้ ตัวอย่างเช่นการขยายโครงสร้างของลิงค์เดต้า (linkeddata.org) นั่นคือใช้โครงสร้างเดียวกันสามารถขยายโดเมนการท่องเที่ยวระดับอำเภอเป็นทั่วประเทศได้และขยายไปสัมพันธ์กับโดเมนอื่นที่เกี่ยวข้องกันได้
การขยายมุมมองของชุมชนการสื่อสาร	การจัดเก็บเนื้อหาและการจัดการอยู่ที่ส่วนกลาง	จัดการและจัดเก็บเนื้อหาข้อมูลแบบกระจายได้ แต่อยู่ในมุมมองรูปแบบเดียวกัน

ประเด็นปัญหา	การออกแบบเว็บในปัจจุบันที่มีอยู่ ตัวอย่างคือเว็บเทศบาล เว็บการท่องเที่ยวแห่งประเทศไทย และเว็บแนะนำที่พักแรม	เว็บเชิงความหมายจากงานวิจัย มีลักษณะคือเว็บที่ตอบผลการค้นหาจากผู้ใช้งาน โครงสร้างออนโทโลยี
การรวมทรัพยากรที่กระจายอยู่	ผู้ดูแลระบบต้องหาข้อมูลของแต่ละเว็บ แต่ละชุดเอกสาร เพื่อปรับปรุงที่ละหน้าผ่านฟอร์ม เมื่อต้องการบำรุงรักษาต้องทำทีละสำเนา	เนื่องจากผู้พัฒนาผลิตโครงสร้างและข้อมูลที่นำกลับมาใช้ใหม่ได้ ผู้ดูแลระบบสามารถอัปเดตข้อมูลที่เดียวแล้วกระจายข้อมูลหลายเว็บเพจได้ เช่นจุดบริการจอดรถใหม่สามารถเพิ่มคุณสมบัติการค้นหาสถานที่ใกล้เคียง เพื่อเชื่อมต่อกับสถานที่ๆ ตั้งบนถนนใกล้เคียงกันได้
การรวมข่าวสารข้ามเว็บ	เป้าหมายเว็บอยู่ที่การเข้าชมของคน เมื่อต้องการใช้ข้อมูลร่วมกันจะเป็นลิงค์เชื่อมโยงไปหา เช่นลิงค์แหล่งท่องเที่ยวจากเว็บเทศบาลไปยังเว็บอื่นๆ	เมื่อมีการอนุญาตระหว่างเจ้าของเว็บเครื่องสามารถสื่อสารเพื่อเข้าถึงข่าวสารโครงสร้างเดียวกันผ่านบริการได้ เช่นการดึงข้อมูลราคาของสถานที่พักแรมจากเว็บชื่อดังมาเปรียบเทียบกัน

2.7 เทคโนโลยีของ Hadoop

ปัจจุบันปัญหาเกี่ยวกับการจัดการข้อมูลขนาดใหญ่ กำลังได้รับความสนใจเป็นอย่างมากทั้งทางภาคธุรกิจและทางด้านงานวิจัย เนื่องจากข้อมูลต่างๆ ที่จัดเก็บมีความสำคัญ เพื่อใช้ในการสืบค้น วิเคราะห์ รวมทั้งประมวลผลข้อมูล อีกทั้งข้อมูลเหล่านั้นมีความล้าบakyงยากมากขึ้น ดังนั้นปัญหาที่สำคัญในการจัดการกับข้อมูลขนาดใหญ่ คือ การเก็บข้อมูลและการค้นหาข้อมูล เพื่อแก้ปัญหาในการจัดเก็บข้อมูลขนาดใหญ่ จึงได้มีการนำเทคโนโลยีที่เรียกว่า Hadoop เข้ามาช่วย

Hadoop เป็นซอฟต์แวร์เฟรมเวิร์ค (Framework) ที่ถูกออกแบบมาเพื่อทำงานบนระบบคอมพิวเตอร์แบบกระจาย และสนับสนุนการทำงานแบบขนาน โดยมีชุดคำสั่ง API เพื่อช่วยอำนวยความสะดวกแก่นักพัฒนาแอปพลิเคชันที่จะสร้างระบบค้นหาหรือวิเคราะห์ข้อมูลขนาดใหญ่ อาจกล่าวได้ว่า Hadoop เป็นส่วนพื้นฐานที่ทำหน้าที่ทั้งประมวลผลแบบขนาน และจัดเก็บข้อมูลแบบกระจาย

Hadoop ใช้สำหรับทำงานกับแอปพลิเคชันบนคลัสเตอร์ของคอมพิวเตอร์ขนาดใหญ่ ซึ่งเหมาะสำหรับแอปพลิเคชันที่ต้องการความน่าเชื่อถือ โดย Hadoop ถูกสร้างจากแบบจำลองการคำนวณคือ Map Reduce ซึ่งจะทำให้การแบ่งแอปพลิเคชันออกเป็นชิ้นงานเล็กๆ โดยชิ้นงานแต่ละชิ้นอาจถูกทำซ้ำ โดยโหนดหลักจากนั้นก็ทำการรวบรวมงานที่ถูกกระจายออกไปในแต่ละโหนดถูกกลับเข้ามาเป็นชิ้นเดียว นอกจากนี้ Hadoop ยังมีระบบแฟ้มแบบกระจาย (Distributed File System) ซึ่งเรียกว่า Hadoop Distributed File System (HDFS) โดยจะรองรับการส่งรับข้อมูลจำนวนมากระหว่างคลัสเตอร์ และจัดเก็บแฟ้มข้อมูลขนาดใหญ่ที่มีความเชื่อถือได้ ซึ่งมีผลต่อการจัดการอินเทอร์เน็ต และแบนด์วิดท์ โดย HDFS จะทำการสำรองข้อมูลบนโหนดคำนวณ (compute node) และเตรียมช่องทางการสื่อสารไว้สำหรับการส่งข้อมูลข้างคลัสเตอร์ ซึ่งทั้ง Map/Reduce และ HDFS ได้ถูกออกแบบมาเพื่อให้เป็นเฟรมเวิร์กที่สามารถจัดการกับโหนดที่ทำงานล้มเหลวหรือไม่เสถียรสามารถทำงานต่อไปได้โดยอัตโนมัติด้วย

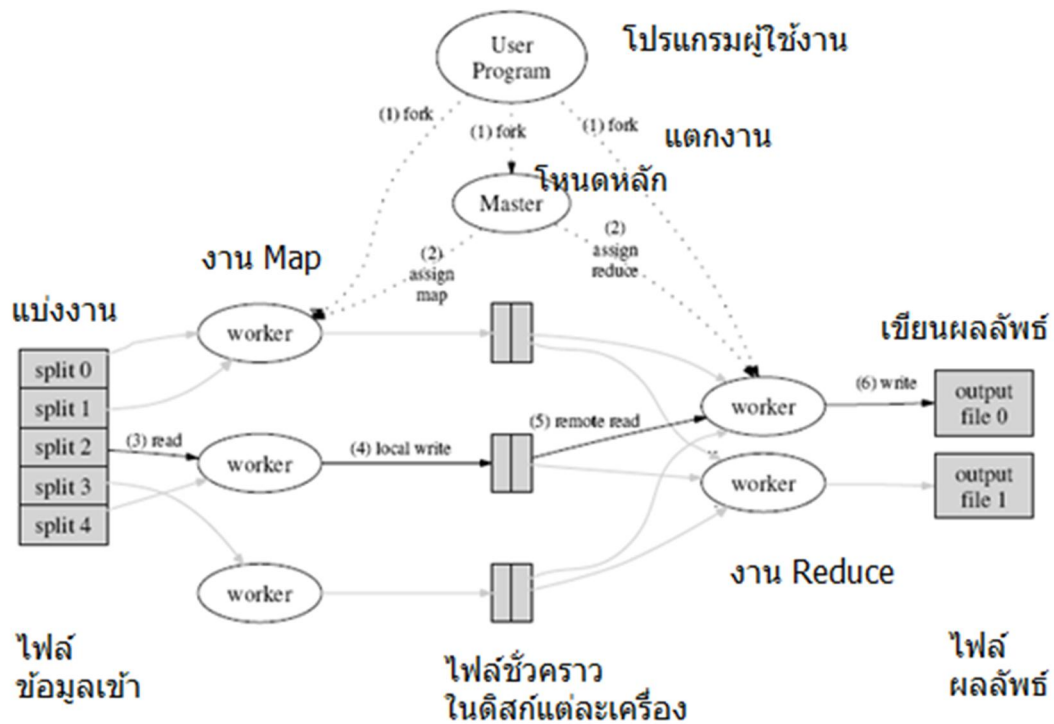
ตัว Hadoop ประกอบไปด้วยองค์ประกอบ 2 ส่วนหลักๆ ดังนี้

- **Hadoop Distributed File System (HDFS)** ซึ่งมีความสามารถในการจัดเก็บข้อมูลขนาดใหญ่ โดยมีหลักการทำงานคือ แบ่งข้อมูลขนาดใหญ่ที่ต้องการจะเก็บออกเป็นส่วนย่อยๆ แล้วกระจายข้อมูลส่วนย่อยๆ นั้นไปยังเครื่องคอมพิวเตอร์ต่างๆ ที่เชื่อมต่อสำหรับการแก้ปัญหาในส่วนของ การค้นหาข้อมูลภายในเทคโนโลยี HDFS วิธีที่เหมาะสมที่สุดเรียกว่าวิธี MapReduce
- **MapReduce Framework** ซึ่งเป็นส่วนของการปฏิบัติงานในการประมวลผลข้อมูลขนาดใหญ่ ซึ่งเป็นการส่งคำสั่งค้นหาแบบกระจายไปยังโหนดต่างๆ ในระบบ โดยที่ไม่จำเป็นต้องมีการย้ายข้อมูลระหว่างการประมวลผล ซึ่งจะทำให้การค้นหาข้อมูลมีประสิทธิภาพทางด้านความเร็วในการค้นหามากยิ่งขึ้น

MapReduce เป็นเฟรมเวิร์กในการเขียนโปรแกรมแบบหนึ่งที่ช่วยในงานประมวลผลที่มีชุดข้อมูลจำนวนมาก หลักการทำงานจะเป็นการทำงานแบบขนานซึ่งจะอาศัยหลายๆ โหนดช่วยกันทำงาน โดยที่ผู้ใช้งานนั้นไม่ต้องสนใจเบื้องหลังการทำงานเช่น วิธีการทำงานแบบขนาน วิธีการแบ่งข้อมูล การจัดการสมดุลของงาน การจัดการความคงทน เป็นต้น ในการทำงานแล้วผู้ใช้งาน MapReduce จะสนใจแค่ส่วนของ Map และส่วนของ Reduce ซึ่ง Map จะทำการจับคู่ของ Key และ Value ที่ต้องการ แล้วก็จะส่งไปให้ Reduce ทำการประมวลผลเพื่อให้ได้ผลลัพธ์รวมที่ต้องการ

MapReduce และ Google's File System ได้คิดการประมวลผลที่เน้นการประมวลกับระบบข้อมูลขนาดใหญ่ โดยเน้นที่การค้นหาข้อมูลลักษณะเดียวกับ Search Engine โดยข้อมูลจะมีรูปแบบเดียวกันทำงานกับหลายๆ ตัว Mapper แต่เพื่อเอื้ออำนวยกับการโปรแกรมแบบขนานกรณีอื่นๆ อาจจะมีการเพิ่มการสร้างฟังก์ชัน Merge ขึ้นมา เพื่อรวมข้อมูลต่างๆ กันได้

การทำงานของ MapReduce คือจะกระจายงานต่างๆ ไปให้งาน Map อยู่บนแต่ละเครื่องทำงาน ซึ่งผู้ที่ควบคุมการกระจายงานก็คือโหนดหลัก Master โดยหลังจากที่ worker ทำงานเสร็จแล้วก็จะแจ้งโหนดหลักเพื่อที่โหนดหลักจะส่งต่อผลการการทำงานที่ได้จากการ Map ให้กับ Reduce worker เพื่อทำงานให้ได้ผลลัพธ์ต่อไปสามารถแสดงแผนผังการทำงานของ MapReduce ได้ดังภาพ 2.6

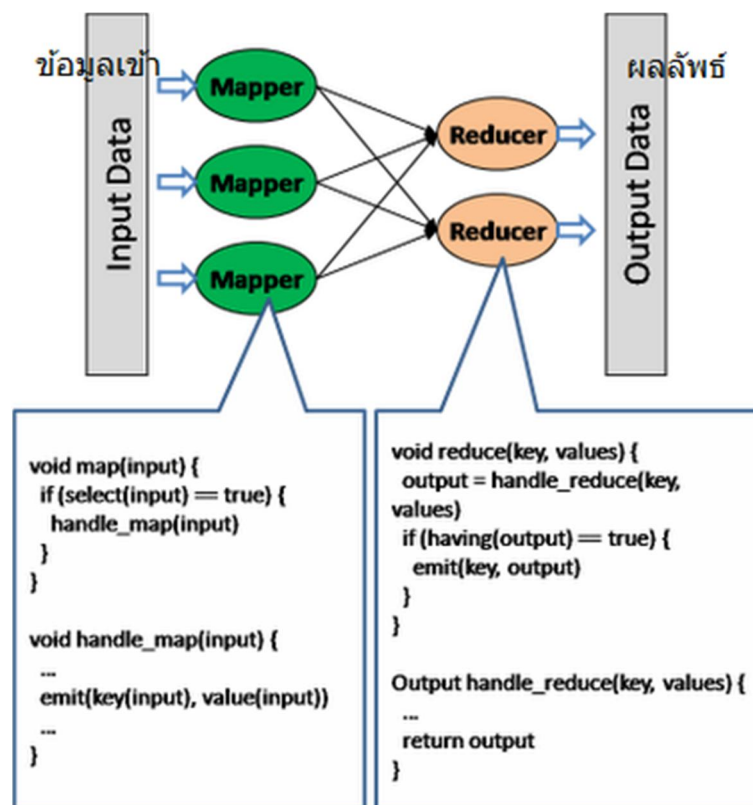


ภาพที่ 2.6 แผนผังการทำงานของ MapReduce [21]

จากภาพที่ 2.6 สามารถอธิบายการทำงานของ MapReduce ได้ดังต่อไปนี้

- เริ่มจากการ Map จะแบ่งแฟ้มข้อมูลที่เป็นข้อมูลเข้าเป็นส่วนย่อยๆ ขนาดประมาณ 16 MB ถึง 64 MB
- โหนดหลักหรือ Master จะพิจารณาว่ามี worker ตัวไหนว่างอยู่บ้าง ก็จะมอบงาน (assign task) ให้ ซึ่งก็จะมีทั้งงาน Map และ Reduce
- สำหรับ worker ที่ถูกมอบงาน Map ก็จะทำอ่านข้อมูลจากอินพุตแฟ้มข้อมูลที่ถูกแบ่งแล้ว Map Key และ Value ตามที่โปรแกรมผู้ใช้งานไว้
- ผลของการ Map จะถูกเก็บไว้ที่เครื่องที่ทำการ Map ซึ่งจะมีการแบ่งเป็นส่วนๆ ตาม Key ที่กำหนด และจะมีการส่งข้อความไปยัง Master เพื่อให้ Master ส่งไปทำงาน Reduce ต่อไป
- เมื่อ Reduce worker ได้รับแจ้งจาก Master จะทำการอ่านแฟ้มข้อมูลจากตำแหน่งที่ Master กำหนดให้มาแล้วทำการจัดเรียงข้อมูลตามลำดับ Key
- เมื่อเรียงข้อมูลตาม Key เสร็จแล้วก็จะทำตามฟังก์ชัน Reduce ตามที่ผู้ใช้งานกำหนดไว้ โดยจำนวนผลลัพธ์ที่ออกมาจะขึ้นอยู่กับจำนวนของ Reduce worker แล้วสุดท้ายก็จะมีผลรวมแฟ้มผลลัพธ์เป็นแฟ้มเดียว
- เมื่อ MapReduce เสร็จแล้วก็จะแจ้งให้โหนด Master รู้ เพื่อแจ้งต่อไปยังโปรแกรมว่าทำ MapReduce เสร็จแล้ว

โมเดลการเขียนโปรแกรมสามารถแสดงได้ดังภาพที่ 2.7

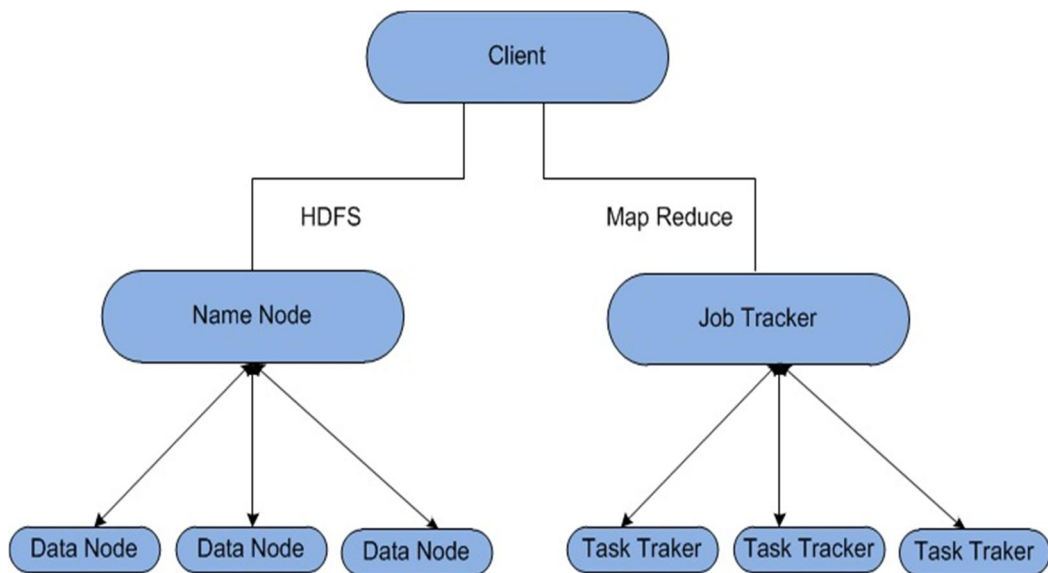


ภาพที่ 2.7 โมเดลการเขียนโปรแกรม [27,28]

เริ่มต้นด้วยการแบ่งงานที่รับเข้ามาเป็นส่วนๆ และทำการส่งแต่ละส่วนของงานนั้นไปยังเครื่องต่างๆ ซึ่งแต่ละเครื่องจะทำงานตาม Map script บนส่วนของข้อมูลที่บันทึกอยู่ โดย Map script จะป้อนข้อมูลบางส่วน และ Map ไปยังคู่ <key,value> ตามข้อกำหนด ตัวอย่างเช่น ถ้าต้องการนับความถี่ของคำในข้อความหนึ่ง key และ value อาจหมายถึง <word, count> เป็นต้น จากนั้น Map script จะเก็บคู่ของ <word,1> สำหรับ word ใดในที่เก็บอินพุตชั่วคราว (input stream) จุดประสงค์ของ Map script คือตัวแบบข้อมูลที่กำหนดให้ Reduce ทำการรวม จากนั้นคู่ <key,value> จะถูกสลับ (shuffled) โดยทั่วไปหมายถึงการที่คู่กับ key เดียวกันถูกจัดกลุ่มและส่งผ่านไปยังเครื่องเดียว ซึ่งจะทำงานในส่วนของ Reduce script ต่อไป นั่นคือ Reduce script จะเก็บรวบรวมคู่ <key,value> และทำการ Reduce ตาม Reduce script ที่ผู้ใช้ระบุ ในตัวอย่างของการนับคำนี้ ต้องการนับจำนวนของคำที่เกิดขึ้นเพื่อให้สามารถได้มาซึ่งความถี่ ดังนั้นจึงทำการ Reduce script เพียงค่าผลรวมค่า 1 ในส่วนของ value ภายใน <key,value> ที่มี key เดียวกัน

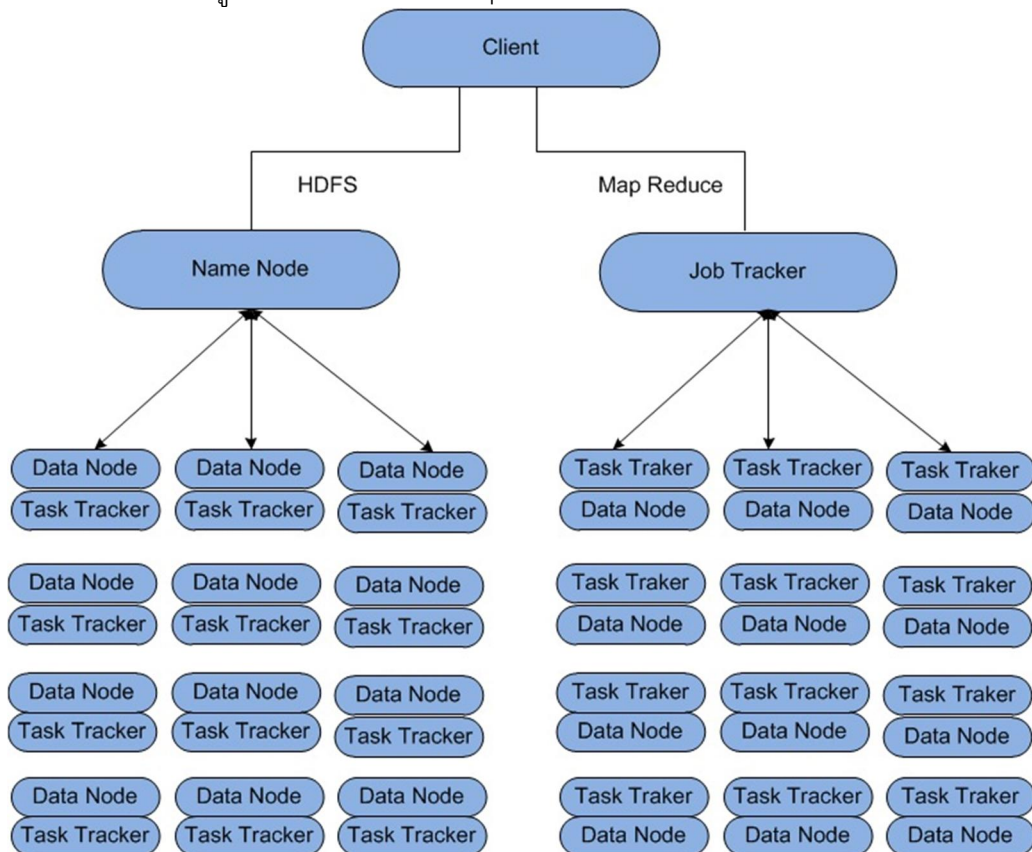
จากภาพที่ 2.8 ใน Hadoop จะสร้างจากโหนดหลักต่างๆ ดังนี้

- Name Node (Master)
- Data Node (Slave)
- Job Tracker
- Task Tracker



ภาพที่ 2.8 สถาปัตยกรรมระดับ Logical ของ Hadoop [29]

จากภาพที่ 2.8 เป็นสถาปัตยกรรมเชิง Logical ของ Hadoop แต่โดยกายภาพแล้ว Data Node และ Task Tracker สามารถถูกกำหนดบนเครื่องต่างๆ กันได้ ดังแสดงในภาพ 2.9



ภาพที่ 2.9 การวาง Data node และ Task Tracker บนเครื่องๆ เดียว [30]

การทำงานของ MapReduce ตาม Job Tracker และ Task Tracker

- ทำการรับงาน MapReduce ที่ถูกกำหนดโดยผู้ใช้
- การ Map งาน และ Reduce งานไปยัง Task Tracker
- ตรวจสอบสถานะของงาน และ Task Tracker และเริ่มทำงานใหม่อีกครั้งหากงานนั้นไม่สำเร็จ

ในส่วน Task Tracker โหนดลูกของ MapReduce ทำงานดังนี้

- การ Map และ Reduce งานภายในคำสั่งที่ได้จาก Job Tracker
- จัดการพื้นที่และการส่งผลลัพธ์ชั่วคราว

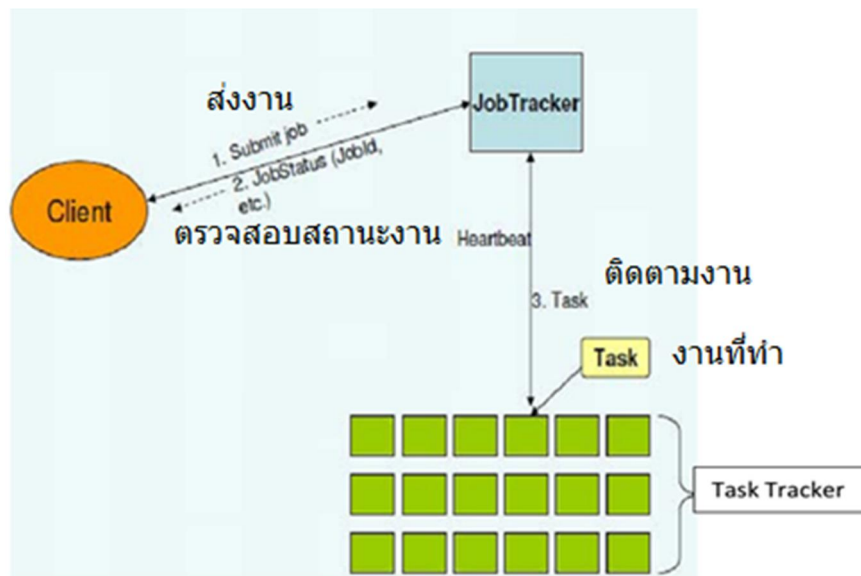
สำหรับ Job Tracker มีหน้าที่การทำงานได้แก่

- ดูแลงานทั้งหมด
- ทำการวางแผนตารางเวลาการทำงานทั้งหมด
- แบ่งงานออกเป็นงานย่อยๆ แล้วทำตามลำดับ
- จัดตารางงานบนโหนดที่ข้อมูลใกล้เคียงกัน โดยการแบ่งงานที่รับเข้ามา
- ตรวจสอบแต่ละงาน
- ยกเลิกและเริ่มทำงานนั้นใหม่ในกรณีที่งานนั้นล้มเหลวหรือสูญหายไป

สำหรับ Task Tracker มีหน้าที่ดังต่อไปนี้

- ร้องของานใหม่
- ทำงาน
- ตรวจสอบสถานะของงาน
- รายงานผลการทำงานกลับโหนดหลัก

ซึ่งสามารถแสดงลำดับการทำงานได้ดังภาพ 2.10



ภาพที่ 2.10 ลำดับการทำงานของงาน MapReduce ใน Hadoop [30]

2.8 Java Concurrency [32]

การทำงานแบบพร้อมกัน (Concurrency) หมายถึง การทำงานหลายๆ งานในเวลาเดียวกัน สำหรับทั้งนี้ Java ได้ออกแบบมาเพื่อรองรับการทำงานหลายๆ งานในเวลาเดียวกัน ซึ่งทั้งหมดนี้จะอยู่ใน Package ของ `java.util.concurrent`

ในการพัฒนาโปรแกรมให้ทำงานแบบพร้อมกัน (Concurrent programming) ต้องเข้าใจพื้นฐานของโปรเซส และเธรด โดยสามารถอธิบายได้ดังนี้

- **โปรเซส** คือกระบวนการทำงานที่มีหน่วยความจำของตัวเอง หรือสามารถทำงานด้วยตัวเองได้เลย เช่น แอปพลิเคชันต่างๆ ไปนั่นเอง หรือ อาจจะเป็นโปรเซสที่ไม่มีหน้าต่างกราฟฟิกส์ หรือ User interface แต่อาจจะทำงานด้านหลังตลอดเวลา
- **เธรด** คือโปรเซสขนาดเล็ก (lightweight process) ซึ่งการใช้เธรด จะใช้ทรัพยากรของเครื่องคอมพิวเตอร์น้อยกว่าในการรันโปรเซส โดยทุกแอปพลิเคชันจะมีอย่างน้อยหนึ่งเธรดเสมอ ซึ่งอาจเรียกได้ว่าเธรดหลัก (main thread)

ปัญหาสำคัญที่เกิดจากการทำงานแบบพร้อมกันของหลายๆ เธรด ได้แก่

- การแทรกแซงระหว่างเธรด ตัวอย่างเช่นการเขียนค่าที่ผิดพลาดระหว่างสองเธรด สมมติมีตัวแปร a และมี 2 เธรดมาทำการเขียนค่าพร้อมกัน ค่าที่ได้จะเปลี่ยนไป นั่นคือ เริ่มจากค่า a มีค่าเป็น 0 จากนั้นถูกเปลี่ยนเป็น 1 โดย Thread1 ต่อมา Thread2 มาทำการลบออกไปทำให้ค่ากลับเป็น 0 อีก เป็นต้น

- ความสอดคล้องของค่าในหน่วยความจำ เป็นการดึงข้อมูลมาอ่านผิดพลาด เช่นจากตัวอย่างข้างต้น Thread1 ต้องการดึงข้อมูลมาดูเป็น 1 แต่ว่า Thread2 ได้ลบเป็น 0 แล้ว เป็นต้น

เทรตเป็นลำดับการประมวลผลคำสั่งที่มีอยู่ภายในโพรเซส เนื่องจากเทรตนั้นมีคุณสมบัติบางส่วนของโพรเซสอยู่ด้วย จึงอาจเรียกได้ว่าเป็นโพรเซสขนาดเล็ก ในโพรเซสหนึ่งอาจประกอบด้วยเทรตจำนวนมากที่กำลังประมวลผลอยู่ โดยทั่วไปเทรตเป็นวิธีการที่นิยมนำมาใช้ในการเพิ่มประสิทธิภาพของการประมวลผลโพรเซสผ่านกลไกการประมวลผลแบบขนาน โดยที่หน่วยประมวลผลกลางจะสามารถสลับการทำงานระหว่างเทรตต่างๆ ได้อย่างรวดเร็ว ทำให้เกิดภาพเสมือนคล้ายกับว่าเทรตทั้งหมดนั้นกำลังประมวลผลอยู่ในเวลาเดียวกัน โพรเซสแบบเก่าจะเป็นโพรเซสที่ประกอบด้วยเทรตเพียงเทรตเดียว ซึ่งอาจจะอยู่ในสถานะใดๆ ก็เป็นไปได้ (running, blocked, ready หรือ terminated) เทรตแต่ละตัวจะมีสแตค (stack) เป็นของตนเองเนื่องจากเทรตมักจะเรียกใช้โพรซีเยอร์หรือฟังก์ชันต่างๆ กันจึงทำให้มีรายการทำงานที่แตกต่างกัน ระบบปฏิบัติการที่สนับสนุนการทำงานของเทรตจะกำหนดให้การประมวลผลในระดับพื้นฐานของหน่วยประมวลผลกลาง นั่นคือเทรตด้วย เนื่องจากเทรตไม่มีความเป็นอิสระจากเทรตอื่น (ภายในโพรเซสเดียวกัน) ดังนั้นเทรตในกลุ่มเดียวกัน (ภายในโพรเซสเดียวกัน) จึงใช้งานโค้ดและข้อมูล และทรัพยากรที่จัดสรรให้โดยระบบปฏิบัติการ เช่น การใช้งานแฟ้มข้อมูลต่างๆ ร่วมกัน

โปรแกรมภาษาจาวา (Java) จะทำงานภายในโพรเซสของตนเองภายใต้ Java Virtual Machine ซึ่งสนับสนุนการใช้งานได้หลายเทรตในหนึ่งโปรแกรมและถือเป็นส่วนหนึ่งของภาษาด้วย นอกจากนี้ Java ได้มีการจัดเตรียมการล็อกโค้ดบางส่วนเพื่อป้องกันหลายๆ เทรตทำงานโค้ดส่วนเดียวกัน ทั้งนี้ใช้คำสำคัญ Synchronized ป้องกันส่วนของเมธอดหรือทั้งเมธอด การใช้วิธีการนี้ เรียกว่าเป็นการ Synchronization เป็นวิธีการหนึ่งของการสื่อสารระหว่างเทรต เพื่อสร้างลำดับการเข้าถึงทรัพยากรที่เหมาะสมที่ให้ผลลัพธ์การทำงานที่ถูกต้อง

ปัญหาที่ควรระวังในการใช้ Synchronization ก็คือ Deadlock นั่นคือเป็นการรอที่ไม่มีที่สิ้นสุดนั่นเอง เนื่องจากเวลาการใช้ Synchronized นั้นให้เทรตทำงานทีละตัว เพราะฉะนั้น ถ้าเมธอดนั้นไปเรียกอีกเมธอดหนึ่งที่ถูกล็อก Synchronized อยู่อาจจะทำให้เกิด Deadlock ได้

ที่กล่าวมาข้างต้นเป็นการทำงานในแอปพลิเคชันขนาดเล็กเท่านั้น ในการทำงานแอปพลิเคชันขนาดใหญ่จะมีการใช้ Executor ซึ่งเป็น Java Interface ที่เตรียมไว้สำหรับทำหน้าที่ไปเรียกวัดูให้ทำงาน มีความสามารถพิเศษในการจัดการกับวัตถุ Runnable จำนวนมากๆ จึงนิยมไว้ใช้กับโปรแกรมที่ต้องการรับมือกับเทรตจำนวนมาก และมีการสร้างกลุ่มของเทรต (Thread Pool) ก็เป็นการกำหนดจำนวนของเทรตที่ใช้ทำงานโดยสามารถกำหนดเทรตที่เกี่ยวข้องเท่านั้นให้ทำงาน

เฟรมเวิร์ก Fork-Join ใน Java 7

ใน Java เวอร์ชัน 7 ได้แนะนำกลไกการทำงานแบบขนานแบบใหม่เพื่อใช้ในการคำนวณงานที่เข้มข้นขึ้น นั่นคือเฟรมเวิร์ก Fork-Join โดยช่วยให้สามารถแจกจ่ายงานบางงานไปยังหลาย worker และคอยผลลัพธ์โดย Fork-Join เป็นโครงสร้างที่สนับสนุนการเขียนโปรแกรมแบบขนาน ซึ่งเป็นเทคนิคที่ออกแบบมาง่ายและมีประสิทธิภาพมากที่สุดเพื่อให้ได้การทำงานแบบขนานที่มีประสิทธิภาพดี โดยอัลกอริธึมของ Fork-Join เป็น

อัลกอริธึมแบบขนานที่อยู่ในกลุ่มของอัลกอริธึมแบบแบ่งแยกและเอาชนะ (divide-and-conquer) สามารถแสดงได้ตามโค้ดดังต่อไปนี้

```
Result solve(Problem problem) {  
    if (problem is small)  
        directly solve problem  
    else {  
        split problem into independent parts  
        fork new subtasks to solve each part  
        join all subtasks  
        compose result from subresults  
    }  
}
```

จากโค้ดข้างต้นขั้นตอน Fork เริ่มทำงานโดยทำการแบ่งงานออกเป็นงานย่อยๆ ส่วน Join จะเกิดขึ้นเมื่องานปัจจุบันไม่สามารถดำเนินการได้จนกระทั่งงานย่อยๆ ที่ถูก Fork สำเร็จ โดยวิธีการ Fork-Join จะทำงานคล้ายกับการแบ่งแยกและเอาชนะและเป็นการเรียกตัวเองซ้ำ (Recursive) แบ่งงานออกเป็นงานย่อยๆ วนซ้ำไปเรื่อยๆจนกระทั่งปัญหานั้นเล็กเพียงพอที่จะสามารถแก้ได้ง่ายโดยวิธีการแบบลำดับขั้นๆ

เฟรมเวิร์ก Fork/Join เป็นรูปแบบใหม่ที่ถูกเพิ่มในแพลตฟอร์ม Java 7 และเป็นหนึ่งในข้อดีหลักใน Java 7 โดยในเฟรมเวิร์ก Fork/Join มีการเพิ่มคลาส 2 คลาสหลักๆเข้าไปใน Package ของ java.util.concurrent นั่นคือ

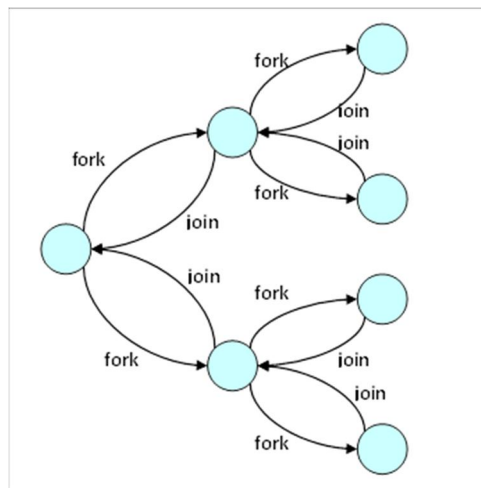
- ForkJoinPool
- ForkJoinTask

คลาส ForkJoinPool เป็นคลาส Executor และรันอินสแตนซ์ของ ForkJoinTask หลักการทำงานของ ForkJoinPool คือทำงานแบบ “work stealing” ซึ่งเทรดจะพยายามค้นหาและดำเนินการกับงานย่อยซึ่งถูกสร้างโดยงานอื่นที่ยังคงทำงานอยู่ โดยอินสแตนซ์ของ ForkJoinTask เป็นเทรดขนาดเล็กมากเมื่อเทียบกับเทรด Java ปกติ เมื่อ ForkJoinTask จะเริ่มต้นทำงาน จะเริ่มที่งานย่อย แล้วรอให้งานย่อยที่กำลังทำงานอยู่เสร็จสิ้น ซึ่งคลาส Executor ForkJoinPool เป็นผู้รับผิดชอบสำหรับการกำหนดงานย่อยนี้ไปยังอีกเทรดหนึ่งใน ThreadPools ซึ่งขณะนี้จะคอยให้บางงานทำงานเสร็จ โดยเทรดที่กำลังทำงานอยู่ใน ThreadPools จะพยายามทำงานย่อยอื่นที่ถูกสร้างโดยงานอื่นๆ ซึ่งจะขึ้นอยู่กับจำนวนของโพรเซสเซอร์ที่สามารถหามาได้ ForkJoinPool พยายามที่จะทำงานร่วมกันในระดับการทำงานแบบขนาน สำหรับ ForkJoinTask มีฟังก์ชันหลักหรือเมธอดที่เพิ่มเติมไป คือ

-fork() เป็นเมธอดที่ตัดสินใจบนการทำงานแบบ Asynchronous ของ ForkJoinTask โดย เมธอดนี้สามารถสร้างงานใหม่ได้

-join() เป็นเมธอดที่รับผิดชอบในการส่งคืนค่าผลลัพธ์ที่ผ่านการคำนวณสำเร็จสมบูรณ์แล้ว โดยยอมให้งานหนึ่งคอยให้อีกงานหนึ่งทำงานสำเร็จก่อน

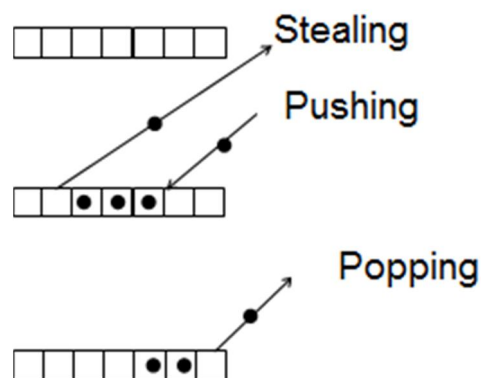
ซึ่งสามารถแสดงภาพการทำงานของ Fork/Join โดยรวมได้ดังภาพ 2.11



ภาพที่ 2.11 แสดงภาพรวมของการทำงานของ Fork/Join

(ที่มา - Java 7 Fork/Join Framework. Accessed December 25.
<http://www.developer.com/java/java-7-forkjoin-framework.html>)

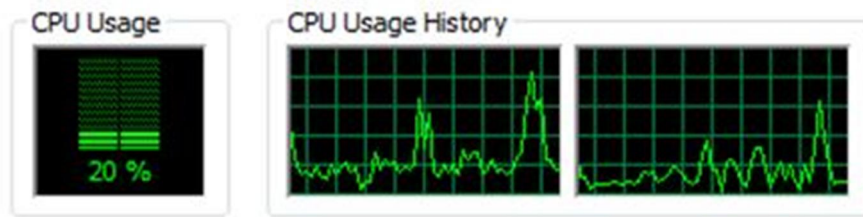
หัวใจสำคัญของ Fork/join อยู่ที่กลไกการจัดตารางงาน แต่ละเธรดทำงาน โดยมีการเก็บงานที่จะทำไว้ในคิวงานของตัวเอง งานย่อยถูกทำตามลำดับแบบ LIFO (Last In First Out) ถ้าในคิวนั้นไม่มีงานที่จะต้องทำ มันจะชิงงานจากโหนดอื่นในแบบ FIFO (First In First Out) ดังภาพ 2.12 ที่เรียกว่า Work Stealing [32]



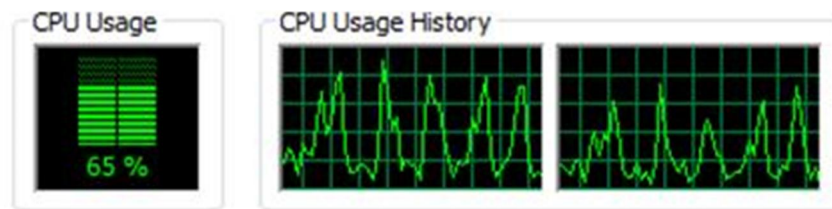
ภาพที่ 2.12 กระบวนการ Work Stealing

นอกจากนี้สิ่งที่สำคัญมากที่จะทราบว่าการใช้งานหน่วยประมวลผลกลางในกรณีของ Single thread จะมีประสิทธิภาพน้อยกว่าเมื่อเปรียบเทียบกับการใช้เฟรมเวิร์ก Fork/Join ซึ่งสามารถแสดงได้ดังภาพ 2.13

เทรดเดี่ยว (Single Threaded Model)



เฟรมเวิร์ก Fork/Join



ภาพที่ 2.13 การทำงานของหน่วยประมวลผลกลางระหว่างเทรดเดี่ยว (Single Threaded Model) กับเฟรมเวิร์ก Fork/Join

(ที่มา - <http://www.developer.com/java/java-7-forkjoin-framework.html>)

การเขียนโปรแกรมแบบขนานในภาษา Java พัฒนามาจากตัวแบบเทรดเดี่ยวมาเป็นเฟรมเวิร์ก Fork/Join ที่ถูกแนะนำใน Java 7 โดยตัว Executor ที่เป็นการปรับปรุงที่สำคัญที่เกิดขึ้นกับ Java ตั้งแต่ Java 5 แต่ไม่สามารถถูกใช้สำหรับการคำนวณแบบขนานได้ ดังนั้นเฟรมเวิร์ก Fork/Join ใน Java 7 จึงเป็นพัฒนาการที่สำคัญ ซึ่งในงานวิจัยนี้ได้ใช้ Java Concurrency Package ในการพัฒนาส่วนการของการดึงข้อมูลแบบขนานด้วย

บทที่ 3 การดำเนินการวิจัย

เป้าประสงค์ของการทำงานโครงการย่อยที่ 1 ประกอบด้วย การสำรวจซอฟต์แวร์ที่ใช้ในปัจจุบันของหน่วยงาน แหล่งที่มาของข้อมูลของการท่องเที่ยวเชิงสุขภาพ ขั้นตอนหรือกระบวนการรวบรวมข้อมูล การศึกษางานวิจัยที่เกี่ยวข้องด้านการประมวลผลแบบขนานบนกลุ่มเมฆ พัฒนาอัลกอริธึมแบบขนาน และการแทนค่าข้อมูลสำหรับการประมวลผล การออกแบบหน้าจอ การติดต่อผู้ใช้ การเชื่อมต่อกับข้อมูล การประเมินผล ผลการค้นหาด้านประสิทธิภาพ รวมทั้งประเมินผลการใช้งานกับผู้ใช้ในบทบาทต่างๆ กัน สรุปเป็นขั้นตอนการดำเนินการดังนี้

1. การสำรวจลักษณะซอฟต์แวร์ ข้อมูลที่ใช้อยู่ปัจจุบันของเทศบาล

คณะผู้วิจัยได้ สอบถามข้อมูลซอฟต์แวร์ที่ใช้อยู่ในเทศบาล และการทำงานของหน่วยงาน แหล่งข้อมูลต่างๆ ของการท่องเที่ยวที่ทางเทศบาลใช้อยู่ ด้วยการสัมภาษณ์และแบบสอบถาม โดยรายละเอียดอยู่ในรายงานแผนการวิจัย หัวข้อ 4.1 และ 4.2 และในภาคผนวกของแผนการวิจัย

2. การศึกษางานวิจัยและทฤษฎีที่เกี่ยวข้องของการประมวลผลบนกลุ่มเมฆ และการทำงานแบบขนาน

คณะผู้วิจัยได้ศึกษารายละเอียดต่างๆ ดังในบทที่ 2 และ บทที่ 3

3. การพัฒนาอัลกอริธึมแบบขนาน การแทนค่าข้อมูลในการประมวลผล

คณะผู้วิจัยได้ทำการพัฒนาอัลกอริธึมแบบขนานในส่วนของการประมวลผลข้อมูลเข้า และส่วนของการค้นหา

4. การออกแบบหน้าจอต่อผู้ใช้งาน การเชื่อมต่อข้อมูลต่างๆ

คณะผู้วิจัยได้ทำการออกแบบหน้าจอ การเชื่อมต่อกับฐานข้อมูล เซอร์วิสจากภายนอก รายละเอียดหน้าจออยู่ในภาคผนวก ข. ของแผนงาน คู่มือผู้ใช้งาน

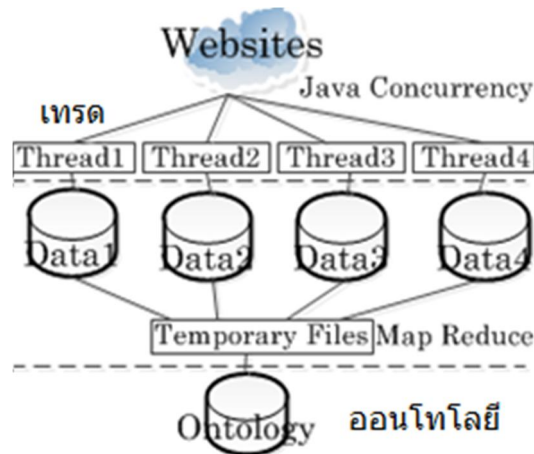
5. การประเมินประสิทธิภาพ

คณะผู้วิจัยได้ทำการประเมินประสิทธิภาพด้านเวลาใน 2 แบบ ได้แก่ เวลาที่ใช้ในการประมวลผลคิวรี และเวลาที่ใช้ในการประมวลผลแบบขนานเมื่อนำข้อมูลตัวอย่างเข้าไปในสถาปัตยกรรมแบบขนาน เช่น หน่วยประมวลผลกราฟิก และทดสอบเวลาการค้นหา

6. การประเมินผู้ใช้งาน

คณะผู้วิจัยได้ทำการประเมินผลจากผู้ใช้งานความเห็นต่อหน้าจอแสดงผลเมื่อระบบพัฒนาเสร็จสิ้นความเห็นต่อการนำเสนอข้อมูลและความเห็นต่อการอบรมผู้ใช้งานในส่วนของเทศบาลตั้งในภาคผนวก ฅ. ของแผนงานวิจัย

3.1 สถาปัตยกรรมโครงสร้างสำหรับการนำข้อมูลเข้าแบบขนาน



ภาพที่ 3.1 สถาปัตยกรรมโครงสร้างการนำข้อมูลเข้าแบบขนาน

สถาปัตยกรรมโครงสร้างการนำข้อมูลเข้าแบบขนานประกอบไปด้วย 3 ส่วน

ส่วนที่ 1: การใช้ Java Concurrency ทำงานแบบหลายเธรดพร้อมกัน เพื่อไปทำการค้นหาข้อมูลจากเว็บไซต์ที่กำหนดภายในเวลาพร้อมกัน จำนวนเธรดถูกกำหนดไว้ในโปรแกรม ข้อมูลที่ได้จากแต่ละเธรดถูกเก็บแยกกัน เพื่อนำไปประมวลผลในขั้นตอนต่อไป

ส่วนที่ 2: การใช้ MapReduce ทำการประมวลผลข้อมูลของแต่ละเธรด และรวบรวมผลลัพธ์ไว้ในแฟ้มข้อมูล

ส่วนที่ 3: การใช้ออนโทโลยีจัดเก็บข้อมูลที่ได้นำเข้ามา

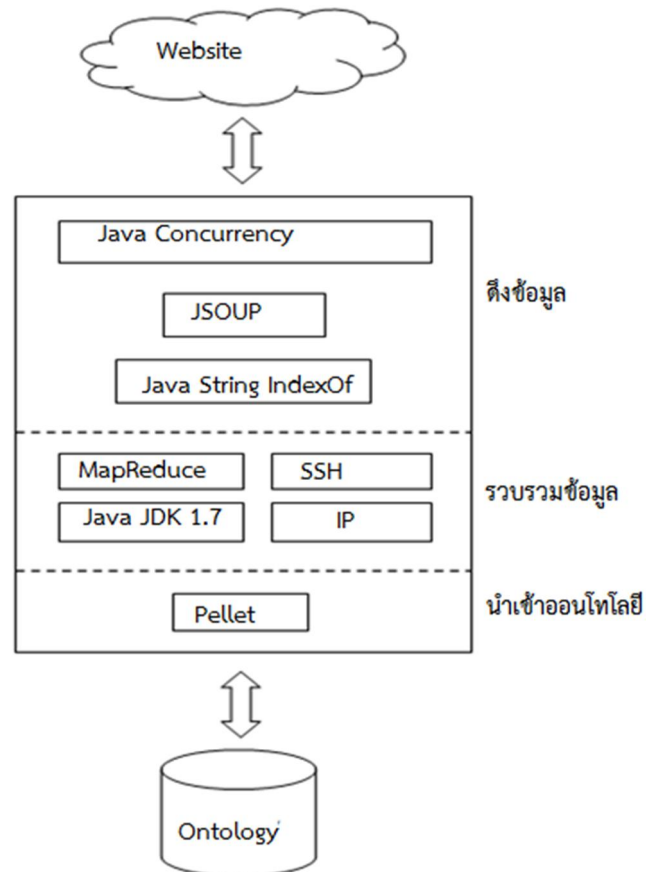
การดึงข้อมูลเข้าจากเว็บไซต์เพื่อจัดลงออนโทโลยีเป็นการทำงานโดยผู้ดูแลระบบ โดยทำการดึงข้อมูลจากเว็บไซต์ เลือกข้อมูลส่วนที่ต้องการ จากนั้นจึงรวมข้อมูลเข้าด้วยกันเพื่อจัดรูปแบบข้อมูล และบันทึกข้อมูลลงตามโครงสร้างออนโทโลยีที่สร้างไว้

ข้อมูลที่ได้รับจากแต่ละเว็บไซต์มีดังในภาคผนวก ข. ของแผนงานวิจัยต่อไปนี้

- ข้อมูลจาก Agoda คือ โรงแรม ที่อยู่ จำนวนห้อง และพิกัดทางภูมิศาสตร์
- ข้อมูลจาก AtSiam คือ โรงแรม ราคา หมายเลขโทรศัพท์ พิกัดภูมิศาสตร์และสิ่งอำนวยความสะดวก
- ข้อมูลจาก Hey Huahin คือ ชื่อร้านหมวดกลุ่มสปาและนวดแผนไทย ที่อยู่และพิกัดทางภูมิศาสตร์
- ข้อมูลจาก TripAdvisor คือโรงแรม ที่อยู่ ราคา ระดับดาว และสิ่งอำนวยความสะดวก

ตัวอย่างสิ่งอำนวยความสะดวก อาทิเช่น สปาและนวดแผนไทย ที่จอตรล ไลฟ์ บริการตามห้อง บาร์ ที่จอตรล ชักริต สระว่ายน้ำ สถานที่ออกกำลังกาย ห้องประชุม ห้องพยาบาล สวน และสินค้าที่ระลึก เป็นต้น

สถาปัตยกรรมซอฟต์แวร์ของการดึงข้อมูลแสดงในภาพที่ 3.2



ภาพที่ 3.2 สถาปัตยกรรมซอฟต์แวร์

ภาพที่ 3.2 แสดงสถาปัตยกรรมซอฟต์แวร์ที่ใช้ในระบบ จะเห็นว่าสถาปัตยกรรมแบ่งออกเป็น 3 ระดับชั้น คือ การค้นหาและการเข้าถึงข้อมูล การรวบรวมข้อมูล และการนำข้อมูลเข้าออนไลน์

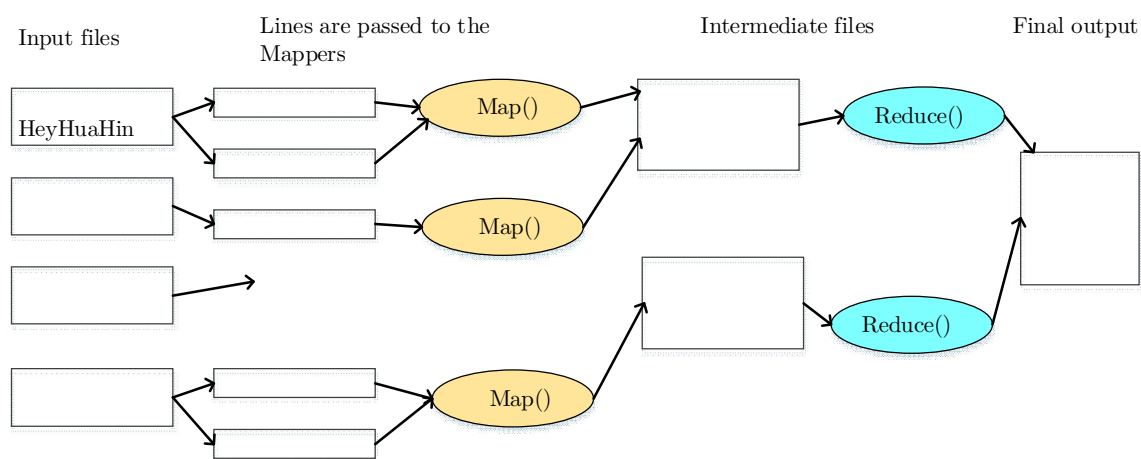
ระดับชั้นที่ 1: การค้นหาและการเข้าถึงข้อมูล

คณะผู้วิจัยใช้ Java Concurrency ช่วยในการเข้าถึงข้อมูลโรงแรมจากเว็บไซต์ที่กำหนด เพื่อช่วยให้ระบบสามารถทำงานด้วยความรวดเร็ว และใช้ Jsoup ช่วยในการดึงข้อความจากหน้าเว็บที่เป็นภาษา HTML จำนวน 3 แหล่งข้อมูล และบันทึกลงแฟ้มข้อมูล (text file) เนื่องจากแต่ละแหล่งข้อมูลอาจจะมีลิงค์ย่อย คณะผู้วิจัยใช้เทรต Java ในการดึงข้อมูลของแต่ละลิงค์ และนำข้อมูลที่ดึงจากแหล่งข้อมูลเดียวกันบันทึกไว้ในแฟ้มข้อมูลเดียวกัน เมื่อได้แฟ้มข้อมูล 3 แฟ้ม จากแต่ละแหล่งข้อมูลเรียบร้อยแล้ว คณะผู้วิจัยได้ใช้คำสั่ง

จาก Java String คำสั่ง indexOf Parsing ช่วยในการค้นหาและตัดเฉพาะข้อความที่ต้องการจากตัวเอกสารที่ดึงมาได้ทั้งหมด

การทำงานในสถาปัตยกรรมระดับขั้นนี้ จะได้ข้อมูลของโรงแรมที่ประกอบไปด้วย ชื่อภาษาไทย ชื่อภาษาอังกฤษ ที่อยู่ รหัสไปรษณีย์ ราคาเริ่มต้น จำนวนดาวโรงแรม จำนวนห้องพัก คะแนนความนิยม หมายเลขโทรศัพท์ ละติจูด ลองจิจูด และสิ่งอำนวยความสะดวกของโรงแรม โดยที่ระบบจะบันทึกสิ่งอำนวยความสะดวกก็ต่อเมื่อโรงแรมมีสิ่งอำนวยความสะดวกนั้นๆ แต่ถ้าหากโรงแรมไม่มีสิ่งอำนวยความสะดวกนั้นๆ ระบบจะบันทึกเป็นช่องว่าง และเพื่อให้่ายในการจัดการข้อมูลจำนวนมาก ระบบจะทำการบันทึกสิ่งอำนวยความสะดวกต่างๆ เรียงตามลำดับ ดังนี้ บริการรถรับ-ส่งสนามบิน บาร์/เลาจ์ บาร์บีคิว ชายหาด บริการรถโดยสาร ร้านกาแฟ บริการรับเลี้ยงเด็ก สวน อินเทอร์เน็ต บริการซักรีด มุมนั่งเล่น นวด ห้องประชุม มินิบาร์ ภูเขา ในท์คลับ/คาราโอเกะ ที่จอดรถ ตู้เย็น ร้านอาหาร รมเซอร์วิส ตู้เซฟ ทะเล บริการจัดเลี้ยง สปา ครูขนาดเล็ก สระว่ายน้ำ โทรศัพท์ และอาหารเช้า ข้อมูลแต่ละส่วนถูกคั่นด้วยจุลภาค (,)

ข้อมูลที่บันทึกในแฟ้มข้อมูลทั้ง 3 แฟ้ม มีการเรียงลำดับข้อมูลเหมือนกันตามที่อธิบายไว้ข้างต้น หากไม่มีข้อมูลจะบันทึกเป็นช่องว่างแทน ทั้งนี้เพื่อความสะดวกในการจัดการข้อมูลจำนวนมาก คณะผู้วิจัยได้เพิ่มสัญลักษณ์และเครื่องหมายทวิภาค (:) ไว้ข้างหน้าข้อความเพื่อระบุข้อความส่วนนี้เป็นอะไร เช่น room: 20 ให้ความหมายได้ว่ามีจำนวนห้อง 20 ห้อง



ภาพที่ 3.3 ขั้นตอน MapReduce

ระดับขั้นที่ 2: การรวบรวมข้อมูล

คณะผู้วิจัยใช้กระบวนการทำงานของ MapReduce ในการปรับรวบรวมข้อมูลที่ดึงจากแหล่งข้อมูล 3 แหล่ง ซึ่งอาจจะมีบางโรงแรมที่ได้ข้อมูลออกมาเหมือนกัน โดยที่ MapReduce แบ่งการทำงานเป็น 2 ส่วน คือ Map จะอ่านข้อมูลจากแฟ้มข้อมูลในขั้นแรก แล้วระบุให้ชื่อโรงแรมภาษาไทยเป็น key (ชื่อภาษาไทย คือ ข้อมูลส่วนแรกซึ่งได้ประโยชน์จากการแบ่งข้อมูลด้วยจุลภาค) เพื่อใช้ในการจำแนกว่ามีโรงแรมใดบ้างที่ใช้ชื่อเดียวกัน ในขั้นตอนนี้จะระบุได้เฉพาะชื่อที่เขียนเหมือนกันทุกตัวอักษรเท่านั้น เช่น คำว่า “ริ

สอร์ท” กับ “รีสอร์ท” การ Map จะมองเป็นคนละชื่อกัน และเพื่อช่วยให้การทำงานง่ายขึ้น จึงมีการตัดการเว้นวรรคภายในชื่อโรงแรมออกไปด้วย เพราะหากมีการเว้นวรรคกับไม่มีการเว้นวรรค Map ก็จะมองเป็นชื่อที่ต่างกันอีกเช่นกัน ส่วนชื่อภาษาอังกฤษและข้อมูลโรงแรมส่วนที่เหลือจะถูกกำหนดให้เป็นค่า value ผลลัพธ์ของการ Map ระบุเป็นคู่ <key, value> และส่งออกไปเป็นอินพุตให้ Reduce จำนวน 1 แพ้มนั้น ไม่ว่าจะมีจำนวนแพ้มข้อมูลเท่าไร การ Map จะรวบรวมข้อมูลออกมาได้เพียง 1 แพ้ม

ในขั้นตอนการ Reduce จะนำค่า value ที่มี key เดียวกันมาต่อกัน คือ หากมีชื่อโรงแรมที่เขียนเหมือนกันทุกประการแล้ว ข้อมูลต่างๆ (ค่า value) จะถูกนำมาต่อกันเท่ากับที่หาได้ทั้งหมดในแพ้มข้อมูล (ผลลัพธ์ของ Map) เมื่อต่อกันของ 1 แพ้มแล้วจะจบด้วยเครื่องหมายอัฒภาค(;) ทั้งนี้หากไม่พบข้อมูลชื่อโรงแรมที่เขียนเหมือนกัน ก็จะมีแต่ข้อมูลจากหน้าเว็บเดียวที่ต่อท้ายด้วยเครื่องหมายอัฒภาค (สามารถรู้ได้ว่ามีข้อมูลจากแหล่งเดียว ไม่มีการต่อกัน) ในขั้นตอนนี้ไม่สามารถบอกได้ว่ามีลำดับการต่อกันของข้อมูลอย่างไร เป็นกระบวนการภายในของ Reduce เอง บางครั้งเราจะพบว่ามี การต่อกันของข้อมูลจาก Atsiam; Tripadvisor; Agoda; หรือ Atsiam; Agoda; Tripadvisor; หรือหากมีข้อมูลเพียงแหล่งเดียวก็จะเป็น Tripadvisor; หรือ Agoda; เพียงอย่างเดียว ผลลัพธ์ของ Reduce ถูกระบุเป็นคู่ <key, value> เช่นเดียวกับการ Map

การทำงานในขั้นตอนนี้ คณะผู้วิจัยได้ทำการลดจำนวนข้อมูลที่ซ้ำซ้อนกันได้ระดับหนึ่งเท่านั้น กล่าวคือ จะได้ข้อมูลของโรงแรมที่ไม่ซ้ำซ้อนกัน เฉพาะกรณีที่มีการเขียนชื่อโรงแรมเดียวกันเหมือนกันเท่านั้น ซึ่งเป็นข้อจำกัดจากการที่ได้ข้อมูลมาตั้งแต่แรกแล้วเพราะแต่ละหน้าเว็บมีการเขียนแตกต่างกัน

ต่อมาจึงทำการเติมข้อมูลที่ไม่มีจากข้อมูลที่มี กล่าวคือ ข้อมูลโรงแรมจากเว็บหนึ่งมีบ้าง ไม่มีบ้าง ทำให้ข้อมูลขาดหายไป คณะผู้วิจัยจึงทำการเติมข้อมูลส่วนที่ไม่มีด้วยข้อมูลที่มี ที่ได้จากการดึงข้อมูลจากหน้าเว็บอื่น เช่น จาก Agoda ไม่มีข้อมูลหมายเลขโทรศัพท์ ก็จะไปหาข้อมูลหมายเลขโทรศัพท์จากข้อมูลที่ต่อท้ายมาเติม (ผลลัพธ์ของ MapReduce) หากไม่มีก็จะถือว่าข้อมูลในส่วนนั้นไม่สามารถระบุได้จริงๆ ในการไปหาข้อมูลมาเติมนั้น ทำการเติมเข้าไปในข้อมูลชุดแรก (แต่ละชุดค้นด้วยเครื่องหมายอัฒภาค) โดยทำเฉพาะส่วนที่ไม่มีอยู่แล้วตั้งแต่แรก (พิจารณาข้อความหลังเครื่องหมายจุลภาค) เมื่อพบเป็นช่องว่าง ก็จะได้ว่าเป็นส่วนที่ขาดจากสัญลักษณ์ใด แล้วจึงไปหาในข้อมูลที่ต่อท้ายที่มีชื่อสัญลักษณ์เดียวกัน หากพบครั้งแรกก็นำมาเติมทันที ไม่ต้องหาในส่วนถัดไป หากไม่พบก็จะค้นหาในลำดับถัดไป (หาเฉพาะที่เป็นชื่อสัญลักษณ์เดียวกัน) ผลลัพธ์จากขั้นตอนนี้จะได้ข้อมูลโรงแรมที่มากขึ้น และมีลำดับเช่นเดียวกับตอนที่นำเข้ามา แต่เหลือเพียงส่วนเดียว ไม่มีเครื่องหมายอัฒภาค ได้ผลลัพธ์ออกมา 1 แพ้มข้อมูล

3.2 การพัฒนาอัลกอริธึมการประมวลผลวิธีแบบขนาน และการแทนค่าข้อมูล

งานวิจัยนี้ทำการประมวลผลวิธีแบบขนานโดยใช้หน่วยประมวลผลกราฟิกส์ที่มีการทำงานแบบพร้อมกัน และใช้หน่วยความจำร่วมกัน จำนวนหน่วยประมวลผลกราฟิกส์ที่ใช้ขึ้นกับลักษณะของการจัดจอ (Graphic card) ที่มีอยู่ การนำข้อมูลเข้าไปให้หน่วยประมวลผลกราฟิกส์ช่วยคำนวณนั้นจะมีขั้นตอนการแปลงข้อมูลก่อนนำเข้า ประกอบด้วยขั้นตอนย่อยดังต่อไปนี้ (ภาพรวมเป็นดังภาพที่ 3.4 (ก))

- ข้อมูลนำเข้า คือ ออนโทโลยีด้านการท่องเที่ยวเชิงสุขภาพ

- ข้อมูลคำถาม คือ คิวรีคำถามจากผู้ใช้งาน
- ข้อมูลผลลัพธ์ คือ คำตอบหลังจากผ่านการคิวรี ส่งจากระบบแก่ผู้ใช้งาน

ภาพที่ 3.4 (ก) แสดงให้เห็นการทำงานของระบบเว็บเชิงความหมายแบบเดิมที่สามารถคิวรีระหว่างโปรแกรมเว็บและ RDF ที่เป็นโครงสร้าง XML ผ่านมิดเดิลแวร์ดังภาพที่ 3.4 (ข) และแปลงเป็นเว็บแอปพลิเคชันสำหรับผู้ใช้ได้เลย

ภาพที่ 3.4 (ค) และ (ง) โปรแกรมเบื้องหลังที่ได้รับการปรับปรุง จะทำงานระหว่างหน่วยประมวลผลกลาง หน่วยประมวลผลกราฟิกส์และชุดข้อมูลเพื่อลดขนาดของข้อมูลให้สามารถนำเข้าไปประมวลผลด้วยหน่วยประมวลผลกราฟิกส์ได้ และเร่งความเร็วด้วยการค้นหาแบบขนานในการคิวรี

ในภาพที่ 3.4 (ง) มีรายละเอียดดังนี้

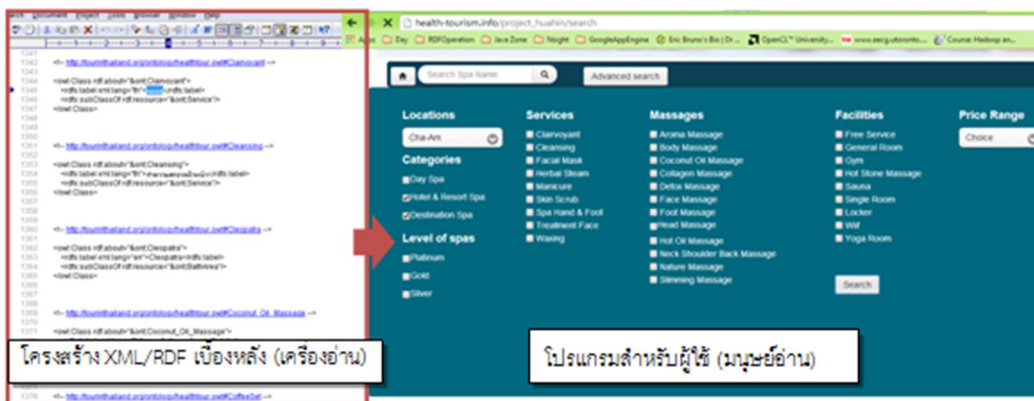
ขั้นตอนที่ 1 นำข้อมูลออนโทโลยีเข้าไปโดยแปลงเป็นในโครงสร้าง RDF และแปลงสู่รูปแบบย่อแบบไบนารีก่อน

ขั้นตอนที่ 2 นำส่วนเนื้อหามาแปลงเป็นเลขสำหรับดัชนีแบบแบ่งเป็นชุดๆ จากนั้นแปลงเป็นลักษณะดัชนี 1-0 เพื่อเข้าสู่หน่วยความจำของประมวลผลกราฟิกส์ เพื่อประมวลผล โดยนำข้อมูลคิวรีที่ต้องการค้นหาที่รับจากหน้าจอผู้เข้ามาแปลงเป็นลักษณะเดียวกันเพื่อค้นหาในหน่วยความจำดังกล่าว

ขั้นตอนที่ 3 นำข้อมูลที่ค้นหาพบออกมาจากหน่วยความจำของหน่วยประมวลผลกราฟิก แล้วแปลงกลับด้วยสคริปต์ต่าง ๆ เพื่อให้ผู้ใช้ระบบสามารถอ่านได้ผ่านเว็บเบราว์เซอร์



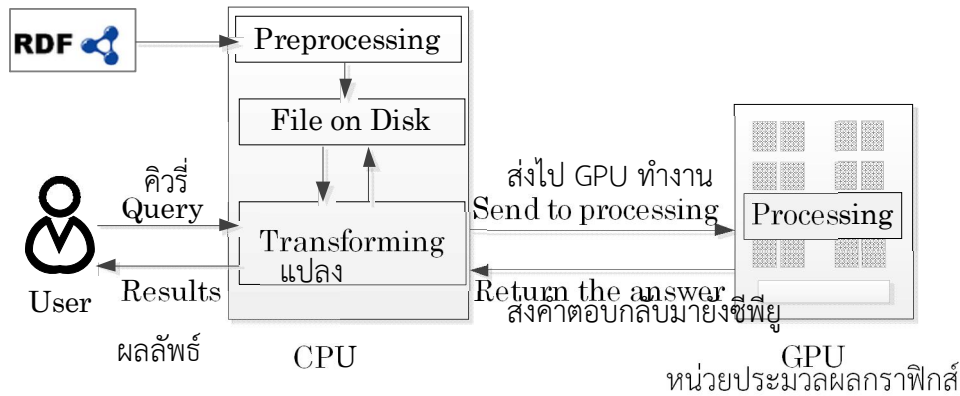
(ก) ขั้นตอนโดยรวม



(ข) จากหน้าเว็บด้านขวาจะสามารถแปลงเป็น RDF ด้านซ้าย



(ค) ขั้นตอนหลังจากนำ RDF มาเปลี่ยนเป็นชุดตัวเลขไบนารี



(ง) ภาพรวมเมื่อทำงานร่วมกับ GPU

ภาพที่ 3.4 ภาพรวมการแปลงข้อมูลออนโทโลยีเป็นข้อมูลเข้าประมวลผลคิวรีแบบขนาน

จากภาพสรุปได้ว่า ระบบทำการประมวลผลการทำงานบน 2 ส่วน ได้แก่ส่วนที่อยู่บนหน่วยประมวลผลกลาง และส่วนที่อยู่บนหน่วยประมวลผลกราฟิกส์ โดยเริ่มจากผู้ใช้ป้อนคิวรีเข้าไป ซึ่งจะถูกแปลงภายในฝั่งซีพียู และในขณะเดียวกันข้อมูลออนโทโลยีจะถูกแปลงเป็นลักษณะ RDF เช่นกัน จากนั้นคิวรีที่แปลงแล้วจะถูกส่งไปยังหน่วยประมวลผลกราฟิกส์เพื่อค้นหา และเมื่อพบคำตอบก็จะนำคำตอบออกจากหน่วยประมวลผลกราฟิกส์และมาแปลงกลับเพื่อแสดงผลต่อไป

3.2.1 การแปลง RDF เป็นชุดตัวเลขไบนารี

การจัด RDF ด้านการท่องเที่ยงเชิงรูปภาพเป็นชุดข้อมูลที่แต่ละบรรทัดประกอบด้วย 3 ตัวหรือเรียกว่า Triple โดยเรียงจาก

Subject Predicate Object

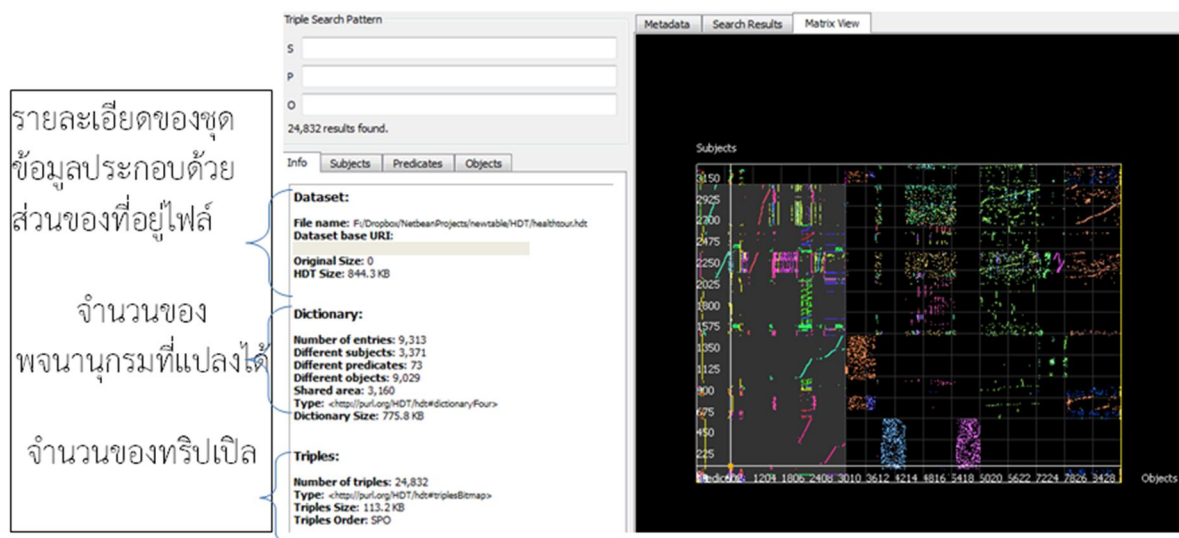
ซึ่งหมายถึงประธาน กริยา และกรรม ตามลำดับ ตัวอย่างเช่น โรงแรม ฮิลตัน มีสิ่งอำนวยความสะดวกห้องพักปรับอากาศ โดยประธานคือ โรงแรม ฮิลตัน กริยาคือ มีสิ่งอำนวยความสะดวก และกรรมคือห้องพักปรับอากาศ

ในโครงสร้าง HDT (Header-Dictionary-Triple) [33] มีการจัดเก็บข้อมูลบรรทัดดังกล่าวเป็นตัวเลข โดยมี unique ID สำหรับประธาน กริยา และกรรมตามลำดับ ถ้ามีหลาย Triple ก็จะมีหลายบรรทัดนั่นเอง ในกรณีทั่วไปประธานหนึ่งตัวอาจจะมีหลายๆ กริยา และกริยาหนึ่งตัวอาจจะมีหลายกรรม ดังนั้นจึงสามารถจัดกลุ่มโดยจัดความสัมพันธ์แบบต้นไม้ม โดยทำการจัดกลุ่มของประธานเข้าด้วยกัน สำหรับแต่ละประธานก็จัดกลุ่มกริยาเข้าด้วยกัน และจัดเก็บเป็นลิสต์ ดังนี้

Subject (กลุ่ม Predicate ของ Subject นั้น (กลุ่ม Object ของ Predicate และ Subject นั้น))

จากนั้นในงานวิจัยนี้ ได้นำเสนอการแทนค่าชุดข้อมูลข้างต้นด้วย ตัวเลข ID ของประธาน กริยา กรรมตามตำแหน่ง และใช้อะเรย์ของบิต 0-1 หรือไบนารีอีกชุดเพื่อกำหนดตำแหน่งเริ่มต้นใหม่ของประธาน กริยา กรรมเป็นต้น

ดังนั้นทั้งประโยคควีรี และชุดข้อมูลหลัก จะต้องถูกแปลงเป็นข้อมูลในรูปแบบเดียวกันนี้เพื่อค้นหา จากนั้นใช้อัลกอริธึมการค้นหาแบบใดก็ได้ที่รันบนหน่วยประมวลผลกราฟิกส์ช่วยค้นหาแบบขนาน ซึ่งในที่นี้ได้ทดสอบกับการค้นหาแบบบรูซฟอร์ซ (Brute force) นั่นคือการเปรียบเทียบเลข ID โดยดูตำแหน่งจากเลขไบนารีของชุดข้อมูลหลัก และเปรียบเทียบเลข ID ที่สัมพันธ์กันกับเลข ID ของประโยคที่ต้องการควีรี ภาพที่ 3.5 แสดงการแปลงจากข้อมูล RDF เป็น HDT



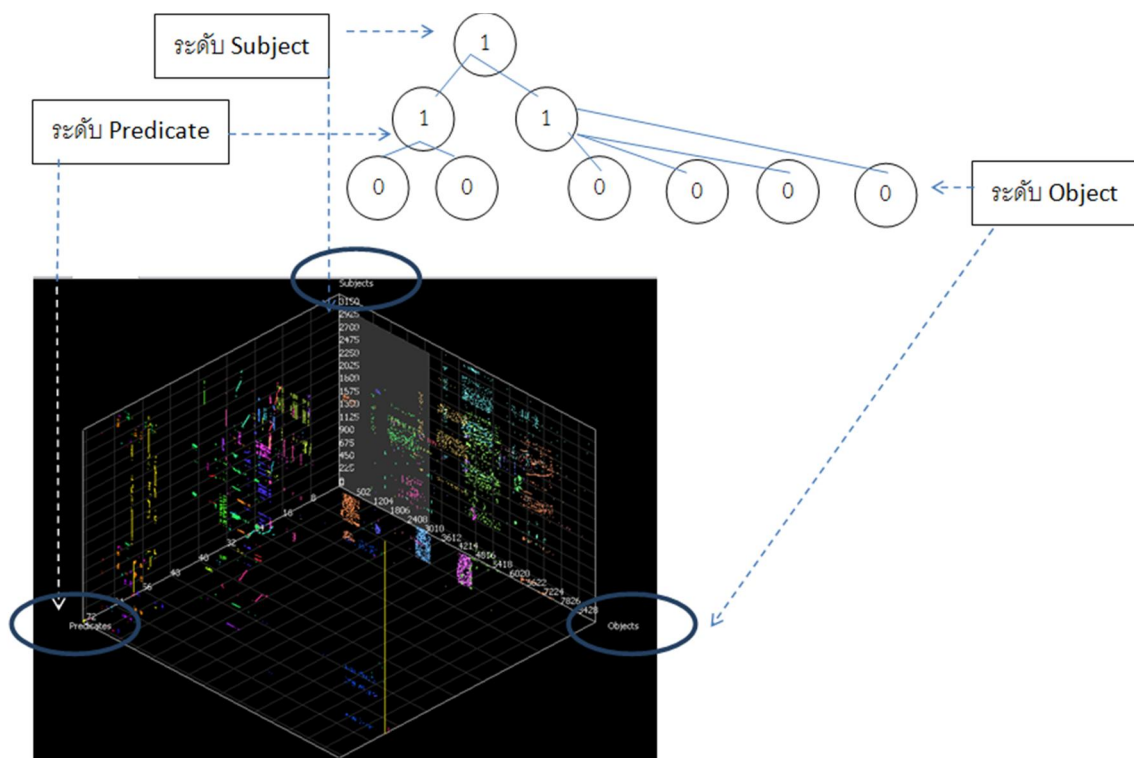
ภาพที่ 3.5 ข้อมูลที่ได้จากการแปลงแฟมออนโทโลยี RDF เป็น HDT

ขั้นตอนการแปลงชุดข้อมูลซึ่งเป็นการเตรียมชุดข้อมูลนั้น เป็นการแปลงชุดข้อมูลให้อยู่ในรูปแบบต้นไม้ม ในการจัดเก็บแบบลิสต์ และอะเรย์ 1-0 (รวมสองอย่างนี้ เรียกว่า 1-0 Tree) เพื่อจะได้นำผลลัพธ์ที่ได้เข้าสู่กระบวนการค้นหาต่อประธาน 1 ตัวต่อไป การทำงานประกอบไปด้วยขั้นตอนต่อไปนี้

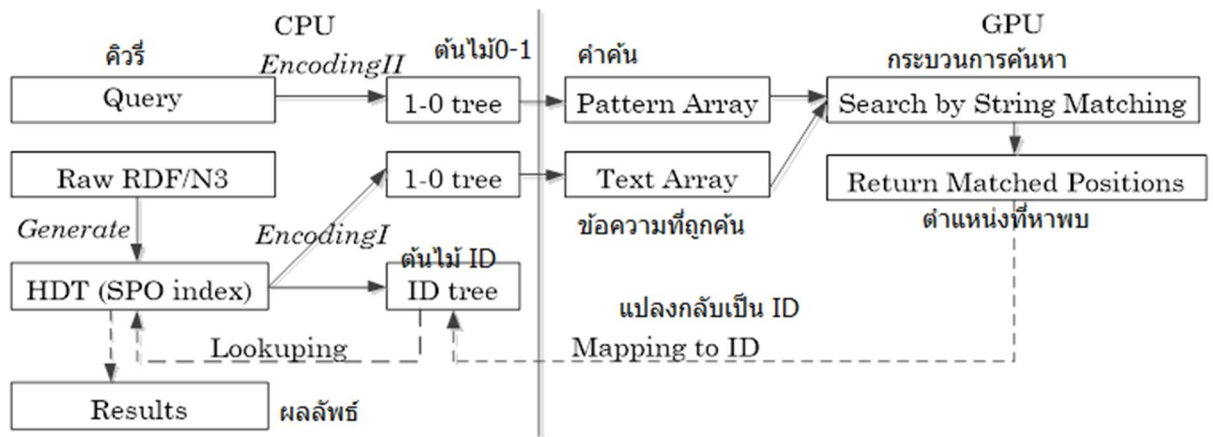
- 1) ตำแหน่งของชุดข้อมูลที่เป็นประธานให้แปลงเป็น 1
- 2) ตำแหน่งของชุดข้อมูลที่เป็นกริยาให้แปลงเป็น 1
- 3) ตำแหน่งของชุดข้อมูลที่เป็นกรรมให้แปลงเป็น 0

จากภาพ 3.6 สามารถเทียบขั้นตอนการแปลงเป็นต้นไม้ของโครงงาน (บน) ได้กับโครงสร้างเดิมของ HDT (ล่าง) ได้ คือ ระดับรากของต้นไม้คือ Subject ชุดเลข 1 สามารถเทียบได้กับแกนของ ของ HDT ในขณะที่ชุดเลข 1 ถัดมาเทียบได้กับแกนของ Predicate ของ HDT และชุดเลข 0 เทียบได้กับแกน Object ของ HDT

โครงสร้าง HDT จะมีโครงสร้าง 3 ส่วนได้ ได้แก่ Header, Dictionary, และ Triples ซึ่ง Triple นี้ อยู่ในรูป S, P,O (Subject, Predicate, Object) ที่จะถูกนำมาใช้สร้างต้นไม้ 1 และ 0 ดังภาพ 3.6 โดยกำหนดให้ Subject S เป็นรากของต้นไม้ จากนั้นระดับถัดไปเป็น Predicate P และระดับต่อไปเป็น Object O ซึ่งจะมีต้นไม้หลายต้นก็ได้ ต้นไม้นี้จะถูกมองเป็นสายอักขระของ ID ที่สัมพันธ์กับสายอักขระของเลข 0-1 โดย ID จะหมายถึงรหัสตัวแทนของแต่ละ Subject Predicate และ Object



ภาพที่ 3.6 ชุดข้อมูลที่ได้หลังจากแปลงตัวเลข SPO ของ HDT



ภาพที่ 3.7 รายละเอียดขั้นตอนการแปลงและการค้นหา

ภาพที่ 3.7 แสดงขั้นตอนการแปลงและการค้นหา โดยเพิ่มรายละเอียดการแปลงข้อมูล ที่เป็นการเปลี่ยนรูปแบบการแทนค่าข้อมูลเพื่อให้สามารถนำเข้าสู่หน่วยประมวลผลกราฟิกส์เพื่อค้นหาได้เร็ว โดยวิธีจะถูกแปลงเป็น RDF ซึ่งจากนั้นจะถูกแปลงให้อยู่รูปของ HDT และแปลงเป็นสายอักขระของ ID และแปลงเป็นรูปแบบต้นไม้ 1 และ 0 (ดังภาพที่ 3.8 และภาพที่ 3.9) และนำเข้าไปในหน่วยความจำของหน่วยประมวลผลกราฟิกส์เพื่อการค้นหาต่อไป ดังนั้นการค้นหาจะเป็นการค้นหาแบบการค้นสายอักขระ 0-1 ซึ่งทำให้การค้นหาได้เร็วและจะใช้อัลกอริธึมค้นหาแบบใดก็ได้มาทำงาน

ไฟล์ชุดตัวเลขแบบ
ต้นไม้ ทำขึ้นเพื่อ
เป็นตัวแปลง
ระหว่างไฟล์ชุด 1-0
และ RDF

ภาพที่ 3.8 เพิ่มข้อมูลชุดตัวเลขของ Subject (Predicate(Object))

3.2.2 การค้นหาแบบขนานด้วยวิธีบรัชฟอร์ซของหน่วยประมวลผลกราฟิกส์

วิธีการค้นหาแบบขนานด้วยวิธีบุรุษพอร์ชของหน่วยประมวลผลกราฟิกส์นั้นแตกต่างจากการค้นหาผ่านหน่วยประมวลผลกลางที่เป็นการค้นหาแบบลำดับ โดยการค้นหาแบบลำดับเป็นการค้นโดยเลื่อนตำแหน่งระหว่างชุดข้อมูล กับประโยคควิรีทีละลำดับ และสามารถค้นหาแบบต่อเนื่องจากหัวชุดข้อมูลไปยังท้ายชุดข้อมูลโดยใช้เทรดแบบเดียว

การค้นหภายในเป็นแบบระบุตำแหน่งคงที่ สามารถแบ่งตามลำดับของชุดการแปลงแบบ
Subject, Predicate, Object ได้โดยขั้นตอนต่อไปนี

- 1) การค้นหา Subject มี 1 แบบ
คือเลข 11 หมายถึงค้นหาตำแหน่งของเลข 1 ที่ตามด้วยเลข 1 ที่เป็นตำแหน่งของ Predicate
- 2) การค้นหา Predicate มี 1 แบบ
คือเลข 10 หมายถึงค้นหาตำแหน่งของเลข 1 ที่ตามด้วยเลข 0 ที่เป็นตำแหน่งของ Object

3) การค้นหา Object มี 3 แบบ

คือ การค้นหาเลข 00 คือลำดับของ Object ที่ตามด้วย Object ใน Predicate เดียวกันและการค้น แบบ 01 คือการค้นหา Object ที่ตามด้วย Predicate ใน Predicate เดียวกัน หรือการค้นหา Object ที่ตามด้วย Subject ชุดใหม่ และสุดท้ายเป็นการค้นหาค่า Object ที่เพิ่มข้อมูลตามด้วยเครื่องหมายจบเพิ่มข้อมูล

ในการคืนค่าคำตอบจากการคิวรีที่ได้จากหน่วยประมวลผลกราฟิกส์ จะได้ชุดของตำแหน่งตามลำดับของ Subject, Predicate และกลุ่มของ Object ในที่นี้ยกตัวอย่างการค้นหาทั้งหมดได้ดังต่อไปนี้ Subject, Predicate และกลุ่มของ Object ในที่นี้ยกตัวอย่างการค้นหาทั้งหมดได้ดังภาพที่ 3.10-ภาพที่ 3.13 โดยที่การค้นหาหน่วยประมวลผลกราฟิกส์เกิดขึ้น เมื่อได้รับประโยคคิวรีจากผู้ใช้งานผ่านหน้าจอและมีการแปลงประโยคสั้นนี้แบบทันที เทียบกับชุดข้อมูลที่มีการแปลงไว้แล้ว จากนั้นทำการคัดลอกจากโฮสต์ไปยังดีไวส์ คือหน่วยประมวลผลกราฟิกส์ ที่เป็นการค้นหาแบบขนานบน

ข้อมูลที่ได้อีกกลับมาจะเป็นสายอักขระของ ID ตามตำแหน่งพอดค่า ‘11’ (ภาพ 3.10) ซึ่งเป็นตำแหน่ง Subject จากนั้นจะนำเอา Subject นี้มาเทียบกับ Subject ที่ผู้ใช้ต้องการหา จากนั้นจะไปหา Predicate ที่สัมพันธ์กับ Subject ID นี้ จากนั้นจะไปหา Object ที่สัมพันธ์กับ Predicate ที่ผู้ใช้ต้องการส่งคืนลิสต์ของ Object ID ที่ต้องการคืนไป เป็นต้นส่งคืนกลับผู้ใช้ และทำการแปลงข้อมูลกลับเป็นคำตอบเพื่อแสดงผล

	0	10	20	30	40	50	60	70	80	90	100	110	120	130	140															
1	0	9	18	370	407	439	796	837	1015	1359	1467	1509	1629	1670	2009	2296	2512	2874	2969	3323	3520	3552	3586	3626	3731	3880	3985	4038	4081	4113
	4414	4775	5126	5233	5258	5392	5738	5884	6248	6269	6625	6652	6986	7322	7461	7521	7621	7638	7957	8352	8618	8698	8715	8771	8833	8891	9098	9138		
	9349	9349	9483	9608	9752	9874	9980	10125	10232	10330	10376	10428	10508	10541	10605	10689	10791	10884	10916	11068	11158	11207	11253	11328						
	11498	11553	11601	11791	11827	11835	11916	12052	12106	12155	12353	12456	12794	12867	13059	13178	13387	13466	13646	13850	13952	14058	14138							
	14330	14407	14545	14627	14685	14718	14812	14880	14952	15033	15094	15217	15297	15352	15437	15522	15612	15724	15929	16010	16101	16217	16257							
	16354	16433	16492	16645	16715	16835	16940	17077	17210	17318	17413	17511	17586	17593	17600	17607	17614	17621	17628	17658	17669	17676	17683							
	17690	17697	17708	17719	17726	17814	17884	17949	17956	17963	17970	17977	17984	17991	18060	18126	18226	18233	18261	18325	18332	18339	18346							
	18417	18516	18630	18689	18767	18872	18941	18973	19065	19144	19240	19346	19456	19539	19625	19731	19812	19876	19971	20064	20103	20148	20204							
	20344	20447	20514	20551	20639	20718	20769	20826	20839	21030	21210	21401	21468	21529	21635	21774	21783	21792	21805	21814	21827	21836	21845							
	21854	21867	21876	21885	22009	22018	22027	22036	22045	22054	22063	22072	22104	22117	22126	22135	22148	22157	22166	22171	22184	22235	22244							
	22253	22262	22271	22280	22289	22302	22315	22324	22333	22342	22355	22364	22377	22386	22395	22404	22413	22422	22435	22444	22457	22466	22475							
	22484	22493	22502	22511	22520	22529	22542	22551	22557	22566	22575	22584	22597	22611	22620	22629	22638	22651	22660	22707	22716	22725	22734							
	22743	22752	22761	22770	22779	22788	22797	22810	22819	22828	22837	22846	22855	22864	22873	22882	22891	22904	22945	22954	22963	22976	22985							
	22994	23048	23057	23066	23075	23084	23097	23106	23111	23120	23129	23138	23147	23160	23173	23186	23195	23239	23248	23257	23270	23319	23332							
	23341	23350	23363	23376	23385	23448	23457	23466	23475	23484	23497	23510	23519	23532	23541	23592	23601	23610	23623	23681	23690	23699	23708							
	23717	23726	23735	23740	23749	23758	23767	23800	23809	23818	23827	23836	23845	23854	23863	23872	23881	23894	23903	23916	23925	23934	23947							
	24127	24136	24149	24202	24211	24224	24237	24246	24259	24300	24309	24322	24331	24344	24388	24397	24406	24415	24424	24581	24590	24599	24608							
	24617	24630	24673	24682	24799	24808	24817	24826	24835	24844	24853	24866	24875	24888	24901	24910	24919	24932	24945	24954	24959	25018	25027							
	25071	25080	25093	25137	25146	25155	25164	25173	25182	25191	25200	25213	25222	25267	25276	25285	25298	25307	25316	25325	25334	25356								
	25369	25423	25432	25441	25454	25461	25470	25479	25493	25502	25515	25566	25575	25584	25593	25602	25615	25624	25637	25650	25659	25668	25677							
	25686	25695	25704	25713	25726	25735	25744	25753	25766	25779	25788	25801	25810	25819	25832	25841	25854	25863	25876	25885	25898	25911	25920							
	25969	25982	25991	26035	26044	26090	26099	26131	26140	26149	26158	26167	26181	26195	26204	26213	26226	26235	26286	26295	26304	26313	26322							
	26335	26344	26353	26362	26371	26384	26429	26438	26447	26460	26469	26478	26487	26496	26505	26514	26523	26532	26541	26550	26555	26564	26573							
	26582	26595	26642	26651	26664	26673	26721	26730	26739	26752	26761	26774	26783	26792	26805	26814	26819	26828	26837	26846	26855	26868	26877							
	26932	26941	26955	26964	26973	26986	27033	27046	27094	27103	27116	27125	27134	27143	27152	27213	27222	27231	27240	27249	27262	27271	27280							
	27293	27306	27315	27328	27337	27346	27355	27361	27370	27376	27385	27398	27446	27455	27464	27477	27490	27503	27512	27521	27530	27543	27552							
	27601	27610	27619	27678	27687	27700	27713	27763	27772	27785	27823	27832	27841	27850	27859	27864	27873	27882	27891	27904	27913	27922	27935							
	27944	27953	27962	27971	27984	27993	28025	28034	28047	28060	28069	28078	28091	28104	28117	28186	28195	28208	28217	28230	28280	28289	28298							
	28307	28320	28329	28342	28351	28396	28405	28418	28431	28440	28453	28462	28471	28480	28524	28533	28546	28555	28564	28577	28586	28595	28604							
	28617	28726	28735	28744	28757	28766	28779	28826	28839	28848	28857	28866	28879	28888	28901	28914	28951	28960	28969	28978	28987	28996	29005							
	29014	29023	29032	29041	29050	29059	29068	29111	29124	29133	29142	29151	29164	29173	29182	29195	29208	29221	29230	29239	29248	29257	29266							
	29275	29288	29323	29332	29341	29350	29359	29368	29377	29390	29399	29412	29464	29473	29482	29495	29504	29517	29526	29535	29544	29610	29619							
	29628	29641	29650	29693	29702	29715	29724	29737	29746	29791	29800	29813	29822	29831	29844	29853	29862	29871	29880	29889	29898	29907	29916							
	29925	29934	29943	29952	29961	29970	29983	29992	30001	30014	30023	30036	30045	30054	30063	30072	30081	30090	30099	30108	30121	30130	30254							
	30263	30276	30285	30298	30307	30320	30329	30342	30351	30360	30369	30378	30391	30404	30413	30426	30435	30444	30457	30466	30479	30488	30532							
	30541	30555	30564	30573	30659	30668	30681	30724	30733	30742	30751	30760	30769	30782	30827	30836	30845	30854	30863	30872	30881	30894	30903							

Line 1, Col 1

Val: 48 30h 00110000b: 35c: 14490

ANSI/DOS: Tab4 LIT W IN

ภาพที่ 3.10 การคืนค่าของ Subject ID ตามจำหน่งที่เจอ “11” ที่สามารถตรวจสอบได้จากตำแหน่งแรก

[illegible]

ภาพที่ 3.12 การค้นหาของ Object ID ตาม “01” ที่สามารถตรวจสอบได้จากตำแหน่งแรกที่เป็นลำดับที่สาม

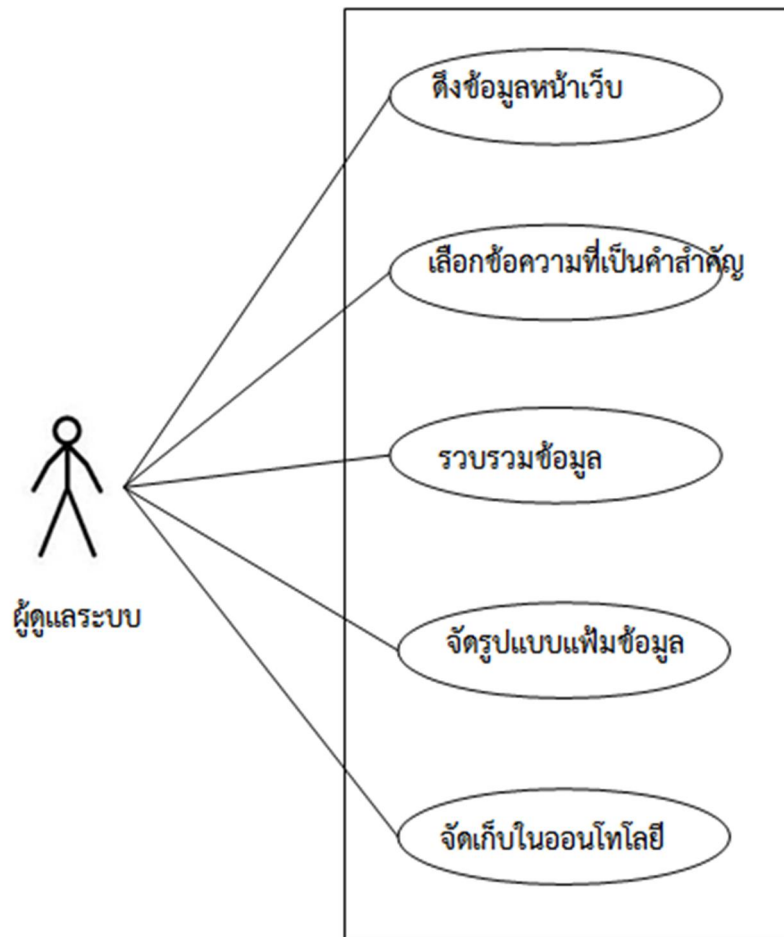
6	10	20	30	40	50	60	70	80	90	100	110	120	130	140																				
39027 39028 39029 39030 39031 39032 39033 39034 39035 39046 39047 39048 39049 39050 39051 39052 39053 39054 39055 39058 39059 39060 39061	39062 39063 39066 39067 39068 39069 39070 39075 39076 39077 39078 39079 39080 39081 39082 39083 39084 39085 39086 39087 39088 39089 39090	39091 39092 39093 39094 39095 39096 39097 39098 39099 39100 39101 39102 39103 39104 39105 39106 39107 39108 39109 39110 39111 39112 39113	39114 39115 39116 39117 39118 39119 39120 39121 39122 39123 39124 39125 39126 39127 39128 39129 39130 39131 39132 39133 39134 39135 39136	39137 39140 39141 39144 39145 39146 39147 39148 39149 39150 39151 39152 39153 39154 39155 39156 39157 39158 39159 39160 39161 39162 39163	39164 39165 39166 39167 39168 39171 39172 39173 39174 39175 39176 39177 39178 39179 39184 39185 39186 39187 39188 39189 39194 39341 39344	39359 39360 39363 39364 39365 39366 39367 39368 39369 39370 39373 39376 39381 39483 39606 39615 39616 39617 39620 39621 39622 39623 39624	39629 39634 39635 39643 39650 39651 39652 39653 39654 39665 39666 39667 39668 39669 39670 39671 39672 39673 39674 39675 39676 39677 39678	39679 39682 39685 39690 39691 39692 39704 39705 39706 39707 39716 39717 39720 39721 39722 39723 39724 39725 39726 39727 39728 39729 39730	39731 39732 39733 39734 39735 39736 39737 39738 39739 39742 39743 39744 39747 39752 39766 39767 39770 39773 39774 39775 39776 39777 39784	39785 39786 39787 39788 39789 39790 39791 39792 39793 39794 39795 39796 39797 39798 39799 39800 39801 39802 39803 39804 39805 39806 39807	39808 39809 39810 39813 39816 39817 39818 39819 39820 39821 39822 39823 39828 39829 39834 39835 39991 40080 40081 40082 40083 40084 40089	40094 40095 40165 40176 40177 40178 40179 40184 40189 40213 40214 40215 40216 40217 40222 40227 40228 40269 40270 40271 40272 40273 40274	40283 40286 40287 40288 40289 40290 40291 40292 40293 40294 40295 40296 40301 40306 40307 40308 40309 40334 40343 40344 40345 40346 40347	40348 40351 40352 40353 40354 40355 40356 40357 40358 40359 40360 40361 40362 40363 40364 40365 40366 40367 40368 40369 40370 40371 40372	40373 40374 40377 40380 40385 40386 40387 40388 40389 40390 40391 40392 40447 40460 40463 40464 40465 40466 40467 40468 40469 40474 40479	40493 40504 40505 40506 40507 40508 40509 40510 40511 40512 40513 40514 40515 40520 40525 40526 40527 40528 40529 40530 40556 40557 40558	40559 40560 40561 40562 40563 40568 40573 40574 40828 40998 40999 41000 41001 41012 41013 41014 41015 41016 41017 41018 41019 41024 41029	41030 41051 41052 41053 41054 41055 41056 41065 41066 41067 41070 41071 41072 41073 41074 41075 41076 41077 41082 41087 41112 41123 41124	41125 41126 41129 41132 41133 41138 41152 41153 41164 41165 41166 41167 41168 41171 41174 41179 41180 41214 41215 41224 41225 41226 41229	41230 41231 41232 41233 41234 41235 41236 41237 41238 41241 41244 41249 41250 41251 41263 41264 41273 41274 41275 41278 41279 41280 41281	41282 41283 41284 41285 41286 41287 41292 41297 41320 41321 41330 41331 41332 41333 41334 41335 41338 41339 41340 41341 41342 41343 41344	41345 41346 41347 41348 41349 41354 41359 41360 41361 41373 41382 41383 41384 41385 41388 41389 41390 41391 41392 41393 41394 41395 41396	41397 41402 41407 41421 41422 41431 41432 41433 41436 41437 41438 41439 41440 41441 41442 41443 41444 41445 41450 41455 41467 41476 41477	41480 41481 41482 41483 41484 41485 41486 41487 41488 41489 41490 41491 41492 41493 41494 41495 41500 41505 41517 41518 41519 41528 41531	41532 41533 41534 41535 41536 41537 41538 41539 41540 41541 41542 41543 41544 41545 41550 41555 41556 41557 41750 41753 41762 41763 41764	41767 41768 41769 41770 41771 41772 41773 41774 41775 41776 41777 41778 41779 41780 41781 41782 41783 41784 41787 41788 41789 41792 41797	41798 41799 41813 41814 41815 41816 41817 41818 41819 41820 41821 41822 41825 41826 41835 41836 41837 41838 41839 41842 41843 41844 41845	41846 41847 41848 41849 41850 41851 41852 41853 41854 41855 41856 41857 41858 41859 41860 41861 41864 41865 41868 41873 41874 41875 41876	41877 41893 41894 41895 41896 41897 41898 41899 41900 41901 41902 41903 41906 41907 41916 41917 41918 41919 41922 41923 41924 41925 41926	41927 41928 41929 41930 41931 41932 41933 41934 41935 41936 41937 41938 41939 41942 41943 41944 41947 41952 41953 41954 41968 41969 41970	41971 41974 41975 41984 41985 41988 41989 41990 41991 41992 41993 41994 41995 41996 41997 41998 41999 42000 42001 42002 42003 42004 42005	42006 42007 42008 42009 42010 42011 42014 42015 42016 42017 42018 42019 42020 42023 42028 42029 42030 42042 42043 42044 42055 42060 42061	42062 42063 42064 42065 42066 42067 42068 42069 42070 42071 42072 42073 42074 42075 42076 42079 42082 42083 42088 42089 42090 42091 42092	42136 42139 42140 42141 42142 42143 42144 42145 42146 42147 42148 42149 42150 42151 42152 42153 42154 42155 42156 42157 42158 42159 42160

Line1, Col1 Val: 50 32h 00110010b Size: 123942 ANSI(DOS) Tab4 LIT W IN

ภาพที่ 3.13 การคืนค่าของ Object ID ตาม “00” ที่สามารถตรวจสอบได้จากตำแหน่งสุดท้าย และสามารถคำนวณค่าสุดท้ายของ Object หลังสุดได้

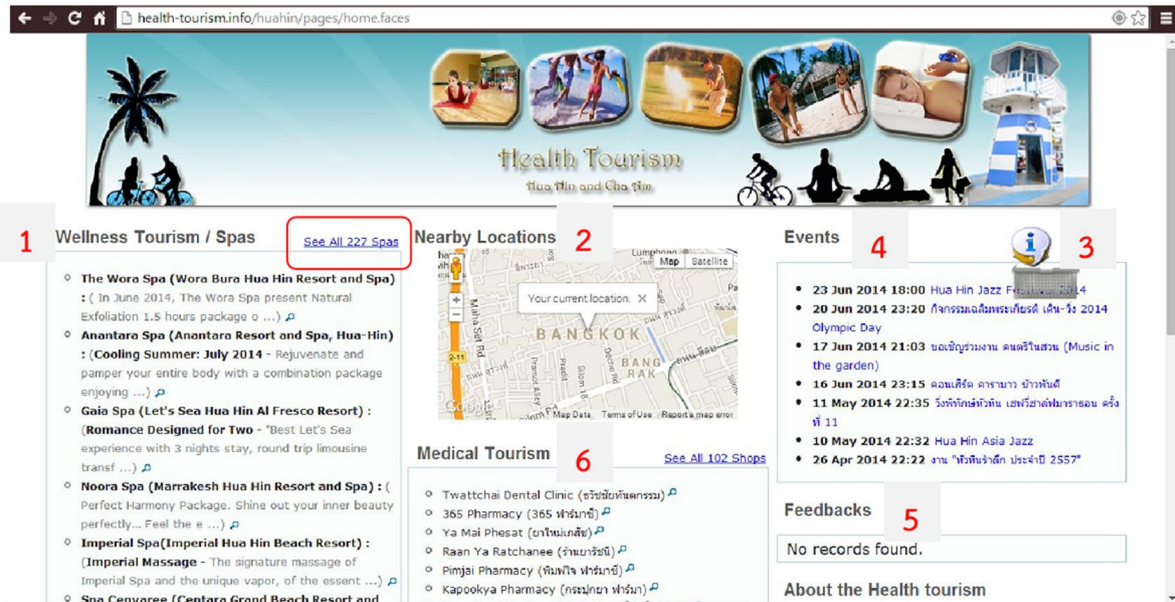
ในขั้นสุดท้ายการแปลงค่าคืนเพื่อส่งคำตอบให้ผู้ใช้สามารถทำได้โดยการเทียบตำแหน่งที่พบกับชุดตัวเลขในหน่วยประมวลผลกลาง แล้วจึงจับคู่กับลำดับตาม HDT เพื่อหาข้อความของ Triple ที่ต้องการออกมา แล้วจึงแปลงมาเป็นคำตอบบนเว็บแอปพลิเคชัน

3.3 Use Case Diagram และหน้าจอผู้ใช้งาน



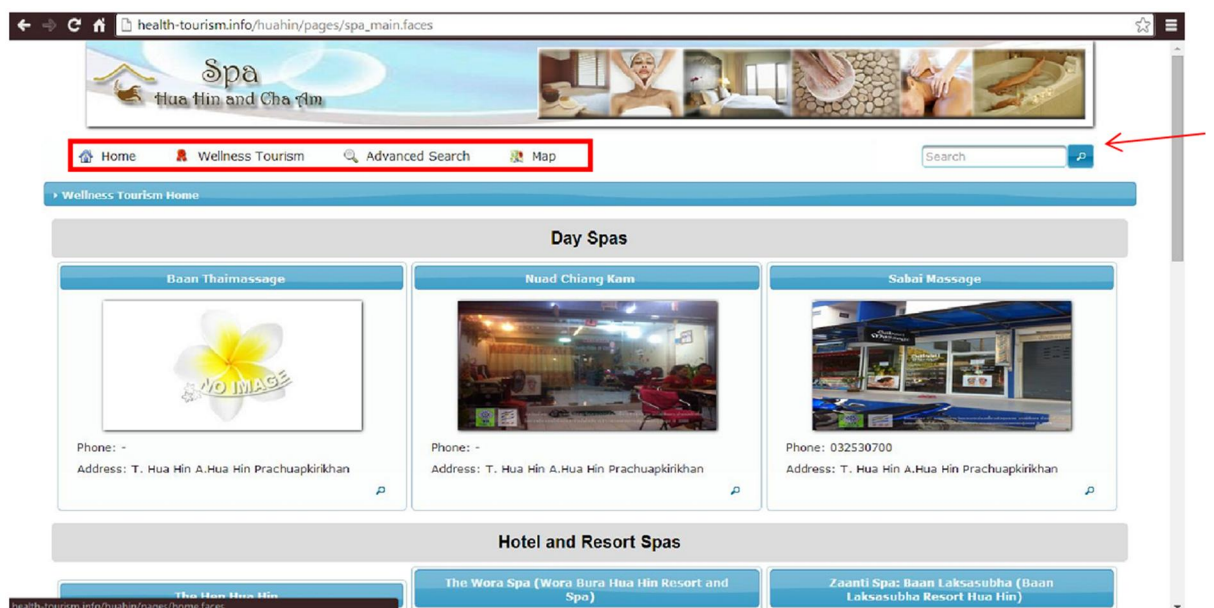
ภาพที่ 3.14 หน้าทีของผู้ดูแลระบบ

ภาพที่ 3.14 แสดงหน้าที่ของผู้ดูแลระบบที่ทำหน้าที่นำข้อมูลจากเว็บไซต์มาเพิ่มในออนไลน์ โดยที่ผู้ดูแลระบบมีการทำงานในแต่ละขั้นตอนเพื่อทำการเตรียมฐานความรู้เพื่อให้ผู้ใช้ระบบที่เป็นนักท่องเที่ยวทำการค้นหาข้อมูลเพื่อประโยชน์ในการประชาสัมพันธ์ด้านการท่องเที่ยวเชิงสุขภาพ



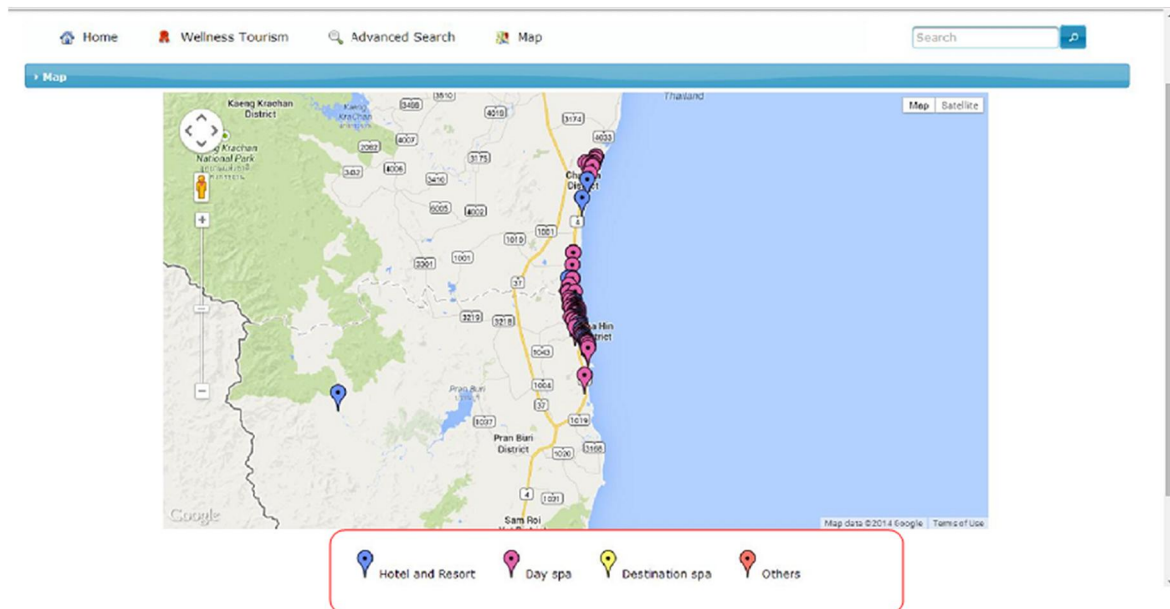
ภาพที่ 3.15 หน้าแรกของเว็บไซต์

ภาพที่ 3.15 แสดงหน้าจอหลักสำหรับการแสดงข้อมูล ผู้ใช้งานสามารถกดปุ่ม See All ที่อยู่ทางด้านขวบนของจอภาพ ซึ่งจะทำได้หน้าจอแสดงในภาพที่ 3.16



ภาพที่ 3.16 หน้า Wellness Tourism

ภาพที่ 3.16 แสดงประเภทของสปา 3 ประเภท คือ Day Spa, Hotel and Resort Spa, Destination Spa, Medical Spa หมายเลข 1 ใช้ค้นหาร้านสปาโดยใช้คำสำคัญ (Keyword) และในกรอบสี่เหลี่ยมมีปุ่มเพื่อกดกลับไปหน้า HOME, Wellness Tourism, Advanced Search, Map การแสดงสถานบริการเหล่านี้จะทำการแสดงแบบสุมมา 3 สถานบริการในแต่ละหมวด



ภาพที่ 3.17 หน้าเว็บเมื่อกดเลือกที่ Map

จากเมนูหลักเมื่อเลือกแท็บ Map จะเป็นการเรียกหน้าเว็บที่แสดงแผนที่ ดังแสดงในภาพที่ 3.17 โดยที่ในการแสดงแผนที่นั้น สีของหมุดที่ปักหมายถึงประเภทของสถานบริการที่ต่างกัน (Hotel and Resort, Day Spa, Destination Spa, Others) นอกจากนี้ผู้ใช้งานสามารถค้นหารายละเอียดลักษณะของสปาในรูปแบบที่ผู้ต้องการด้วยการเลือกแท็บ Advanced Search ทำให้ได้หน้าจอ ดังภาพที่ 3.18

หากต้องการดูแผนที่ว่ามีร้านสปาอยู่ตำแหน่งใดบ้าง ให้เลือกแท็บ Map จะนำมายังหน้าถัดไปดังภาพที่ 3.17 ภาพที่ 3.17 ในการแสดงแผนที่นั้นสีของหมุดที่ปักนั้นจะหมายถึงประเภทสถานบริการที่ต่างกัน ดังในกรอบด้านล่าง

health-tourism.info/huahin/pages/spa_advanced_search.faces

Spa Hua Hin and Cha Am

Home Wellness Tourism Advanced Search Map

Advance Search

Search by Type Facility Massage Service Other Service Advanced Therapy and Product Miscellaneous

Type of Wellness Tourism/Spa

- ☒ Hotel and Resort
- ☐ Destination Spa
- ☐ Day Spa
- ☐ Medical Spa

Type of Health Business According to Act

- ☐ Health Spa
- ☐ Massage for Beauty
- ☐ Massage for Health

Search

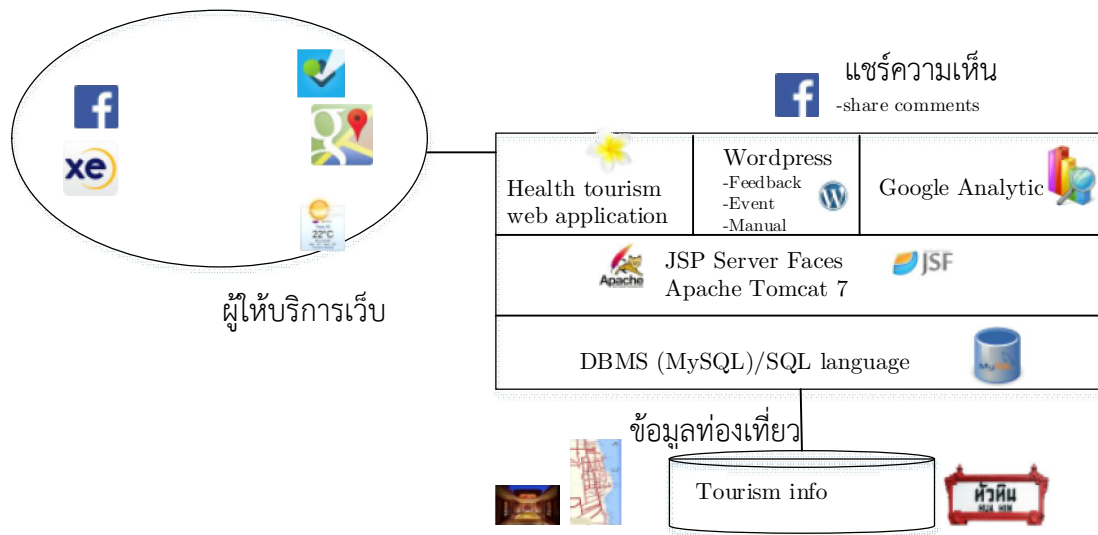
Results

	Name	Address/Contact	Facebook Likes	Foursquare Rating	Type	Register Status	Package	Detail
1	AKA Spa (AKA Resort Hua Hin)	Addr: 152 Moo 7 Baan Nhong Hiang T. Hin Lek Fai A. Hua Hin Tel: 032618900 Email: svn@akaresorts.com	3,155	6.42	Hotel and Resort		☒	Detail
2	AYURAH SPA (Aleenta Resort and Spa)	Addr: 183 Moo 4 Paknampran A. Pranburi Tel: 032618 333 svn.hhc@aleenta.com	4,691	7.09	Hotel and Resort		☒	Detail

ภาพที่ 3.18 หน้าเว็บเมื่อกดเลือก Advanced Search

ในภาพที่ 3.18 ประกอบด้วยการค้นหาตามรายละเอียดต่างๆ ตามแท็บ (1) ในที่นี้แท็บต่างๆ ให้ใช้กำหนดคุณสมบัติในการค้นหา เมื่อคลิกจนขึ้นเครื่องหมายถูกแล้ว จะกดปุ่ม Search (หมายเลข 2) ผลของการค้นหาจะปรากฏในแท็บ Results (หมายเลข 3) โดยแสดงชื่อ (Name) ที่อยู่ (Address) จำนวน Like บนเฟซบุ๊กของสถานบริการ Rating ของ สถานบริการจาก Foursquare ประเภทสถานบริการ สถานบริการนั้นมี package หรือไม่ และปุ่มแว่นขยายเมื่อคลิกจะแสดง รายละเอียด (Detail) ของสถานบริการ รายละเอียดอื่นๆ สามารถอ่านเพิ่มเติมได้จากภาคผนวก ข. แผนงานวิจัย

3.4 การเชื่อมต่อข้อมูลกับระบบอื่นๆ



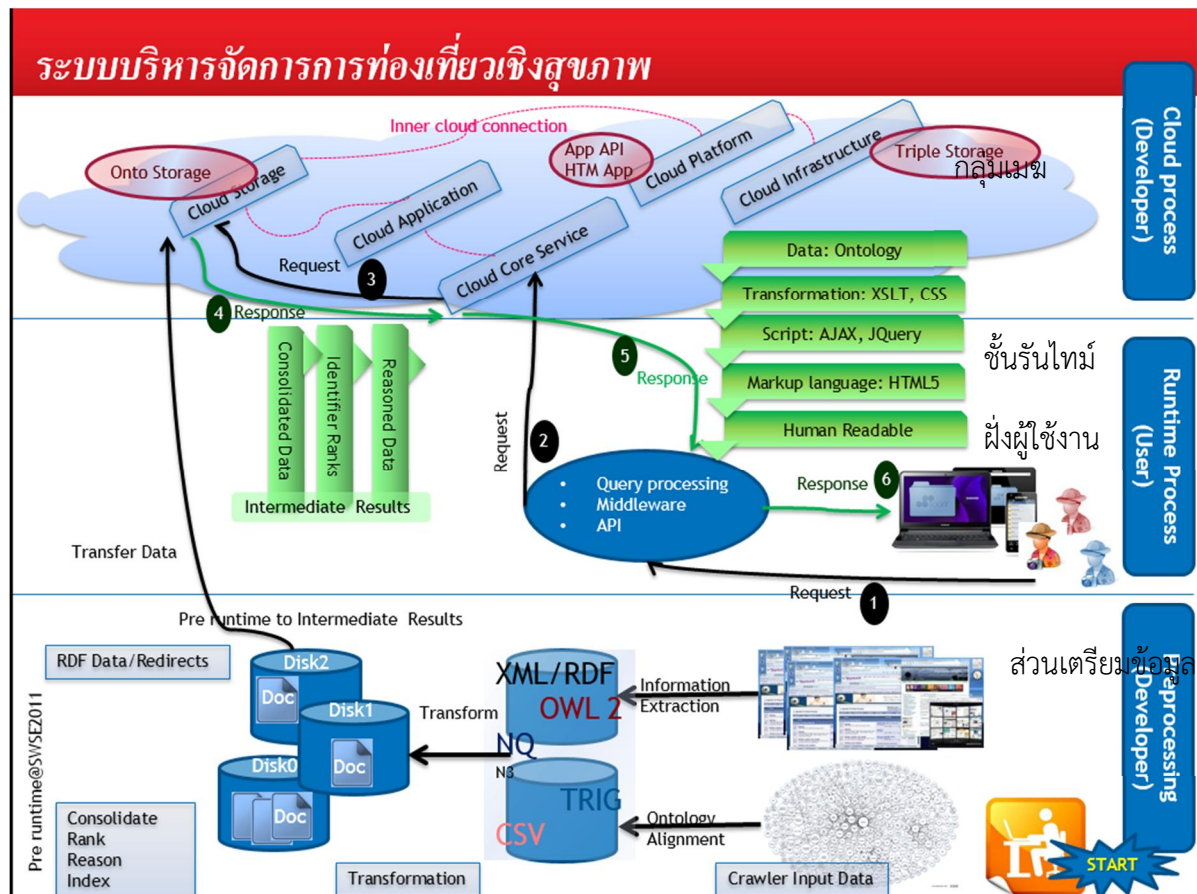
ภาพที่ 3.19 การเชื่อมต่อข้อมูลกับระบบข้อมูลอื่นๆ

สำหรับข้อมูลออนไลน์นั้น สามารถนำออกมาและนำไปใส่ในฐานข้อมูลแบบสัมพันธ์ เพื่อให้ง่ายต่อการดูแลรักษาต่อไป ในต้นแบบระบบนี้ ได้มีการนำฐานข้อมูลที่ได้ มาเชื่อมต่อกับส่วนอื่นๆ โดยมีโครงสร้างซอฟต์แวร์ที่เกี่ยวข้องในลำดับขั้น ได้แก่ ชั้นล่างสุดเป็นข้อมูลที่ได้เก็บมาจากเว็บไซต์และการสำรวจภาคสนาม โดยผ่านการคัดเลือกปรับแต่งรูปแบบให้เหมาะสม นำมาจัดเก็บในฐานข้อมูลที่ออกแบบไว้ ในตัวอย่างนี้ใช้ ฐานข้อมูล MySQL เชื่อมต่อด้วยซอฟต์แวร์ Apache Tomcat 7 และใช้ร่วมกับเฟรมเวิร์กในการพัฒนาเว็บได้แก่ JSP Server Faces ดังภาพที่ 3.19

ในส่วนของเว็บแอปพลิเคชันนั้น พัฒนาโดยมีหน้าจหลักดังจะไดกล่าวต่อไป มีการเชื่อมต่อกับโมดูลต่างๆ เพื่อเก็บข้อมูลข่าวสาร (Event) ความเห็น (Feedback) ข้อเสนอแนะรวมทั้งการแชร์ข้อมูลผ่านเครือข่ายสังคม รวมทั้งคู่มือการใช้งานต่างๆ (Manual) ผ่าน Wordpress และ ใช้ Google analytics เพื่อการตรวจนับสถิติผู้เยี่ยมชม ในหมวดหมู่ต่างๆ

นอกจากนี้ ยังเชื่อมต่อกับระบบภายนอกผ่านเว็บเซอร์วิส เช่นจำนวน Like ของเฟซบุ๊ก Rating จาก Foursquare อัตราแลกเปลี่ยนจาก XE แผนที่จาก Google Map และพยากรณ์อากาศจาก Yahoo

3.5 โครงสร้างต้นแบบระบบบนกลุ่มเมฆ



ภาพที่ 3.20 ระบบบริหารจัดการการท่องเที่ยวยั่งยืนแบบ

จากภาพที่ 3.20 การทำงานของระบบบริหารจัดการการท่องเที่ยวยั่งยืนประกอบด้วย 3 ระดับชั้น คือ ชั้นการจัดเตรียมข้อมูล (Preprocessing) ชั้นรันไทม์ (Runtime Process) และชั้นกลุ่มเมฆ (Cloud Process) โดยที่ชั้นจัดเตรียมข้อมูลนั้น ผู้ดูแลระบบเป็นผู้ดำเนินการเตรียมข้อมูลไว้ให้ผู้ใช้คิวรี ประกอบด้วยโปรแกรมแปลงข้อมูลมาใส่ออนโทโลยี จากนั้นนำไปเก็บไว้ที่เก็บข้อมูลบนกลุ่มเมฆ ร่วมกับแพลตฟอร์ม โครงสร้างและแอปพลิเคชัน เพื่อรอให้ผู้ใช้งานขอการคิวรีจากชั้นรันไทม์มายังบริการในชั้นบนกลุ่มเมฆ การตอบกลับจากชั้นบนกลุ่มเมฆไปยังผู้ใช้แบ่งเป็นสองส่วนคือส่วนข้อมูลที่อยู่ในโครงสร้าง RDF และส่วนการแสดงผลทางหน้าเบราเซอร์

ภายในระดับชั้นของระบบการประมวลผลบนกลุ่มเมฆที่เกี่ยวข้องกับโครงงานย่อยที่ 1 แบ่งได้ 3 ระดับได้แก่ SaaS (Software as a Service) PaaS (Platform as a Service) และ IaaS (Infrastructure as a Service) มีรายละเอียดแต่ละชั้น ดังต่อไปนี้

- 1) **ชั้น IaaS** ประกอบด้วยเซิร์ฟเวอร์ ระบบเครือข่ายและส่วนเก็บออนโทโลยี
- 2) **ชั้น PaaS** ประกอบด้วยส่วนมิดเดิลแวร์ส่วนจัดการออนโทโลยี การใส่รายละเอียดหรือคำบรรยายหน้าเว็บและโปรแกรมส่วนรันไทม์

- 3) **ชั้น SaaS** ประกอบด้วย เว็บแอปพลิเคชัน HTML5 เว็บเซอร์วิสของผู้ให้บริการต่างๆ เช่น สกู๊ปเงิน สภาพอากาศ แผนที่ การท่องเที่ยวเชิงสุขภาพ และหน้าจอบควบคุมแอปพลิเคชันเพื่อการส่งออกไปยังเครื่องมือต่างๆ

3.6 การประเมินความพึงพอใจ

คณะผู้วิจัยได้สร้างแบบสอบถามเพื่อสำรวจความพึงพอใจในด้านการให้บริการข้อมูล โดยทำการเก็บรวบรวมข้อมูลจากแบบสอบถามผู้ใช้งานระบบ และใช้ค่าทางสถิติที่ทำการวัดแนวโน้มเข้าสู่ส่วนกลางด้วยการวัดค่ากลางของข้อมูลโดยใช้ค่าเฉลี่ย (Mean) และวัดการกระจายของข้อมูลโดยใช้ส่วนเบี่ยงเบนมาตรฐาน (Standard Deviation)

ค่าตัวกลางเลขคณิต (Arithmetic Mean) หรือค่าเฉลี่ย

$$X = \frac{\sum_{i=1}^n x_i}{n}$$

โดยที่ X แทนตัวกลางเลขคณิตหรือค่าเฉลี่ย

n แทนจำนวนข้อมูลทั้งหมด

x_i แทนข้อมูลดิบตัวที่ i

$\sum_{i=1}^n x_i$ แทนผลรวมทั้งหมดของข้อมูล

ค่าส่วนเบี่ยงเบนมาตรฐาน (Standard deviation)

$$SD = \sqrt{\frac{\sum_{i=1}^N (x_i - X)^2}{N}}$$

โดยที่ SD แทนค่าเบี่ยงเบนมาตรฐานของกลุ่มตัวอย่าง

X แทนค่าเฉลี่ยของกลุ่มตัวอย่าง

x_i แทนค่าของข้อมูล

i แทน 1,2,3,..., N

N แทนค่าของจำนวนข้อมูล (กลุ่มตัวอย่าง)

คณะผู้วิจัยได้กำหนดเกณฑ์การให้คะแนนในแบบประเมินตามตารางที่ 3.1 เกณฑ์การให้คะแนนของแบบประเมินและกำหนดเกณฑ์การแปลความหมายข้อมูลและพิจารณาจากค่าตัวกลางเลขคณิต(ค่าเฉลี่ย) ดังแสดงในตาราง 3.2

ตารางที่ 3.1 เกณฑ์การให้คะแนนของแบบประเมิน

ระดับเกณฑ์	ความหมาย
5	ประสิทธิภาพของระบบอยู่ในระดับดีมาก
4	ประสิทธิภาพของระบบอยู่ในระดับดี
3	ประสิทธิภาพของระบบอยู่ในระดับปานกลาง
2	ประสิทธิภาพของระบบอยู่ในระดับน้อย
1	ประสิทธิภาพของระบบอยู่ในระดับควรปรับปรุง

ตารางที่ 3.2 เกณฑ์การแปลความหมายข้อมูลและพิจารณาจากค่าตัวกลางเลขคณิต (ค่าเฉลี่ย)

ระดับเกณฑ์	ความหมาย
4.50 – 5.00	ประสิทธิภาพของระบบอยู่ในระดับดีมาก
3.50 – 4.49	ประสิทธิภาพของระบบอยู่ในระดับดี
2.50 – 3.49	ประสิทธิภาพของระบบอยู่ในระดับปานกลาง
1.50 – 2.49	ประสิทธิภาพของระบบอยู่ในระดับน้อย
1.00 – 1.49	ประสิทธิภาพของระบบอยู่ในระดับควรปรับปรุง

บทที่ 4 ผลการวิจัย

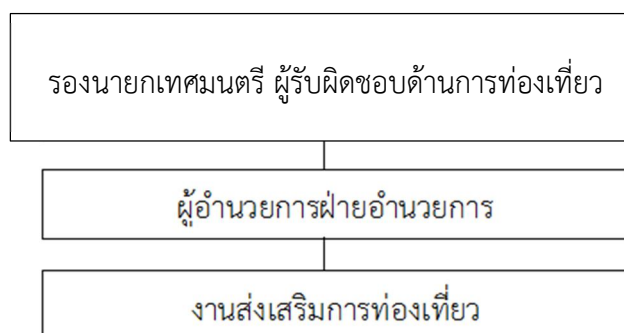
ผลการวิจัยดำเนินงานแบ่งเป็น 3 ส่วน ดังนี้ ส่วนแรกเกี่ยวกับการสำรวจลักษณะของซอฟต์แวร์ที่ใช้ในปัจจุบันของหน่วยงาน และข้อมูลที่หน่วยงานใช้อยู่ ส่วนที่ 2 เป็นผลการประเมินประสิทธิภาพด้านเวลาของระบบเว็บเชิงความหมาย การประมวลผลแบบขนาน ส่วนที่ 3 เป็นผลการประเมินประสิทธิภาพด้วยเวลาสำหรับการประมวลผลการดึงข้อมูล และค้นหาข้อมูลแบบขนาน ส่วนการประเมินผู้ใช้งานนั้นอยู่ในภาคผนวกของรายงานแผนงานวิจัย จึงไม่ขอกล่าวถึงในที่นี้

4.1 การสำรวจข้อมูลของเทศบาลด้านลักษณะหน่วยงาน ข้อมูลที่มีและซอฟต์แวร์ต่างๆ

คณะผู้วิจัยได้ สอบถามข้อมูลซอฟต์แวร์ที่ใช้ภายในหน่วยงานเทศบาล และการทำงานของหน่วยงาน แหล่งข้อมูลต่างๆ ของการท่องเที่ยวที่ทางเทศบาลใช้อยู่ ด้วยการสัมภาษณ์และแบบสอบถาม โดยรายละเอียดอยู่ในรายงานแผนการวิจัย หัวข้อ 4.1 และ 4.2 และในภาคผนวก ก. ของรายงานแผนการวิจัย

4.1.1 ข้อมูลเกี่ยวข้องกับหน่วยงานเทศบาลเมืองหัวหิน

ผู้ที่เกี่ยวข้องกับโครงการนี้ในหน่วยงานของเทศบาลเมือง คือ นายกเทศมนตรีมอบหมายงานให้รองนายกเทศมนตรีด้านการท่องเที่ยว คือคุณมนตรี ชูภู เป็นผู้อนุมัติการประเมินระบบนี้



ภาพที่ 4.1 แผนผังหน่วยงานที่เกี่ยวข้องกับงานวิจัย

รองนายกเทศมนตรี ผู้รับผิดชอบด้านการท่องเที่ยว ได้มอบหมายฝ่ายอำนวยความสะดวก คือคุณเบญจวรรณ เมธาวิกุล มีหน้าที่ควบคุมดูแลและรับผิดชอบการปฏิบัติงานในหน้าที่ที่เกี่ยวข้องกับงานส่งเสริมการท่องเที่ยว มีหัวหน้างานคือ คุณอุกฤษฏ์ ป้อทองคำ มีพนักงาน 7 คน สำหรับงานส่งเสริมการท่องเที่ยว มีหน้าที่เกี่ยวกับ

- งานส่งเสริมและเผยแพร่การท่องเที่ยวของเทศบาลภายในประเทศ
- งานควบคุมกิจการบริการท่องเที่ยวภายในประเทศของเทศบาลและหน่วยงานอื่น
- งานต้อนรับ แนะนำ อำนวยความสะดวกและบริการนำเที่ยว รวมทั้งดำเนินธุรกิจการท่องเที่ยว

- งานประชาสัมพันธ์ ให้คำแนะนำเผยแพร่ความรู้เกี่ยวกับแหล่งท่องเที่ยวหรือกิจการ ที่เกี่ยวกับการท่องเที่ยวภายในจังหวัดและเทศบาล
- งานจัดทำรายงานและสถิติการเคลื่อนไหวเกี่ยวกับธุรกิจการท่องเที่ยว
- งานอื่นที่เกี่ยวข้องหรือตามที่ได้รับมอบหมาย

4.1.2 สรุปข้อมูลที่ได้จากการสำรวจความต้องการและข้อมูลที่มีอยู่

รายละเอียดของการสำรวจความต้องการซอฟต์แวร์จากหน่วยงานเทศบาลเมืองหัวหินสำหรับการสำรวจข้อมูลในพื้นที่นั้นได้ทำ 2 ครั้ง นั่นคือ ครั้งที่ 1 เมื่อวันที่ 23 สิงหาคม พ.ศ. 2556 และได้กำหนดวันนัดหมายเพื่อพูดคุยรายละเอียดกันในเดือนกันยายน แต่ได้มีการเลื่อนเป็นครั้งที่ 2 วันที่ 4 ตุลาคม พ.ศ. 2556 ดังรายละเอียดดังต่อไปนี้

ระดับผู้บริหารนั้นมีความต้องการใช้ข้อมูลที่ผ่านการรวบรวมอย่างเป็นระบบ เนื่องจากการนำเสนอข้อมูลจากหน่วยงานต่อสาธารณชนนั้นควรมีรายละเอียดรอบด้านมากขึ้น

การสัมภาษณ์ผู้อำนวยการฝ่ายอำนวยการ และสำรวจจากเทศบาลเมืองหัวหิน หัวหน้างานส่งเสริมการท่องเที่ยว และเจ้าหน้าที่การท่องเที่ยวแห่งประเทศไทยในเขตเมืองหัวหินและพนักงานปฏิบัติการ คุณเกษมศักดิ์ ปัตตอด มีดังแบบสำรวจความต้องการการใช้งานซอฟต์แวร์ของระบบการจัดการท่องเที่ยวเชิงสุขภาพ สรุปผลการสำรวจได้ดังนี้

- ระดับหน่วยงานส่งเสริมการท่องเที่ยวมีศูนย์บริการนักท่องเที่ยว 2 แห่งคือเทศบาลเมืองหัวหินและหอนาฬิกา
- การรวบรวมข้อมูลเป็นการทำงานแบบระบบมือ
- การค้นหาข้อมูลเกี่ยวกับที่พัก ร้านอาหาร สถานที่แอ่งภัย ส่วนมากใช้ระบบเดินสำรวจหรือค้นหาทางเว็บไซต์เฉพาะที่ ในขณะที่สถานพยาบาล สปา นวดแผนไทย อายุรเวช นั้นไม่เคยสำรวจ
- การติดต่อสื่อสารระหว่างหน่วยงานเทศบาลและที่พักชั่วคราวอาทิโรงแรมในพื้นที่มักใช้โทรศัพท์และการเดินเอกสาร สำหรับการใช้งานโปรแกรมจองโรงแรมหรือเว็บไซต์เกี่ยวกับการท่องเที่ยว เช่น Agoda, TripAdvisor, AtSiam นั้น ไม่เคยใช้ สำหรับโปรแกรมการค้นหาเว็บไซต์นั้นใช้กูเกิล ส่วนเครือข่ายสังคมออนไลน์นั้นใช้เฟสบุ๊คเพื่อกระจายข่าวสารเทศบาลและการท่องเที่ยว
- สำหรับเครื่องมือที่ใช้ค้นหาผ่านอินเทอร์เน็ตคือคอมพิวเตอร์ส่วนบุคคลที่เชื่อมต่ออินเทอร์เน็ตทั้งแลนและระบบไร้สาย
- การทำรายงานนำเสนอผู้บริหารประกอบด้วยเชื้อชาติของนักท่องเที่ยวและจำนวนผู้สอบถามเป็นรายวันรายเดือนและรายปีผ่านโปรแกรมไมโครซอฟต์เอ็กเซล

- สำหรับเอกสารเผยแพร่ข้อมูลจัดทำเป็นแผ่นพับสองภาษาคือภาษาไทยและภาษาอังกฤษ และมีการปรับปรุงประมาณปีละครั้ง
- สิ่งที่ต้องการจากระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพคือข้อมูลโรงแรมเบื้องต้นเพิ่มเติมจากการลงพื้นที่เองและจากกูเกิล ประโยชน์ที่จะนำไปใช้นอกจากรายงานสำหรับผู้บริหารคือ เป็นแหล่งข้อมูลสำหรับนักวิจัย นักเรียน นักศึกษา ในกรณีที่มีการร้องขอข้อมูลอย่างละเอียด
- ข้อมูลที่ร้องขอเพิ่มเติมคือหมายเลขโทรศัพท์และอีเมลของสถานทูต ข้อมูลสำนักงานตรวจคนเข้าเมืองเพื่อต่อวีซ่า
- ต้องการใช้ระบบบริหารจัดการการท่องเที่ยวเชิงสุขภาพ เพื่อรวบรวมข้อมูลร้อยละ 45 เพื่อตรวจสอบข้อมูลร้อยละ 45 ใช้ประกอบรายงานสำหรับผู้บริหารร้อยละ 10
- หากโปรแกรมเสร็จแล้วมีความต้องการอบรมเพื่อใช้งานโปรแกรม

นอกจากนี้ได้มีการสอบถามข้อมูลการวางแผนท่องเที่ยวจากสำนักงานการท่องเที่ยวแห่งประเทศไทย สำนักงานประจวบคีรีขันธ์วันที่ 23 สิงหาคม พ.ศ. 2556 ในเรื่องการวางแผนการท่องเที่ยวสำหรับนักท่องเที่ยวทั้งชาวไทยและชาวต่างประเทศซึ่งมีการเก็บข้อมูลนักท่องเที่ยวที่คล้ายคลึงกับเทศบาลเมืองคือเชื้อชาติและจำนวนนักท่องเที่ยวที่มาสอบถามรายวัน สำหรับข้อมูลรายได้และนักท่องเที่ยวโดยรวมของภาคกลางจะได้รับมาจากการท่องเที่ยวแห่งประเทศไทยสำนักงานใหญ่การวางแผนท่องเที่ยวในส่วนข้อมูลทั่วไปและกิจกรรมการท่องเที่ยวที่เผยแพร่บนเว็บไซต์ที่เป็นผู้รวบรวมขึ้นเองส่วนเครื่องมือค้นหาเว็บไซต์นั้นใช้งานผ่านกูเกิลสำหรับรายละเอียดของหน้าที่สำนักงาน กิจกรรมการท่องเที่ยว รายละเอียดนักท่องเที่ยวและแผนการท่องเที่ยวมีให้ตามแผ่นพับของ ททท.

4.1.3 ข้อมูลจากหน่วยงานเพื่อตรวจสอบความถูกต้อง

สำหรับสถานบริการสปาหรือนวดแผนโบราณที่เป็นเป้าหมายของงานวิจัย ในกรณีที่เป็นโรงแรมหรือที่พักขนาดใหญ่ที่มีอยู่ใน Agoda หรือ TripAdvisor นั้น จะถูกดึงเข้าสู่ระบบโดยโปรแกรมคอมพิวเตอร์ อย่างไรก็ตามยังมีสถานบริการริมชายหาดและสถานบริการเชิงสุขภาพขนาดเล็ก ซึ่งผู้อำนวยการฝ่ายอำนวยความสะดวกได้กรุณาให้ความรู้ในไว้ดังนี้

- กรณีที่เป็นการนวดริมชายหาดของเอกชนนั้นสามารถตรวจสอบได้กับเทศกิจ
- กรณีที่เป็นร้านขนาดเล็ก จะมีการขึ้นทะเบียนกับกองสาธารณสุข สิ่งที่ได้คือชื่อร้านและที่อยู่
- สำหรับพิกัดที่ตั้งนั้น งานแผนที่ภาษีมีบันทึกไว้ในระบบพิกัดกริดแบบ UTM รายชื่อผู้เสียภาษีและที่อยู่

- เมื่อต้องการนำมาตรวจสอบกับระบบที่ผู้วิจัยเก็บไว้เพื่อให้ได้ชื่อร้านและพิกัดละติจูด/ลองจิจูด ต้องนำข้อมูลมาจากทั้ง 3 หน่วยงานนี้มาตรวจสอบร่วมกัน

4.2 ผลการประเมินประสิทธิภาพ

การวัดประสิทธิภาพด้านเวลา มี 2 แบบ แบบแรกเป็นการวัดประสิทธิภาพด้านเวลาการค้นหาด้วยเว็บเชิงความหมายโดยค้นหาในออนไลน์ และแบบที่สองเป็นการค้นหาสายอักขระด้วยการประมวลผลแบบขนานด้วยหน่วยประมวลผลกราฟิก

4.2.1 เวลาการค้นหาด้วยเว็บเชิงความหมาย

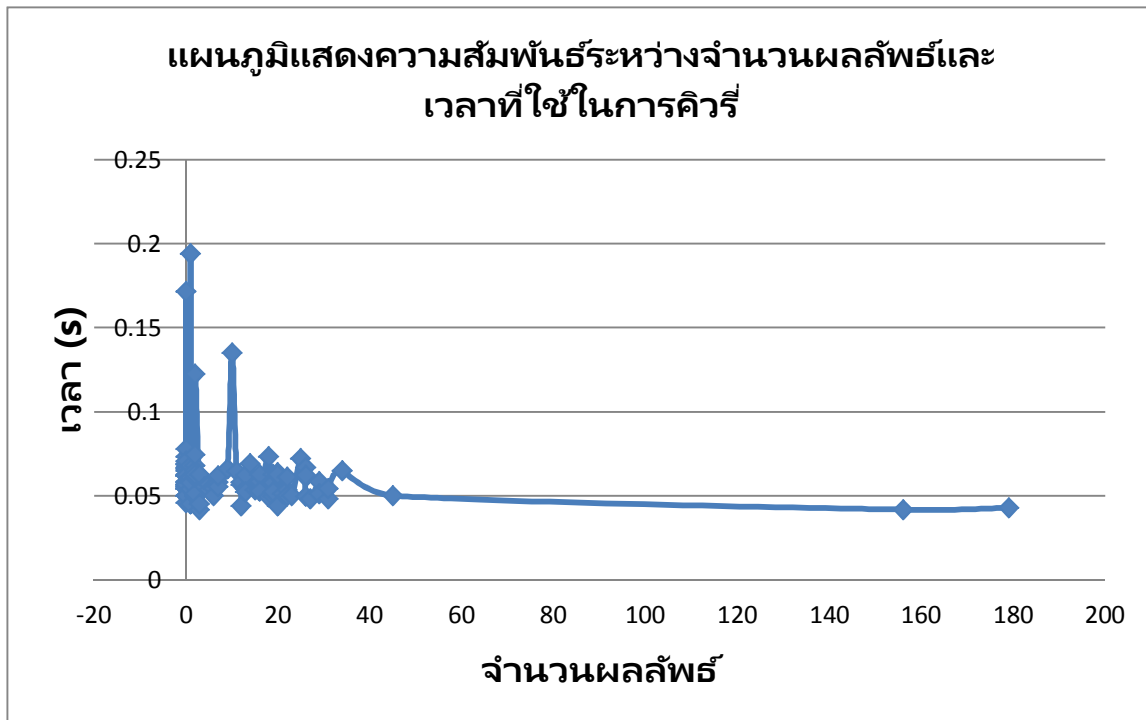
ในการทดสอบประสิทธิภาพ คณะผู้วิจัยได้ทดสอบวัดความเร็วคิวรีจากจำนวนคิวรี 100 ตัวอย่าง ทั้งที่ไม่มี ความซับซ้อน มีความซับซ้อนน้อยไปจนถึงตัวอย่างที่มีความซับซ้อนมากเพื่อวัดเวลาที่ใช้ในการคิวรีข้อมูล ตารางที่ 5.1 แสดงรายละเอียดข้อมูลออนไลน์

ตารางที่ 4.1 ข้อมูลออนโทโลยี

ข้อมูลทั้งหมดในออนโทโลยี	จำนวน
Axiom (ข้อกำหนดความสัมพันธ์)	24663
Logical axiom (ความสัมพันธ์เชิงตรรกะ)	21094
Class (คลาส)	229
Object Property (ความสัมพันธ์)	33
Data Property (คุณสมบัติ)	29
Individual (สมาชิกของคลาส)	3008
ข้อจำกัดความสัมพันธ์ในคลาส	จำนวน
SubClassOf (คลาสย่อย)	227
EquivalentClasses (คลาสเทียบเท่า)	19
DisjointClasses (คลาสที่ไม่เกี่ยวข้องกัน)	236
Hidden GCI (General Concept Inclusions)	13
ข้อจำกัดความสัมพันธ์ใน Object Property	จำนวน
SubObjectPropertyOf (ความสัมพันธ์ย่อย)	2
EquivalentObjectProperties (ความสัมพันธ์เทียบเท่า)	2
InverseObjectProperties (ความสัมพันธ์ที่แบบผกผัน)	16
FunctionalObjectProperty (ความสัมพันธ์ระหว่างกันแบบหนึ่งต่อหนึ่ง)	2
TransitiveObjectProperty (ความสัมพันธ์ที่มีการส่งต่อไปเรื่อยๆ)	4
ObjectPropertyDomain (สมาชิกภายในเซตข้อมูลที่ใช้จับคู่กับอีกเซตหนึ่ง)	33
ObjectPropertyRange (เซตข้อมูลที่ถูกจับคู่โดยความสัมพันธ์)	26
ข้อจำกัดความสัมพันธ์ใน Annotation	จำนวน
AnnotationAssertion (คำอธิบายประกอบคุณสมบัติ ความสัมพันธ์ และข้อมูล)	271
ข้อจำกัดความสัมพันธ์ใน Individual	
ClassAssertion (การกำหนดสมาชิกในคลาส)	3076
ObjectPropertyAssertion (การกำหนดความสัมพันธ์ของข้อมูล)	9943
DataPropertyAssertion (การระบุคุณสมบัติของข้อมูล)	7454

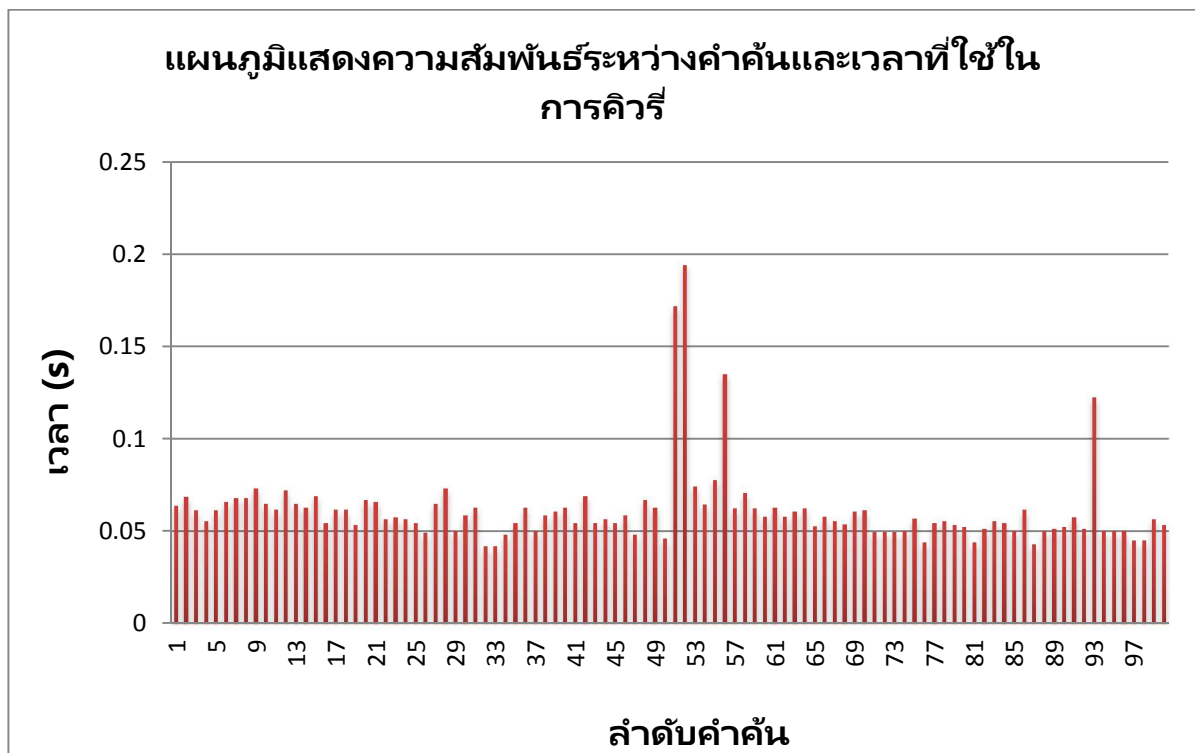
คณะผู้วิจัยทำการทดลองด้วยการสุ่มการค้นหาแบบ Advanced Search ขึ้นมาดังคอลัมน์เงื่อนไขการค้นหาในตารางภาคผนวก ข. ของโครงการ ทั้งนี้เครื่องคอมพิวเตอร์ที่ใช้ในการทดสอบมีคุณสมบัติ Intel Core i5-3230M ความเร็ว 2.60 GHz หน่วยความจำ ขนาด 4 GB DDR3 ฮาร์ดดิสก์ขนาด 640 GB การ์ดอีเทอร์เน็ตและใช้ระบบปฏิบัติการ Windows 7 เป็นต้น

ผลการทดลองวัดเวลาที่ใช้ในการค้นหาพบว่า แต่ละตัวอย่างที่ใช้ทดลองควิรีใช้เวลาในการค้นหาใกล้เคียงกัน และใช้เวลาไม่มาก โดยเฉลี่ยใช้น้อยกว่า 1 วินาที (เวลาโดยเฉลี่ย 0.061618 วินาที) ทั้งนี้ ความซับซ้อนของคำค้นและจำนวนผลลัพธ์ของการสืบค้นไม่ส่งผลให้ใช้เวลาในการค้นหาแตกต่างกันมาก



ภาพที่ 4.2 แผนภูมิแสดงความสัมพันธ์ระหว่างจำนวนผลลัพธ์และเวลาที่ใช้ในการควิรี

ภาพที่ 4.2 แสดงแผนภูมิความสัมพันธ์ระหว่างจำนวนผลลัพธ์และเวลาที่ใช้ในการควิรี แสดงให้เห็นว่า จำนวนผลลัพธ์ของการสืบค้นไม่ส่งผลให้เวลาที่ใช้ในการควิรีแตกต่างกันมาก แต่ละตัวอย่างที่ใช้ทดลองในการควิรีมีเวลาใกล้เคียงกัน และใช้เวลาโดยเฉลี่ยไม่ถึง 1 วินาที (เวลาเฉลี่ย 0.061618 วินาที) ทั้งนี้จากแผนภูมิจะเห็นได้ว่ากราฟที่ได้มีลักษณะเป็นเส้นตรง และมีบางข้อมูลที่มีค่าเวลาที่ใช้ในการควิรีมากกว่าข้อมูลอื่นๆ เนื่องจากมีปัจจัยภายนอกที่ส่งผลให้เวลาที่ใช้ในการควิรีมากขึ้น เช่น ความเร็วของอินเทอร์เน็ต เป็นต้น



ภาพที่ 4.3 แผนภูมิแสดงความสัมพันธ์ระหว่างคำค้นและเวลาที่ใช้ในการควรี

ภาพที่ 4.3 แผนภูมิแสดงความสัมพันธ์ระหว่างคำค้นและเวลาที่ใช้ในการควรี แสดงให้เห็นว่า ความซับซ้อนของคำค้นไม่ส่งผลให้เวลาที่ใช้ในการควรีแตกต่างกันมาก แต่ละตัวอย่างที่ใช้ทดลองในการควรีมี เวลาใกล้เคียงกันและใช้เวลาโดยเฉลี่ยไม่ถึง 1 วินาที (เวลาเฉลี่ย 0.061618 วินาที) ทั้งนี้จากแผนภูมิจะเห็นได้ ว่าความสูงของกราฟที่ได้มีลักษณะเป็นเส้นตรง และมีบางข้อมูลที่มีค่าเวลาที่ใช้ในการควรีมากกว่าข้อมูลอื่นๆ เนื่องจากมีปัจจัยภายนอกที่ส่งผลให้เวลาที่ใช้ในการควรีมากขึ้น เช่น ความเร็วของอินเทอร์เน็ต เป็นต้น

4.2.2 การประเมินเวลาของการนำเข้าข้อมูล

ตารางที่ 4.2 แสดงผลการค้นหาด้วย Java Concurrency โดยการแตกเทร็ด โดยทดสอบกับ เครื่องคอมพิวเตอร์ที่มี 4 คอร์ Intel® Core™ i3 CPU M 350 @ 2.27GHz (64-bit) และ RAM 3.7 GB โดยทดสอบกับเว็บไซต์ในตาราง ซึ่งจะเห็นว่าเมื่อมีจำนวนเทร็ดมากขึ้น จะช่วยให้เวลาที่ใช้ในการควรีลดลง

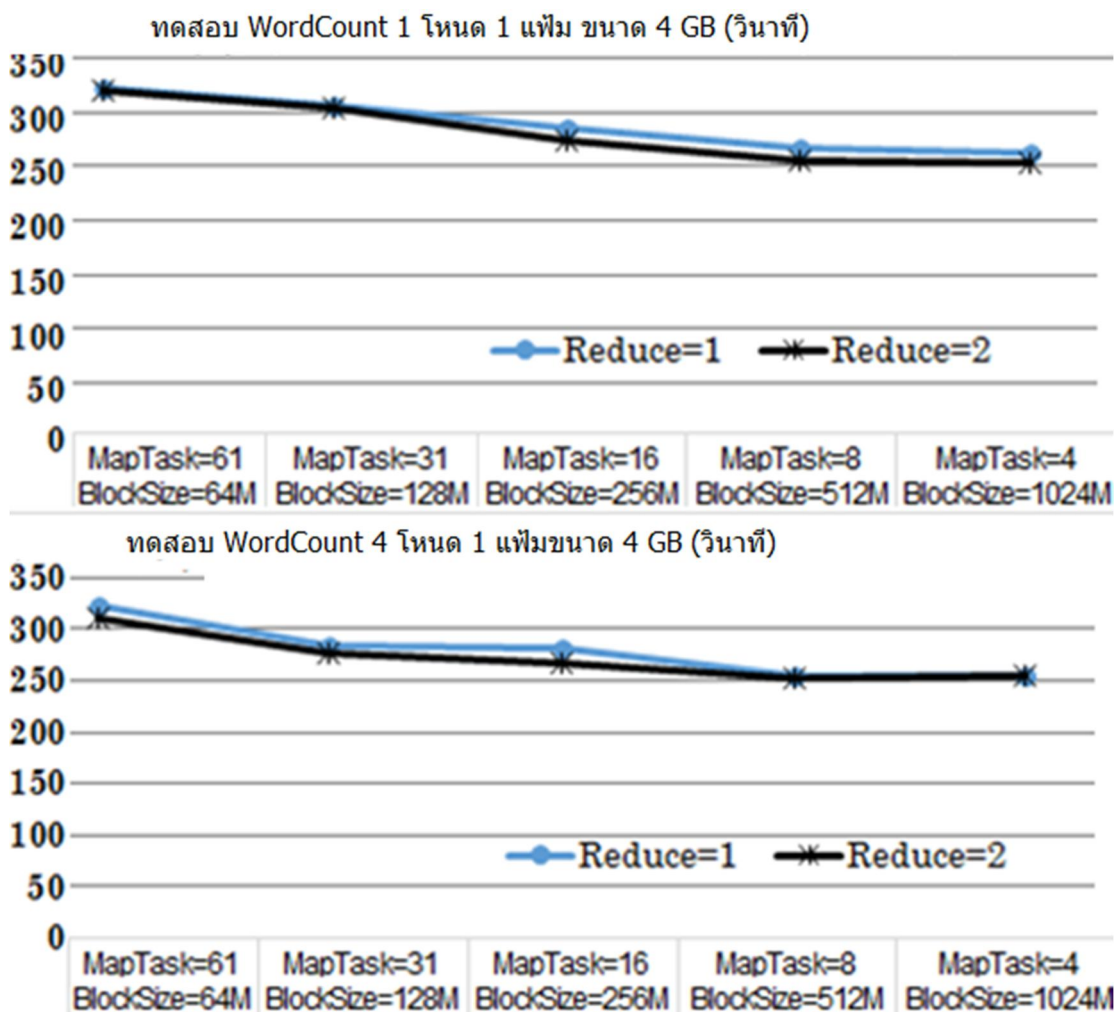
ภาพที่ 4.4 เวลาในการทำงานหลายงานของ MapReduce เมื่อทำการรวมแฟ้มข้อมูล โดยใน เบื้องต้นทดสอบกับข้อมูลจาก DBpedia BTC2012¹² โดยแฟ้มข้อมูลมีขนาด 4GB จำนวน 21,398,608 Triples

¹² <http://km.aifb.kit.edu/projects/btc-2012/>

ตารางที่ 4.2 เวลาการค้นหาด้วย Java Concurrency

Website/Thread	1 เทรด(ms)	2 เทรด (ms)	4 เทรด (ms)
Heyhuahin (82 webs)	4,602	2,529	1,994
Atsiam (146 webs)	100,012	64,728	45,763
Tripadvisor (247 webs)	270,121	135,600	73,981
Agoda (237 webs)	194,443	143,590	74,935

มีการใช้ 3 โหนด Slave และมีโหนดหลัก Master จำนวน 1 โหนด ซึ่งมี ลักษณะดังนี้ โหนด Slave 2 ตัว เป็น Intel ® Core™ 2 Duo CPU E8400 @ 3.00GHz (64-bit), 1.9 GB RAM และอีกโหนด Slave เป็น Intel ® Core™ i5-3330 CPU@ 3.00GHz (64-bit), 7.7 GB RAM และ โหนด Master เป็น Intel ® Core™ i5-3330 CPU@ 3.00GHz (64-bit), 7.7 GB RAM



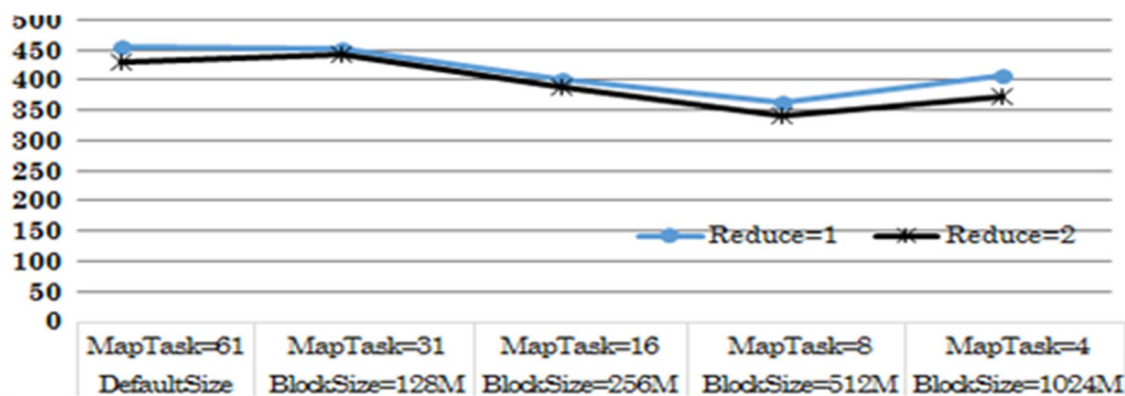
ภาพที่ 4.4 เวลาในการทำงานหลายงานของ MapReduce เมื่อทำการรวมแฟ้มข้อมูลภาพที่ 4.4 ภาพบนคือ เวลาของการใช้เพียง 1 โหนด และภาพล่างคือเวลาของการใช้ 4 โหนด โดยงาน Map มีจำนวนตั้งแต่ 4 ถึง

61 และขนาดบล็อกข้อมูลมีขนาดตั้งแต่ 64, 128, 256, 512 และ 1024 MB และทดสอบกับงาน Reduce 1 หรือ 2 งาน

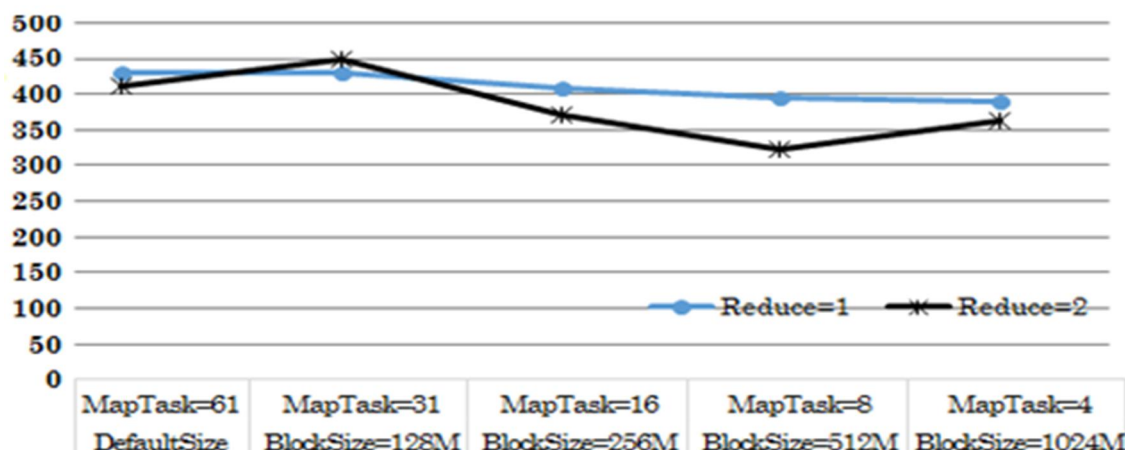
ผลการทดลองพบว่าเวลาที่ใช้ลดลงเมื่อจำนวนงาน Map ลดลง และเมื่อขนาดบล็อกข้อมูลใหญ่ขึ้น สำหรับคลัสเตอร์ขนาดเล็กนี้ จากงานวิจัย [1,2] พบว่าจำนวนงาน Map จะต้องสอดคล้องกับจำนวนคอร์ทั้งหมดที่มี และจำนวนงาน Reduce มีเพียง 1 หรือ 2 งาน ก็เพียงพอ แต่ควรกำหนดขนาดบล็อกข้อมูลให้ใหญ่ให้เหมาะสมเพียงพอเพื่อให้เวลาการทำ I/O ไม่มากเกินไปกว่าเวลาคำนวณในบล็อกนั้นจริงๆ

ภาพที่ 4.5 แสดงเวลาในการสร้าง S(PO) ดัชนีในการจัดกลุ่มค่า Predicate และ Object สำหรับแต่ละ Subject ด้วย MapReduce เช่นกัน เป็นการรวมต่ออักขระ โดยใช้ key เดียวกันคือ Subject ซึ่งจะเป็นการรวมข้อมูลในลักษณะเดียวกับภาพ 4.4

สร้างดัชนีสำหรับ S(PO) ใช้ 1 โหนด แฟ้มขนาด 4GB จำนวน 1 แฟ้ม (วินาที)



สร้างดัชนี S(PO) ใช้ 4 โหนด แฟ้มขนาด 4 GB จำนวน 1 แฟ้ม (วินาที)



ภาพที่ 4.5 เวลาในการสร้าง Index โดยใช้ 1 โหนดและ 4 โหนด

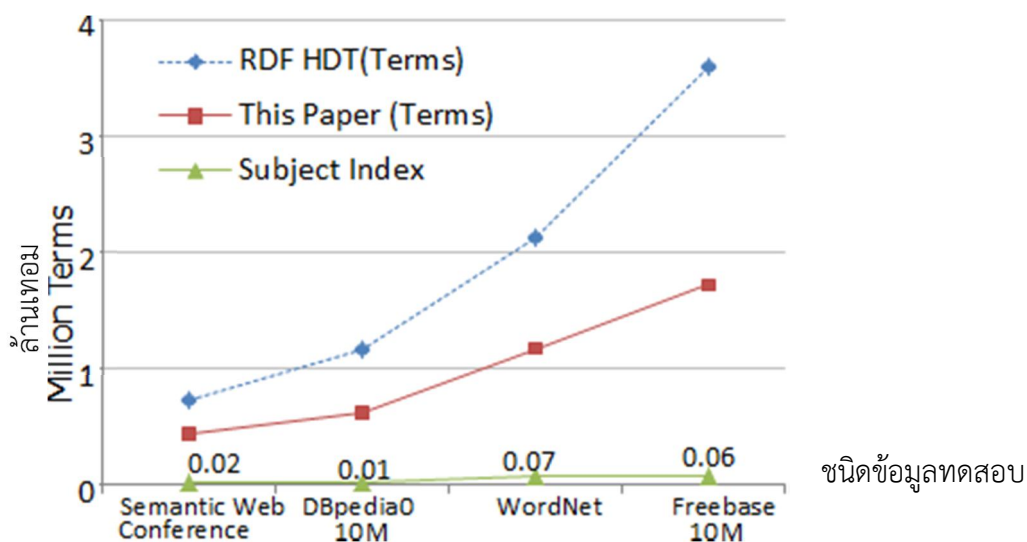
ผลการทดสอบพบว่าประสิทธิภาพด้านเวลาของการใช้ 1 โหนดกับ 4 โหนด ในกรณีนี้พบว่า 4 โหนดทำงานได้เร็วกว่า เมื่อให้จำนวนงาน Map เป็น 8 ถ้ามีจำนวนโหนดมากขึ้น ก็จะมีผลด้านเวลาดีขึ้น และเมื่อกำหนดขนาดของแฟ้มให้ใหญ่ในระดับ GB การทำงานของ MapReduce เป็นการทำงานที่เน้นข้อมูลมากๆ และมีการ Overlap การคำนวณระหว่างงานเมื่องาน Map ทำงานเสร็จ Reduce ก็ทำงานต่อได้ และ

งาน Map อื่นๆ ก็ทำงานต่อไป จึงจะทำให้การทำงานนั้นต่อเนื่องงาน Map จึงต้องประมวลผลกับข้อมูลจำนวนมาก นอกจากนี้เวลาในการเริ่มต้นต่างๆ เช่นการเริ่มต้นของ Java Virtual Machine (JVM) ก็ไม่ต่ำกว่า 0.3-0.5 วินาที^{13,14} ดังนั้น ขนาดข้อมูลในบล็อกควรใหญ่พอเพียงด้วยให้งาน Map ทั้งหลายทำงานพร้อมกันด้วย แต่ก็ทำให้เกิดจำนวนงาน Map มากเกินไปเพราะจะให้ Overhead มากเนื่องจากมีงานในระบบเป็นจำนวนงานเกินไปเทียบกับจำนวนคอร์จริงของเครื่องที่มี

เมื่อเทียบกับเวลาการรวมแฟ้มข้อมูลทั่วไปด้วยโปรแกรม Java แบบลำดับจะใช้เวลามากกว่า 10 นาทีสำหรับข้อมูลขนาด 4GB หรือบางครั้งอาจจะรันไม่ได้เลย การปรับพารามิเตอร์ต่างๆ เช่นจำนวนงาน Map จำนวนงาน Reduce ขนาดของบล็อก ขนาดของแฟ้มข้อมูลทั้งหมด จำนวน VM และอื่นๆ ของ Hadoop ให้เหมาะกับคลัสเตอร์ที่มีค้อยข้างขึ้นกับลักษณะของงานและคลัสเตอร์อย่างมาก ดังได้กล่าวไว้ใน [1-5]

4.2.3 การประเมินเวลาของการค้นหาด้วยต้นไม้ไบนารี

เมื่อใช้การแทนค่าแบบไบนารีหรือต้นไม้ 0-1 จะเห็นว่าจะมีขั้นตอนเพิ่มกว่าการใช้รูปแบบ RDF ทั่วไป เพราะจำต้องมีการแปลงจาก RDF เป็น HDT และจาก HDT เป็นไบนารี จากนั้นจึงนำข้อมูลที่ผ่านการแปลงเรียบร้อยแล้วเข้าหน่วยประมวลผลกราฟิกส์ เพื่อทำการค้นหา การแปลงเป็นไบนารีย่อมได้เปรียบในการย่อขนาดข้อมูลซึ่งจะทำให้สามารถนำข้อมูลเข้าไปในหน่วยประมวลผลกราฟิกส์ได้มากกว่ารูปแบบ RDF ทั่วไป ดังนั้นในขั้นแรกคณะผู้วิจัยทดสอบวัดขนาดของข้อมูลในรูปแบบไบนารี

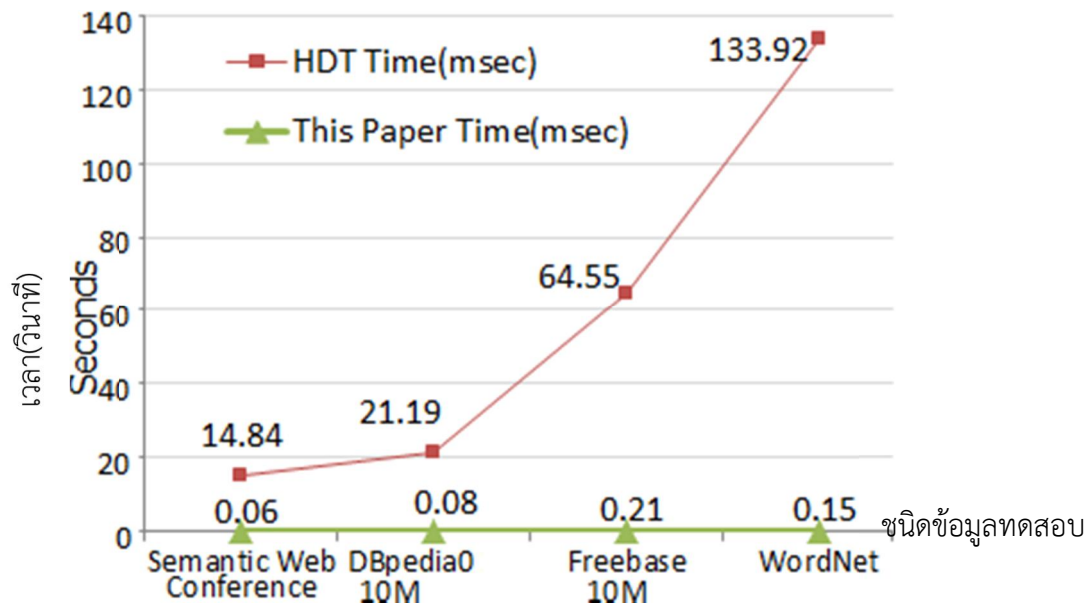


ภาพที่ 4.6 เปรียบเทียบจำนวนเทอม (Term) สำหรับตัวอย่างข้อมูลต่างๆ กันในแกน สำหรับวิธีแทนค่าข้อมูลแบบดั้งเดิม (RDF) แบบ HDT และแบบไบนารีในงานวิจัยนี้

¹³ <http://stackoverflow.com/questions/10862096/hadoop-map-reduce-how-to-speed-up-job-launch-setup>.

¹⁴ <http://blog.griddynamics.com/2010/07/war-story-optimizing-one-hadoop-job.html>.

คณะผู้วิจัยได้ทำการเปรียบเทียบจำนวนทอมในหน่วยหลักล้านของแฟ้มข้อมูลในรูปแบบ RDF รูปแบบ HDT และรูปแบบไบนารีที่นำเสนอ จะเห็นว่าในงานวิจัยนี้นำเสนอรูปแบบข้อมูลที่มีขนาดเล็กกว่าแบบดั้งเดิมและแบบ HDT ในทุกกรณี ดังภาพที่ 4.6



ภาพที่ 4.7 เปรียบเทียบเวลาในการขนย้ายข้อมูลจาก CPU ไปยัง GPU

จากนั้นทำการทดสอบวัดเวลาในการโอนย้ายข้อมูลจากหน่วยประมวลผลกลางไปหน่วยประมวลผลกราฟิกส์ เมื่อข้อมูลเป็นแบบไบนารีจะทำให้สามารถโอนย้ายข้อมูลได้เร็วกว่า ดังภาพ 4.7

ตารางที่ 4.3 และตารางที่ 4.4 แสดงเวลาที่ใช้ในการค้นหาข้อมูลในหน่วยประมวลผลกลาง และหน่วยประมวลผลกราฟิกตามลำดับ เมื่อข้อมูลอยู่ในรูปแบบ HDT ที่ได้กำหนดการสุ่ม ID ของ Subject, Predicate, Object จำนวน 1,024 ตัว และวัดเวลาเฉลี่ยของการค้นหาจำนวน 25 ครั้ง

ตารางที่ 4.5 แสดงเวลาที่ใช้ในการค้นหาข้อมูลในหน่วยประมวลผลกราฟิก เมื่อข้อมูลอยู่ในรูปแบบ SPO โดยใช้ข้อมูลต่อไปนี้ในการทดสอบ swdf, wn3.0, Freebase และ DBPedia ซึ่งแสดงขนาดข้อมูลเป็นจำนวนบิต แสดงจำนวนตำแหน่งของ Subject, Predicate และ Object รวมทั้งจำนวนตำแหน่งทั้งหมดของคำค้นหรือ pattern ที่มี รวมทั้งแสดงข้อมูลเวลาในมิลลิวินาที (ms) และ ข้อมูลของ Kernel ในหน่วยประมวลผลกราฟิกส์ ได้แก่ จำนวนบล็อกต่อกริด และจำนวนเทรตต่อบล็อกจากข้อมูลในตารางแสดงให้เห็นว่าวิธีการแบบไบนารี สามารถนำเข้าข้อมูลได้มากกว่า และค้นหาข้อมูลได้พร้อมกันหลายตำแหน่ง ทำให้ใช้เวลาน้อยกว่า อย่างไรก็ตามวิธีการใช้วิธีนี้จำเป็นต้องมีเวลาที่ต้องใช้ในการแปลงข้อมูลมาพิจารณาด้วย ภาพที่ 4.8 แสดงลักษณะข้อมูลแต่ละตัวอย่างทดสอบใน ตารางที่ 4.5

ตารางที่ 4.3 เวลาในการค้นหาข้อมูลด้วย HDT ในหน่วยประมวลผลกลาง

ชนิดข้อมูล	ขนาดข้อมูล ต้นฉบับ (MB)	จำนวนเซตของ patterns	เวลา (วินาที)
Freebase	109	1024	54
DBpedia	507	1024	35

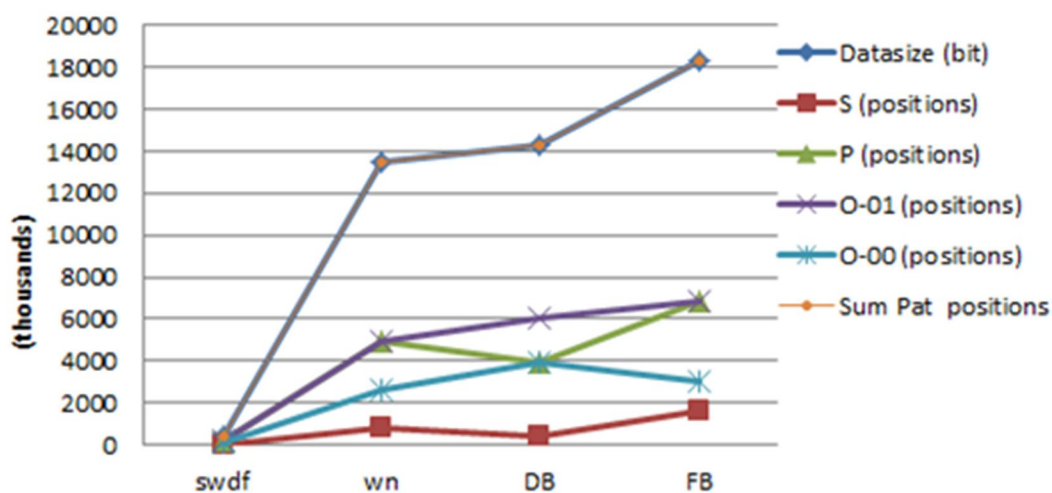
ตารางที่ 4.4 เวลาในการค้นหาข้อมูลในรูปแบบ HDT ในหน่วยประมวลผลกราฟิกส์

ชนิดข้อมูล	ขนาดข้อมูล (B)	Patterns	เวลา (วินาที)	จำนวนบล็อก	จำนวนเทรด
Freebase	1793413	1	7.16	588	1024
DBpedia	1626803	1	6.54	1751	1024

ตารางที่ 4.5 เวลาการค้นหา SPO ในหน่วยประมวลผลกราฟิกส์ (GTX560Ti)

HDT	Data size (bit)	S_11 (positions)	P_10 (positions)	O(00) (positions)	O(01) (positions)	Sum Pat positions	Times (ms)	blocks/ grid	threads block
swdf	433825	23282	168104	168111	73915	410,130	0.87	424	1024
Wn	13,462,923	819,505	4,983,833	4,983,783	2,662,655	13,449,776	26.68	13148	1024
Freebase	14,303,370	386,361	3,932,948	6,037,265	3,932,827	14,289,401	28.32	13969	1024
DBpedia	18,331,369	1,580,553	6,845,230	6,845,180	3,042,504	18,313,467	34.17	17902	1024

Datasize and total of number of pattern



ภาพที่ 4.8 เปรียบเทียบขนาดของข้อมูลและตำแหน่ง SPO ที่พบในแต่ละข้อมูล

4.3 การประเมินความพึงพอใจของผู้ใช้บริการเว็บเชิงความหมาย

ในการวิจัยได้ทำการประเมินความพึงพอใจของผู้ใช้งานระบบเว็บเชิงความหมาย กำหนดให้ ผู้ใช้งานระบบจำนวน 100 คนทำแบบสอบถามเพื่อประเมินประสิทธิภาพการทำงานของระบบใน 3 ด้านคือ

- ความพึงพอใจด้านคุณภาพของข้อมูล
- ความพึงพอใจด้านการออกแบบและการจัดรูปแบบ
- ความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับ

ผลการประเมินความพึงพอใจด้านคุณภาพของข้อมูลเป็นดังตาราง 4.6 จะเห็นได้ว่า ค่าเฉลี่ยของความพึงพอใจด้านคุณภาพของข้อมูลเท่ากับ 3.96 และมีค่าเฉลี่ยของส่วนเบี่ยงเบนมาตรฐานเท่ากับ 0.57 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพในด้านคุณภาพของข้อมูลอยู่ในระดับดี

ตารางที่ 4.6 ผลการประเมินความพึงพอใจด้านคุณภาพของข้อมูล

หัวข้อการประเมิน	ค่าเฉลี่ย (Mean)	ส่วนเบี่ยงเบนมาตรฐาน (SD)	ความหมาย
ความรวดเร็วในการเข้าถึงข้อมูล	3.95	0.53	ดี
การเชื่อมโยงข้อมูลภายในเว็บไซต์	4.13	0.65	ดี
ความครบถ้วนสมบูรณ์ของข้อมูลต่างๆ	3.56	0.76	ดี
ความครอบคลุมในการสืบค้นในเชิงกว้าง และสามารถสืบค้นข้อมูลแบบเจาะจงได้	4.07	0.66	ดี
ผลลัพธ์ที่ได้ตรงกับความต้องการ	4.07	0.25	ดี
ค่าเฉลี่ย	3.96	0.57	ดี

ผลการประเมินความพึงพอใจด้านการออกแบบและการจัดรูปแบบเป็นดังตาราง 4.7 การประเมินความพึงพอใจด้านการออกแบบและการจัดรูปแบบซึ่งได้ผลลัพธ์ว่าความพึงพอใจด้านนี้มีค่าเฉลี่ยเท่ากับ 4.08 และมีค่าส่วนเบี่ยงเบนมาตรฐานโดยเฉลี่ยเท่ากับ 0.54 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพในด้านการออกแบบและการจัดรูปแบบอยู่ในระดับดี

ตารางที่ 4.7 ผลการประเมินความพึงพอใจด้านการออกแบบและการจัดรูปแบบ

หัวข้อการประเมิน	ค่าเฉลี่ย (Mean)	ส่วนเบี่ยงเบนมาตรฐาน (SD)	ความหมาย
ความสวยงามและน่าสนใจของหน้าจอ	3.91	0.58	ดี
การจัดรูปแบบง่ายต่อการอ่านและการใช้งาน	3.88	0.55	ดี
ความง่ายต่อการใช้งานเมนูต่างๆ	4.47	0.56	ดี
ความเหมาะสมของสีสันทันและขนาดตัวหนังสือบนหน้าจอ	4.08	0.48	ดี
ค่าเฉลี่ย	4.08	0.54	ดี

ผลการประเมินความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับในตาราง 4.8 การประเมินความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับซึ่งได้ผลลัพธ์ว่าความพึงพอใจด้านนี้มีค่าเฉลี่ยเท่ากับ 3.90 และมีค่าส่วนเบี่ยงเบนมาตรฐานโดยเฉลี่ยเท่ากับ 0.54 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพในด้านประโยชน์ของข้อมูลที่ได้รับอยู่ในระดับดี

ตารางที่ 4.8 ผลการประเมินความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับ

หัวข้อการประเมิน	ค่าเฉลี่ย (Mean)	ส่วนเบี่ยงเบนมาตรฐาน (SD)	ความหมาย
ให้ข้อมูลสถานที่ให้บริการด้านสุขภาพและสปาที่น่าสนใจ	4.02	0.50	ดี
ให้ข้อมูลเกี่ยวกับสถานที่ใกล้เคียง	3.97	0.55	ดี
แสดงตำแหน่งที่ตั้งและเส้นทางการเดินทาง	4.03	0.57	ดี
ให้ข้อมูลเกี่ยวกับค่าใช้จ่ายโดยประมาณ	3.51	0.50	ดี
ช่วยในการตัดสินใจและการวางแผนการเดินทาง	3.96	0.60	ดี
ค่าเฉลี่ย	3.90	0.54	ดี

ตารางที่ 4.9 ผลการประเมินความพึงพอใจโดยรวม

หัวข้อการประเมิน	ค่าเฉลี่ย (Mean)	ส่วนเบี่ยงเบนมาตรฐาน (SD)	ความหมาย
ความพึงพอใจด้านคุณภาพของข้อมูล	3.97	0.65	ดี
ความพึงพอใจด้านการออกแบบและการจัดรูปแบบ	3.94	0.56	ดี
ความพึงพอใจด้านประโยชน์ของข้อมูลที่ได้รับ	3.97	0.47	ดี
ความพึงพอใจของระบบโดยรวม	3.94	0.44	ดี
ค่าเฉลี่ย	3.96	0.53	ดี

ผลการประเมินความพึงพอใจโดยรวมในตาราง 5.9 ซึ่งได้ผลลัพธ์ว่าความพึงพอใจโดยรวมมีค่าเฉลี่ยเท่ากับ 3.96 และมีค่าส่วนเบี่ยงเบนมาตรฐานโดยเฉลี่ยเท่ากับ 0.53 ดังนั้นระบบที่ได้พัฒนาขึ้นมีประสิทธิภาพโดยรวมอยู่ในระดับดี

บทที่ 5 สรุปผลการวิจัย อภิปรายผล และข้อเสนอแนะ

ในงานวิจัยนี้ คณะผู้วิจัยได้นำเสนอระบบบริหารจัดการข้อมูลท่องเที่ยวเพื่อสุขภาพ โดยในโครงการย่อยนี้ ได้ทำการสำรวจความต้องการของซอฟต์แวร์ปัจจุบัน รวมทั้งลักษณะโครงสร้างการทำงานของหน่วยงานเทศบาลเมือง พบว่าการค้นหาข้อมูลส่วนใหญ่เป็นการค้นหาจากเว็บไซต์ทั่วไปเช่น Agoda, TripAdvisor ผ่าน Google Search Engine เป็นต้น นอกจากนี้ข้อมูลอื่นๆ มาจากการเดินสำรวจ ซอฟต์แวร์ในการจัดเก็บข้อมูล และเก็บสถิติได้แก่ Excel และ Word เพื่อทำรายงานสรุปต่างๆ แหล่งข้อมูลต่างๆ ในการรวบรวมข้อมูลสถานประกอบการได้แก่ จากเว็บไซต์ Agoda, TripAdvisor, AtSiam, HeyHuaHin และจาก Official Website ของสถานบริการ รวมทั้งการเดินสำรวจในพื้นที่หัวหินและชะอำ ข้อมูลเหล่านี้ถูกนำเข้าออนโทโลยี เชื่อมต่อกับหน้าจอผู้ใช้งานที่ได้ออกแบบไว้ เพื่อสามารถค้นหาได้ สำหรับการประมวลผลการค้นหาเมื่อผู้ใช้ระบุเงื่อนไขการค้นหาจะถูกแปลงเป็นคิวรีส่งไปยังมิดเดิลแวร์เพื่อติดต่อกับออนโทโลยี และค้นหาข้อมูลที่ต้องการขึ้นมาแสดงผล

คณะผู้วิจัยการนำข้อมูลเข้าแบบขนานด้วยเทคโนโลยี Java แบบพร้อมกัน(Concurrent) และ MapReduce โดยพบว่าการนำเข้าข้อมูลโดยการค้นหาแบบพร้อมกัน ทำให้ทำงานได้เร็วขึ้น ประมาณ 3-4 เท่าสำหรับกรณีใช้ 4 เทรด และการทำงานแบบ MapReduce จะช่วยทำการรวมข้อมูลจากแฟ้มข้อมูลขนาดใหญ่เช่น 4 GB ได้ดีภายในเวลาไม่เกิน 7 วินาที เทียบกับวิธีปกติโดยโปรแกรมภาษา Java แบบ 1 เทรดที่ใช้เวลาร่วมสิบกว่านาที นอกจากนี้คณะผู้วิจัยได้ทำการวัดประสิทธิภาพการค้นหาจากออนโทโลยี โดยจะพบว่าเวลาเฉลี่ยในการค้นหาไม่เกิน 30 วินาที และทำการนำเสนอโมเดลการค้นหาด้วยวิธีการแบบขนานในโมเดลต้นไม้แบบไบนารีกับหน่วยประมวลผลกราฟิกส์เพื่อประมวลผลแบบขนาน โดยนำข้อมูลออนโทโลยีในรูปแบบ Triple เข้ามาในหน่วยประมวลผลกราฟิกส์เพื่อค้นหาแบบพร้อมกัน ซึ่งสามารถนำเข้าข้อมูลได้มากขึ้นและนำมาค้นหาได้เร็วขึ้นนั่นเอง

คณะผู้วิจัยได้ประเมินความพึงพอใจจากผู้ใช้งาน 200 คนเพื่อวัดความพึงพอใจของข้อมูลและการออกแบบหน้าเว็บอีกด้วย การออกแบบระบบด้วยออนโทโลยีให้ข้อดี ในการค้นหาข้อมูลต่างๆ ได้ตรงตามความต้องการของผู้ใช้มากขึ้น เนื่องจากการค้นหาตามความหมายของคำค้นในบริบทที่ผู้ใช้ต้องการจริง อย่างไรก็ตามเครื่องมือที่สนับสนุนการออกแบบออนโทโลียังมีขีดจำกัดในด้านต่างๆ เช่น กลไกการอนุมาน ลอจิกที่รองรับ ความเสถียรของระบบเชื่อมต่อ และการทำงานที่ช้ากว่าระบบฐานข้อมูลเชิงสัมพันธ์ที่มีอยู่ ดังนั้นการพัฒนาอย่างต่อเนื่องของเทคโนโลยีเหล่านี้อาจจะทำให้ระบบค้นหาเชิงความหมายสมบูรณ์ขึ้น

นอกจากนี้การรองรับการประมวลผลแบบขนานของเครื่องมือซอฟต์แวร์ที่มีอยู่ยังมีขีดจำกัด ทำให้เมื่อต้องการใช้กลไกการค้นหาแบบขนานมาช่วย ต้องทำการออกแบบการเชื่อมต่อ รวมทั้งเวลาการประมวลผลขึ้นก่อนการค้นหาหรือ Preprocess ต่างๆ ค่อนข้างมากซึ่งเป็น Overhead ในการทำงานโดยรวมอยู่ ตาราง 5.1 สรุปกิจกรรมต่างๆ ที่ได้ดำเนินการในโครงการนี้

ตารางที่ 5.1 ตารางเปรียบเทียบวัตถุประสงค์กิจกรรมที่วางแผนไว้และกิจกรรมที่ดำเนินการมา และผลที่ได้รับของ โครงการย่อย 1

กิจกรรมที่ได้วางแผนไว้	ผลที่ได้รับ
1. สำรวจความต้องการซอฟต์แวร์ในปัจจุบัน	รายงานผลสำรวจความต้องการซอฟต์แวร์ จากผู้ใช้งานในหน่วยงาน (บทที่ 4 หัวข้อ 4.1) และภาคผนวก ก ค ง และ จ ของแผนงาน
2. การศึกษางานวิจัยที่เกี่ยวข้องด้านการประมวลแบบขนานบนกลุ่มเมฆ	รายงานงานวิจัยที่เกี่ยวข้อง (บทที่ 2)
3. พัฒนาตัวแบบของระบบการประมวลบนกลุ่มเมฆ	ตัวแบบการประมวลบนกลุ่มเมฆ (บทที่ 3 หัวข้อ 3.5)
4. พัฒนาอัลกอริทึมการประมวลคิวรีแบบขนาน และการแทนค่าข้อมูล	อัลกอริทึมแบบขนาน สำหรับการประมวลคิวรีรูปแบบการแทนค่าข้อมูล (บทที่ 3 หัวข้อ 3.1, 3.2 และภาคผนวก ก)
5. พัฒนาโปรแกรมการประมวลผลคิวรีแบบขนาน	โปรแกรมการประมวลผลคิวรีแบบขนาน
6. การออกแบบหน้าจอ การติดต่อผู้ใช้ ระบบออนไลน์	แบบหน้าจอ การติดต่อผู้ใช้ (บทที่ 3 หัวข้อ 3.3)
7. ออกแบบฐานข้อมูลระบบ ได้แก่ ข้อมูลผู้ใช้ระบบ ข้อมูลสถานที่ท่องเที่ยว พัฒนาระบบการค้นหา พัฒนาโปรแกรมทดลองใช้	ฐานข้อมูลผู้ใช้ระบบและข้อมูลเกี่ยวกับการท่องเที่ยว ระบบค้นหา และการจัดการ ในรายงานแผนงาน ภาคผนวก ซ. และคู่มือนำเข้าข้อมูล ในรายงานโครงการย่อย 1 ภาคผนวก ค
8. เปรียบเทียบผลการค้นหาคิวรี ด้านประสิทธิภาพเวลา	ผลการประเมินด้านประสิทธิภาพ (บทที่ 4 หัวข้อ 4.2 และภาคผนวก ข)
9. ประเมินผลการใช้งานกับผู้ใช้งานหลายๆ คน	ผลการประเมินผู้ใช้งาน (บทที่ 4 หัวข้อ 4.3)

กิจกรรมที่ได้วางแผนไว้	ผลที่ได้รับ
10. ประเมินค่าใช้จ่ายถ้าจะนำระบบไปใช้จริง	ประมาณการค่าใช้จ่าย ในรายงานแผนภาคผนวก ฉ. และ ช.
11. จัดเตรียมการอบรมผู้ใช้ จัดทำคู่มือผู้ใช้งาน	การอบรมผู้ใช้งาน คู่มือการใช้งาน ในรายงานแผน ภาคผนวก ฉ.
12. เขียนสรุปรายงานและเอกสารวิชาการ	สรุปรายงานการวิจัย และเอกสารวิชาการในภาคผนวก ฉ.

5.1 ข้อเสนอแนะ

ในงานวิจัยในอนาคต การพัฒนาโครงสร้างข้อมูลการแทนค่าออนโทโลยีพื้นฐานนี้ให้เหมาะกับการประมวลผลแบบขนานไม่ว่าจะด้วยลักษณะคลัสเตอร์ หรือหน่วยประมวลผลกราฟิกส์ก็ตาม โลบาร์ต่างๆ ที่รองรับจะมีส่วนสำคัญในการเชื่อมต่อการประมวลผลออนโทโลยีกับเว็บเชิงความหมายให้เป็นอัตโนมัติ และด้วยประสิทธิภาพที่ดีต่อไป เมื่อพัฒนาการทางฮาร์ดแวร์ของหน่วยประมวลผลแบบหลายแกนดีขึ้น มีราคาถูกลง การซอฟต์แวร์รองรับการพัฒนาโปรแกรมแบบหลายประมวลผลหลายแกนมากขึ้น รวมทั้งชุดพัฒนาต่างๆ จะทำให้การพัฒนาการค้นหาข้อมูลเชิงความหมายเป็นไปได้ง่ายมากขึ้น

อย่างไรก็ตามองค์ความรู้ในการพัฒนางานประยุกต์ที่ได้นี้ ก็จะมีส่วนสำคัญในการสร้างนักวิจัยแบบทางคอมพิวเตอร์ในการพัฒนางานวิจัยที่รองรับการทำงานแบบพร้อมกันหรือบนสถาปัตยกรรมแบบหลายแกนต่อไปในอนาคตต่อไป

เอกสารอ้างอิง

- [1] Bernstein, D., and Vij,D., “Using semantic web ontology for Inter cloud directories and exchanges,” In *Proceedings of ICOMP'10, the 11th International Conference on Internet Computing*, 2010.
- [2] Benedikt, M., and Teuteberg, F., “Risk and compliance management for cloud Computing services: designing a reference model,” in *Proceedings of AMCIS*, 2011.
- [3] Ming, X. Z., “A semantic grid oriented to etourism”, in *Proceedings of the 1st International Conference on Cloud Computing*. 2009.
- [4] Lee, C.G., “Health care and tourism: evidence from Singapore,” *Tourism Management*, Vol. 31, No. 4, pp. 486-488, 2010.
- [5] Prantner, K., Ding, Y., Luger, M., Yan, Z., and Herzog, C., “Tourism ontology and semantic management system: State-of-the-arts Analysis”, In: *IADIS International Conference WWW/Internet 2007*. Vila Real, Portugal, 2007.
- [6] McCreddie, R., Mcdonald, C., and Ounis, I., “Comparing distributed indexing: to MapReduce or not?,” in *Proceedings of the 7th Workshop on Large-Scale Distributed Systems for Information Retrieval*, 2009.
- [7] Lin, J., and Dyer, C., "Data-Intensive Text Processing with MapReduce" , pre-production manuscript of a book in the Morgan & Claypool Synthesis Lectures on Human Language Technologies. Prepared April 3, 2012.
- [8] Wei, Z. and JaJa, J. , “A fast algorithm for constructing inverted files on heterogeneous platforms,” *Journal of Parallel and Distributed Computing*, vol. 72, no. 5, pp. 728-738, 2011.
- [9] Ming, L., et al., “Authorized private keyword search over encrypted data in cloud computing,” in *Proceedings of the 31st International Conference on Distributed Computing Systems, IEEE Computer Society*, 2011.
- [10] Kyong-Ha, L., et al., “Parallel data processing with MapReduce: a survey,” *SIGMOD Rec.*, Vol. 40, no. 4, pp. 11-20, 2011.
- [11] Lin, J., Metzler, D., Elsayed, T., Wang, L., “Of ivory and smurfs: Loxodontan MapReduce experiments for web search,” in *Proceedings of the Eighteenth Text REtrieval Conference, TREC*, 2009.

- [12] Williams, G. T., Weaver, J., et al., "Scalable reduction of large datasets to interesting subsets," *Web Semantics: Science, Services and Agents on the World Wide Web*, vol. 8, no. 4, pp. 365-373, 2010.
- [13] Urbani, J., et al., "WebPIE: A Web-scale parallel inference engine using MapReduce," *Web Semantics: Science, Services and Agents on the World Wide Web*, vol. 10, pp. 59-75, 2012.
- [14] Mohammad, F.H., McGlothlin, J., Masud, M. M., Khan, L.R., and Thuraisingham, B., "Heuristics-based query processing for large RDF graphs using cloud computing," *IEEE Transactions on Knowledge and Data Engineering*, vol. 23, no. 9, pp. 1312-1327, Sept. 2011, doi:10.1109/TKDE.2011.103.
- [15] Rohloff, K., and Richard, E.S., "High-performance, massively scalable distributed systems using the MapReduce software framework: the SHARD triple-store," in *Programming Support Innovations for Emerging Distributed Applications*, ACM: Reno, Nevada, 2010.
- [16] Hartley, T.D.R., Saule, E., and Çatalyürek, Üm.V., "Improving performance of adaptive component-based dataflow middleware," *Parallel Computing*, Vol. 38, no. 6-7, pp. 289-309, 2012.
- [17] Soma, R., and Prasanna, V.K., "Parallel inferencing for OWL knowledge bases", in *Proceedings of the 2008 37th International Conference on Parallel Processing*, IEEE Computer Society, Washington DC, USA, 2008.
- [18] Nicholas, C. and V.L. Narasimhan, "A Novel Data Distribution Technique for Host-Client Type Parallel Applications," *IEEE Trans. Parallel Distrib. Syst.*, Vol. 13, No. 2, pp. 97-110, 2002.
- [19] Yang, H.-c., Dasdan, A., Hsiao, R.-L., and Parker, D. S, "Map-Reduce-Merge: simplified relational data processing on large clusters," in *Proceedings of the 2007 ACM SIGMOD international conference on Management of data*, pp. 1029-1040, 2007.
- [20] Hiemstra, D. and Hauff, C., "MapReduce for Experimental Search," in *Proceedings of the Nineteenth Text REtrieval Conference (TREC 2010)*, 2010.
- [21] Dean, J., and Ghemawat, S., "MapReduce: simplified data processing on large clusters", *Communications of the ACM*, Vol. 50, No. 1, pp. 137-149, 2008.
- [22] Fukunaga, A., Kishimoto, A., Botea, A., "Iterative Resource Allocation for Memory Intensive Parallel Search Algorithms on Clouds", *Grids, and Shared Clusters*, 2012.

- [23] Rohloff, K., and Schantz, R. E. "Clause-iteration with MapReduce to scalably query datagraphs in the SHARD graph-store," In *Proceedings of the fourth international workshop on Data-intensive distributed computing (DIDC '11)*, pp. 35-44, 2011. DOI=10.1145/1996014.1996021.
- [24] Lin, J., Metzler, D., Elsayed, T., and Wang, L., "Of ivory and smurfs: Loxodontan MapReduce experiments for web search," In *Proceedings of the Eighteenth Text REtrieval Conference (TREC 2009)*, 2009.
- [25] Byun, C., Arcand, W., Bestor, D., Bergeron, D., Hubbell, M., Kepner, J., and Yee, C., "Driving big data with big compute," [Online]. Available: http://www.mit.edu/~kepner/pubs/ByunKepner_2012_BigData_Paper.pdf.
- [26] Rohloff, K., Cloud computing for Scalability:The SHARD Triple-Store, [Online]. Available: https://dist-systems.bbn.com/people/krohloff/papers/2011/Rohloff_Meetup_01_09_2011.pdf
- [27] Apache Hadoop คืออะไร. [Online]. Available: <http://www.wegethosting.com/blog/Apache-hadoop-คืออะไร>
- [28] MacLean, D., "A very brief introduction to MapReduce.", Lecture note, http://hci.stanford.edu/courses/cs448g/a2/files/map_reduce_tutorial.pdf.
- [29] Jain, S., "Hadoop architecture-types of Hadoop nodes in cluster-Part1," [Online]. Available: <http://www.haaram.com/Compleetearicle.aspx?aid=505847&ln=>
- [30] Jain, S., "Hadoop Architecture-Types of Hadoop Nodes in Cluster-Part2," [Online]. Available: <http://www.haaram.com/Compleetearicle.aspx?aid=505847&ln=>
- [31] Vaidya, M., "Parallel Processing of cluster by Map Reduce," *International Journal of Distributed and Parallel Systems (IJDPS)*, vol. 3, no. 1, pp. 167-197, 2012.
- [32] Lea, D., "A Java Fork/Join framework." State University of New York at Oswego, Proceedings of ACM 2000 conference on Java Grande (JAVA'00), pp. 36-43, 2000.
- [33] Fernández, J. D., Martínez-Prieto, M. A., Gutiérrez, C., Polleres, A., and Arias, M., "Binary RDF representation for publication and exchange (HDT)." *Web Semant*, vol. 19, no. 22-41, 2013.

ภาคผนวก โครงการย่อย 1

ภาคผนวก ก. การติดตั้ง Hadoop และการทดสอบประสิทธิภาพ MapReduce

ภาคผนวก ข. วิธีที่ใช้ทดสอบประสิทธิภาพด้านเวลา

การติดตั้ง hadoop บนระบบปฏิบัติการลินุกซ์ Ubuntu แบบ 1 โหนด

1. ติดตั้ง java jdk 7 ลงในเครื่อง

1.1 ดาวน์โหลดจาวาที่มีสกุลไฟล์เป็น .tar.gz* จากนั้นทำการแตกไฟล์ผ่าน terminal โดยจะต้องเข้าไปอยู่ในไดเรกทอรีที่เก็บไฟล์จาวาไว้แล้วใช้คำสั่ง

```
> cd Downloads (เข้าไปในไดเรกทอรีที่เก็บไฟล์จาวา)
```

```
> tar -xvf jdk-7u45*-linux-i586.tar.gz (32 bit)
```

```
> tar -xvf jdk-7u45-linux-x64.tar.gz (64 bit)
```

* ดาวน์โหลดไฟล์จาวาได้จาก <http://www.oracle.com/technetwork/java/javase/downloads/jdk7-downloads-1880260.html>

** 7u45 คือ เวอร์ชันของจาวา ซึ่งจะมีเวอร์ชันใหม่กว่าอีก

1.2 เมื่อแตกไฟล์แล้วจะได้โฟลเดอร์ jdk1.7.0_45 ให้สร้างโฟลเดอร์เพื่อเก็บจาวาโดยเฉพาะ

```
> mkdir -p /usr/lib/jvm (สร้างโฟลเดอร์เองจะไว้ที่ใดก็ได้)
```

```
> sudo mv /home/user/Downloads/jdk1.7.0_45 /usr/lib/jvm/jdk1.7.0 (ย้ายไปโฟลเดอร์ที่สร้างไว้)
```

รูปแบบการใช้คำสั่งย้ายโฟลเดอร์ คือ `sudo mv [ไดเรกทอรีต้นทางมีไฟล์ที่แตกแล้ว] [ไดเรกทอรีปลายทาง]`

1.3 กำหนดลำดับความสำคัญของ JDK โดยรันคำสั่งต่อไปนี้

```
> sudo update-alternatives --install "/usr/bin/java" "java"
```

```
"/usr/lib/jvm/jdk1.7.0/bin/java" 1
```

```
> sudo update-alternatives --install "/usr/bin/javac" "javac"
```

```
"/usr/lib/jvm/jdk1.7.0/bin/javac" 1
```

```
> sudo update-alternatives --install "/usr/bin/javaws" "javaws"
```

```
"/usr/lib/jvm/jdk1.7.0/bin/javaws" 1
```

1.4 แก้ไขความเป็นเจ้าของไฟล์และสิทธิ์ของการปฏิบัติงาน โดยรันคำสั่งต่อไปนี้

```
> sudo chmod a+x /usr/bin/java
```

```
> sudo chmod a+x /usr/bin/javac
```

```
> sudo chmod a+x /usr/bin/javaws
```

```
> sudo chown -R user:user /usr/lib/jvm/jdk1.7.0
```

* java JDK มี executables อื่นๆ อีกมากมายที่สามารถติดตั้งในทำนองเดียวกันกับข้างต้น ซึ่ง java, javac, javaws อาจจะจำเป็นต้องใช้บ่อยที่สุด

```

hduser@hduser-System-Product-Name:~/Downloads$ cd
hduser@hduser-System-Product-Name:~$ cd Downloads
hduser@hduser-System-Product-Name:~/Downloads$ sudo mv jdk1.7.0_45 /usr/lib/jvm/
jdk1.7.0/
hduser@hduser-System-Product-Name:~/Downloads$ cd
hduser@hduser-System-Product-Name:~$ sudo update-alternatives --install "/usr/bin/java" "java"
"/usr/lib/jvm/jdk1.7.0/jdk1.7.0_45/bin/java" 1
update-alternatives: using /usr/lib/jvm/jdk1.7.0/jdk1.7.0_45/bin/java to provide /usr/bin/java
(java) in auto mode.
hduser@hduser-System-Product-Name:~$ sudo update-alternatives --install "/usr/bin/javac" "javac"
"/usr/lib/jvm/jdk1.7.0/jdk1.7.0_45/bin/javac" 1
update-alternatives: using /usr/lib/jvm/jdk1.7.0/jdk1.7.0_45/bin/javac to provide /usr/bin/javac
(javac) in auto mode.
hduser@hduser-System-Product-Name:~$ sudo update-alternatives --install "/usr/bin/javaws" "javaws"
"/usr/lib/jvm/jdk1.7.0/jdk1.7.0_45/bin/javaws" 1
update-alternatives: using /usr/lib/jvm/jdk1.7.0/jdk1.7.0_45/bin/javaws to provide /usr/bin/javaw
s (javaws) in auto mode.
hduser@hduser-System-Product-Name:~$ mkdir ~/.mozilla/plugins/
hduser@hduser-System-Product-Name:~$ ln -s /usr/lib/jvm/jdk1.7.0/jdk1.7.0_45/jre/lib/amd64/libnpjp2.so
~/.mozilla/plugins/
hduser@hduser-System-Product-Name:~$ java -version
java version "1.7.0_45"
Java(TM) SE Runtime Environment (build 1.7.0_45-b18)
Java HotSpot(TM) 64-Bit Server VM (build 24.45-b08, mixed mode)
hduser@hduser-System-Product-Name:~$ █

```

1.5 จากนั้นรันคำสั่ง

> *sudo update-alternatives --config java*

คำสั่งนี้จะปรากฏลำดับเลข 0, 1, 2 และพาธให้เลือก โดยเลือกเลขที่เป็นพาธของ
/usr/lib/jvm/jdk1.7.0/bin/java (พิมพ์ตัวเลขลงใน terminal)

> *sudo update-alternatives --config javac*

เลือกเลขที่เป็นพาธของ /usr/lib/jvm/jdk1.7.0/bin/javac

> *sudo update-alternatives --config javaws*

เลือกเลขที่เป็นพาธของ /usr/lib/jvm/jdk1.7.0/bin/javaws

1.6 ตรวจสอบการติดตั้งจาวาด้วยคำสั่ง

> *java -version*

*** ขั้นตอนการติดตั้งจาวาอ้างอิงจาก <http://askubuntu.com/questions/56104/how-can-i-install-sun-oracles-proprietary-java-jdk-6-7-8-or-jre>

2. Configuring SSH

Hadoop ร้องขอ SSH เพื่อใช้ในการเข้าถึงโหนดต่างๆ เช่น การทำงานแบบ remote จากเครื่องแม่ (master) ไปเครื่องลูก (slave)

2.1 ติดตั้ง SSH

> *sudo apt-get install openssh-server*

2.2 สร้างคู่คีย์ RSA ด้วยรหัสผ่านที่ว่างเปล่า เพราะไม่ต้องป้อนรหัสผ่านทุกครั้งที่ทำการโต้ตอบกับแต่ละโหนด

```
> ssh-keygen -t rsa -P "" (จากนั้นกด Enter)
```

```
> cat $HOME/.ssh/id_rsa.pub >> $HOME/.ssh/authorized_keys
```

3. Disable IPv6

ปิดการใช้งานไอฟีร์วัน 6 เพราะ ubuntu ใช้ IP 0.0.0.0 สำหรับการกำหนดค่าที่แตกต่างกันของ Hadoop

3.1 เปิดไฟล์ sysctl.conf ด้วย Text Editor

```
> sudo gedit /etc/sysctl.conf
```

3.2 เพิ่มข้อความต่อไปนี้ต่อท้ายไฟล์

```
#disable ipv6

net.ipv6.conf.all.disable_ipv6 = 1
net.ipv6.conf.default.disable_ipv6 = 1
net.ipv6.conf.lo.disable_ipv6 = 1

export HADOOP_OPTS=-Djava.net.preferIPv4Stack=true
```

จากนั้นบันทึกและปิดไฟล์ หากติดปัญหา permission ให้ใช้คำสั่งเดียวกันนี้ใน root account

3.3 กำหนดการตั้งค่าเริ่มต้นอีกครั้ง โดยใช้คำสั่ง

```
> sudo sysctl -p
```

3.4 ตรวจสอบการปิด IPv6 ด้วยคำสั่ง

```
> cat /proc/sys/net/ipv6/conf/all/disable_ipv6
```

หากได้ค่า 1 แสดงว่าทำการปิด IPv6 แล้ว

```
hduser@hduser-System-Product-Name:~$ sudo gedit /etc/sysctl.conf
hduser@hduser-System-Product-Name:~$ sudo sysctl -p
net.ipv6.conf.all.disable_ipv6 = 1
net.ipv6.conf.default.disable_ipv6 = 1
net.ipv6.conf.lo.disable_ipv6 = 1
hduser@hduser-System-Product-Name:~$ cat /proc/sys/net/ipv6/conf/all/disable_ipv6
6
1
```

4. ติดตั้ง Hadoop

4.1 ดาวน์โหลด Hadoop ที่มีสกุลไฟล์ .tar.gz ในคู่มือนี้ดาวน์โหลดเวอร์ชัน 0.20.2

4.2 แยกไฟล์โดยเข้าไปในไดเรกทอรีที่เก็บไฟล์ hadoop ไว้ก่อน ในคู่มือนี้เก็บไว้ใน /home

```
> sudo tar xzf hadoop-0.20.2.tar.gz
```

4.3 เมื่อแตกไฟล์แล้วจะได้โฟลเดอร์ hadoop-0.20.2

```
> sudo mv hadoop-0.20.2 hadoop (เปลี่ยนชื่อโฟลเดอร์เป็น hadoop)
```

4.4 Update \$HOME/.bashrc

```
> sudo gedit /home/hduser*/.bashrc (เปิดไฟล์, *hduser คือ ชื่อผู้ใช้ที่ติดตั้ง hadoop)
```

4.5 เพิ่มข้อความต่อไปนี้ต่อท้ายไฟล์

```
# Set Hadoop-related environment variables
```

```
export HADOOP_HOME=/home/hduser/hadoop*
```

```
# Set JAVA_HOME (we will also configure JAVA_HOME directly for Hadoop later on)
```

```
export JAVA_HOME=/usr/lib/jvm/jdk1.7.0**
```

```
# Some convenient aliases and functions for running Hadoop-related commands
```

```
unalias fs &> /dev/null
```

```
alias fs="hadoop fs"
```

```
unalias hls &> /dev/null
```

```
alias hls="fs -ls"
```

```
# If you have LZOP compression enabled in your Hadoop cluster and
```

```
# compress job outputs with LZOP (not covered in this tutorial):
```

```
# Conveniently inspect an LZOP compressed file from the command
```

```
# line; run via:
```

```
#
```

```
# $ lzohead /hdfs/path/to/lzop/compressed/file.lzo
```

```
#
```

```
# Requires installed 'lzop' command.
```

```
#
```

```
lzohead () {
```

```
    hadoop fs -cat $1 | lzop -dc | head -1000 | less
```

```
}
```

```
# Add Hadoop bin/ directory to PATH
export PATH=$PATH:$HADOOP_HOME/bin
```

* ไดรกทอรีที่เก็บไฟล์เดอร์ในข้อ 4.3

** JAVA_HOME จากข้อ 1.2

4.6 แก้ไข JAVA_HOME

```
> sudo gedit /home/hduser/hadoop/conf/hadoop-env.sh (เปิดไฟล์)
```

แก้ไข # export JAVA_HOME=/usr/lib/j2sdk1.5-sun เป็น

```
export JAVA_HOME=/usr/lib/jvm/jdk1.7.0 (ลบเครื่องหมาย # แล้วเปลี่ยนเป็นไดเรกทอรีที่เก็บจาวา)
```

4.7 สร้างโฟลเดอร์สำหรับเก็บไฟล์ข้อมูลและการทำงานต่างๆ ของ hadoop

```
> sudo mkdir /home/hduser/tmp (สร้างชื่อใดก็ได้ ในที่นี้สร้าง tmp เก็บไว้ที่ /home)
```

```
> sudo chown hduser:hduser /home/hduser/tmp (แก้ไขสิทธิ์)
```

```
> sudo chmod 750 /home/hduser/tmp (เพิ่มความปลอดภัย)
```

4.8 แก้ไขไฟล์ **core-site.xml**

```
> sudo gedit /home/hduser/hadoop/conf/core-site.xml (เปิดไฟล์)
```

แล้วเติมข้อความต่อไปนี้ระหว่าง **<configuration>** .. **</configuration>**

```
<!-- In: conf/core-site.xml -->
```

```
<property>
```

```
<name>hadoop.tmp.dir</name>
```

```
<value>/home/hduser/tmp*</value>
```

```
<description>A base for other temporary directories.</description>
```

```
</property>
```

```
<property>
```

```
<name>fs.default.name</name>
```

```
<value>hdfs://localhost:54310**</value>
```

```
<description>The name of the default file system. A URI whose
scheme and authority determine the FileSystem implementation. The
uri's scheme determines the config property (fs.SCHEME.impl) naming
```

the FileSystem implementation class. The uri's authority is used to determine the host, port, etc. for a filesystem.</description>

</property>

* ไดรกทอรี่ในข้อ 4.7

** กำหนด file system ในพื้นที่ทำในเครื่องตัวเอง (1 โหนด)

4.9 แก้ไขไฟล์ **mapred-site.xml**

> `sudo gedit /home/hduser/hadoop/conf/mapred-site.xml` (เปิดไฟล์)

แล้วเติมข้อความต่อไปนี้ระหว่าง <configuration> .. </configuration>

<!-- In: conf/mapred-site.xml -->

<property>

<name>mapred.job.tracker</name>

<value>localhost:54311</value>

<description>The host and port that the MapReduce job tracker runs at. If "local", then jobs are run in-process as a single map and reduce task.

</description>

</property>

4.10 แก้ไขไฟล์ **hdfs-site.xml**

> `sudo gedit /home/hduser/hadoop/conf/hdfs-site.xml` (เปิดไฟล์)

แล้วเติมข้อความต่อไปนี้ระหว่าง <configuration> .. </configuration>

<!-- In: conf/hdfs-site.xml -->

<property>

<name>dfs.replication</name>

<value>1*</value> * จำนวนโหนด (DataNode)

<description>Default block replication.

The actual number of replications can be specified when the file is created.

The default is used if replication is not specified in create time.

</description>

</property>

5. Running Hadoop

5.1 พอร์แมนเนทโหมดโดยใช้คำสั่ง

> `cd Hadoop` (เข้าไปใน HADOOP_HOME)

> `bin/hadoop namenode -format`

หากปรากฏคำถามว่าจะ re-format filesystem หรือไม่ ให้ตอบ Y

```
suphaksa@Suphaksa: ~/hadoop
suphaksa@Suphaksa:~$ rm -rf /home/suphaksa/tmp/dfs/data
suphaksa@Suphaksa:~$ cd hadoop
suphaksa@Suphaksa:~/hadoop$ bin/hadoop namenode -format
14/02/06 09:43:48 INFO namenode.NameNode: STARTUP_MSG:
/*****
STARTUP_MSG: Starting NameNode
STARTUP_MSG:   host = Suphaksa/127.0.1.1
STARTUP_MSG:   args = [-format]
STARTUP_MSG:   version = 0.20.2
STARTUP_MSG:   build = https://svn.apache.org/repos/asf/hadoop/common/branches/b
ranch-0.20 -r 911707; compiled by 'chrisdo' on Fri Feb 19 08:07:34 UTC 2010
*****/
Re-format filesystem in /home/suphaksa/tmp/dfs/name ? (Y or N) Y
14/02/06 09:43:53 INFO namenode.FSNamesystem: fsOwner=suphaksa,suphaksa,adm,cdro
m,sudo,dip,plugdev,lpadmin,sambashare
14/02/06 09:43:53 INFO namenode.FSNamesystem: supergroup=supergroup
14/02/06 09:43:53 INFO namenode.FSNamesystem: isPermissionEnabled=true
14/02/06 09:43:53 INFO common.Storage: Image file of size 98 saved in 0 seconds.
14/02/06 09:43:53 INFO common.Storage: Storage directory /home/suphaksa/tmp/dfs/
name has been successfully formatted.
14/02/06 09:43:53 INFO namenode.NameNode: SHUTDOWN_MSG:
/*****
SHUTDOWN_MSG: Shutting down NameNode at Suphaksa/127.0.1.1
*****/
suphaksa@Suphaksa:~/hadoop$
```

สังเกตการสำเร็จด้วยข้อความ successfully formatted

5.2 สตาร์ทงานทุกประเภทด้วยคำสั่ง

> `bin/start-all.sh`

```
suphaksa@Suphaksa:~/hadoop$ bin/start-all.sh
starting namenode, logging to /home/suphaksa/hadoop/bin/../logs/hadoop-suphaksa-namenode-Suphaksa.out
localhost: starting datanode, logging to /home/suphaksa/hadoop/bin/../logs/hadoop-suphaksa-datanode-Suphaksa.out
localhost: starting secondarynamenode, logging to /home/suphaksa/hadoop/bin/../logs/hadoop-suphaksa-secondarynamenode-Suphaksa.out
starting jobtracker, logging to /home/suphaksa/hadoop/bin/../logs/hadoop-suphaksa-jobtracker-Suphaksa.out
localhost: starting tasktracker, logging to /home/suphaksa/hadoop/bin/../logs/hadoop-suphaksa-tasktracker-Suphaksa.out
suphaksa@Suphaksa:~/hadoop$ jps
3705 NameNode
4197 SecondaryNameNode
2678 org.eclipse.equinox.launcher_1.3.0.v20130327-1440.jar
3966 DataNode
4548 Jps
4279 JobTracker
4500 TaskTracker
suphaksa@Suphaksa:~/hadoop$
```

สังเกตความสำเร็จโดยใช้คำสั่ง

```
> jps
```

ปรากฏงานที่สตาร์ทแล้ว ประกอบด้วย Namenode, Datanode, Jobtracker, Tasktracker และ Secondarynamenode

กรณีเกิดปัญหาสตาร์ท DataNode ไม่ได้ ให้แก้ไขโดยการสตาร์ททีละงานเอง โดยมีขั้นตอนดังนี้

```
> rm -rf /home/hduser/tmp/dfs/data
```

```
> cd Hadoop
```

```
> bin/hadoop namenode -format
```

```
> bin/hadoop-daemon.sh start namenode
```

```
> bin/hadoop-daemon.sh start datanode
```

```
> bin/hadoop-daemon.sh start secondarynamenode
```

```
> bin/hadoop-daemon.sh start tasktracker
```

```
> bin/hadoop-daemon.sh start jobtracker
```

```

hduser@hduser-System-Product-Name:~$ rm -rf /home/hduser/tmp/dfs/data
hduser@hduser-System-Product-Name:~$ cd hadoop
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop namenode -format
13/12/06 08:25:36 INFO namenode.NameNode: STARTUP_MSG:
/*****
STARTUP_MSG: Starting NameNode
STARTUP_MSG:   host = hduser-System-Product-Name/127.0.1.1
STARTUP_MSG:   args = [-format]
STARTUP_MSG:   version = 0.20.2
STARTUP_MSG:   build = https://svn.apache.org/repos/asf/hadoop/common/branches/b
ranch-0.20 -r 911707; compiled by 'chrisdo' on Fri Feb 19 08:07:34 UTC 2010
*****/
Re-format filesystem in /home/hduser/tmp/dfs/name ? (Y or N) Y
13/12/06 08:25:39 INFO namenode.FSNamesystem: fsOwner=hduser,hduser,adm,cdrom,su
do,dip,plugdev,lpadmin,sambashare
13/12/06 08:25:39 INFO namenode.FSNamesystem: supergroup=supergroup
13/12/06 08:25:39 INFO namenode.FSNamesystem: isPermissionEnabled=true
13/12/06 08:25:39 INFO common.Storage: Image file of size 96 saved in 0 seconds.
13/12/06 08:25:39 INFO common.Storage: Storage directory /home/hduser/tmp/dfs/na
me has been successfully formatted.
13/12/06 08:25:39 INFO namenode.NameNode: SHUTDOWN_MSG:
/*****
SHUTDOWN_MSG: Shutting down NameNode at hduser-System-Product-Name/127.0.1.1
*****/
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop-daemon.sh start namenode
starting namenode, logging to /home/hduser/hadoop/bin/../logs/hadoop-hduser-name
node-hduser-System-Product-Name.out
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop-daemon.sh start jobtracke
r
starting jobtracker, logging to /home/hduser/hadoop/bin/../logs/hadoop-hduser-jo
btracker-hduser-System-Product-Name.out
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop-daemon.sh start tasktrack
er
starting tasktracker, logging to /home/hduser/hadoop/bin/../logs/hadoop-hduser-t
asktracker-hduser-System-Product-Name.out
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop-daemon.sh start datanode
starting datanode, logging to /home/hduser/hadoop/bin/../logs/hadoop-hduser-data
node-hduser-System-Product-Name.out
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop-daemon.sh start secondary
namenode
starting secondarynamenode, logging to /home/hduser/hadoop/bin/../logs/hadoop-hd
user-secondarynamenode-hduser-System-Product-Name.out
hduser@hduser-System-Product-Name:~/hadoop$ jps
3114 DataNode
2868 NameNode
3039 TaskTracker
3189 SecondaryNameNode
3278 Jps
2960 JobTracker
hduser@hduser-System-Product-Name:~/hadoop$ █

```

กรณีรันคำสั่ง jps ไม่ได้ ให้แก้ไขโดยรันคำสั่ง

> `sudo update-alternatives --install /usr/bin/jps jps /usr/lib/jvm/jdk1.7.0/bin/jps 1`

5.3 copy file to HDFS

ก่อนรัน MapReduce ต้องคัดลอกไฟล์ข้อมูลจากหน่วยความจำไปไว้ที่ HDFS โดยรันคำสั่ง

```
> bin/hadoop dfs -copyFromLocal /home/hduser/tmp/gutenberg  
/home/hduser/gutenberg
```

รูปแบบคำสั่งคัดลอกไฟล์ คือ bin/hadoop dfs -copyFromLocal [ไดเรกทอรีที่มีไฟล์ข้อมูล]
[HDFS : ไดเรกทอรีปลายทาง]

ดูไฟล์ข้อมูลที่คัดลอกเสร็จแล้วด้วยคำสั่ง

```
> bin/hadoop dfs -ls /home/hduser/gutenberg
```

```
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop dfs -copyFromLocal /home/hduser/tmp/gutenberg /home/hduser/gutenberg  
copyFromLocal: Target /home/hduser/gutenberg/gutenberg is a directory  
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop dfs -ls /home/hduser  
Found 2 items  
drwxr-xr-x - hduser supergroup 0 2013-12-05 01:21 /home/hduser/gutenberg  
drwxr-xr-x - hduser supergroup 0 2013-12-05 00:56 /home/hduser/tmp
```

5.4 รัน MapReduce ในคู่มือนี้ยกตัวอย่างการรัน wordcount เป็นการนับจำนวนคำที่มีในไฟล์ข้อความ

```
> bin/hadoop jar hadoop-examples-0.20.2.jar wordcount /home/hduser/gutenberg  
/home/hduser/gutenberg-output
```

5.5 ดูผลลัพธ์การรันงาน

```
> bin/hadoop dfs -ls /home/hduser/gutenberg-output
```

```
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop dfs -ls /home/hduser/gutenberg  
Found 3 items  
-rw-r--r-- 1 hduser supergroup 661808 2013-12-05 09:21 /home/hduser/gutenberg/pg20417.txt  
-rw-r--r-- 1 hduser supergroup 1540093 2013-12-05 09:21 /home/hduser/gutenberg/pg4300.txt  
-rw-r--r-- 1 hduser supergroup 794290 2013-12-05 09:21 /home/hduser/gutenberg/pg5000.txt  
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop jar hadoop-0.20.2-examples.jar wordcount /home/hduser/gut  
enberg /home/hduser/gutenberg-output  
13/12/05 09:26:40 INFO Input.FileInputFormat: Total input paths to process : 3  
13/12/05 09:26:40 INFO mapred.JobClient: Running job: job_201312050912_0004  
13/12/05 09:26:41 INFO mapred.JobClient: map 0% reduce 0%  
13/12/05 09:26:54 INFO mapred.JobClient: map 66% reduce 0%  
13/12/05 09:27:00 INFO mapred.JobClient: map 100% reduce 0%  
13/12/05 09:27:03 INFO mapred.JobClient: map 100% reduce 22%  
13/12/05 09:27:12 INFO mapred.JobClient: map 100% reduce 100%  
13/12/05 09:27:14 INFO mapred.JobClient: Job complete: job_201312050912_0004  
13/12/05 09:27:14 INFO mapred.JobClient: Counters: 17  
13/12/05 09:27:14 INFO mapred.JobClient: Job Counters  
13/12/05 09:27:14 INFO mapred.JobClient: Launched reduce tasks=1  
13/12/05 09:27:14 INFO mapred.JobClient: Launched map tasks=3  
13/12/05 09:27:14 INFO mapred.JobClient: Data-local map tasks=3  
13/12/05 09:27:14 INFO mapred.JobClient: FileSystemCounters  
13/12/05 09:27:14 INFO mapred.JobClient: FILE_BYTES_READ=2012848  
13/12/05 09:27:14 INFO mapred.JobClient: HDFS_BYTES_READ=2996191  
13/12/05 09:27:14 INFO mapred.JobClient: FILE_BYTES_WRITTEN=3285297  
13/12/05 09:27:14 INFO mapred.JobClient: HDFS_BYTES_WRITTEN=766835  
13/12/05 09:27:14 INFO mapred.JobClient: Map-Reduce Framework  
13/12/05 09:27:14 INFO mapred.JobClient: Reduce input groups=71680  
13/12/05 09:27:14 INFO mapred.JobClient: Combine output records=88291  
13/12/05 09:27:14 INFO mapred.JobClient: Map input records=63766  
13/12/05 09:27:14 INFO mapred.JobClient: Reduce shuffle bytes=1005340  
13/12/05 09:27:14 INFO mapred.JobClient: Reduce output records=71680  
13/12/05 09:27:14 INFO mapred.JobClient: Spilled Records=227899  
13/12/05 09:27:14 INFO mapred.JobClient: Map output bytes=5055388  
13/12/05 09:27:14 INFO mapred.JobClient: Combine input records=522098  
13/12/05 09:27:14 INFO mapred.JobClient: Map output records=522098  
13/12/05 09:27:14 INFO mapred.JobClient: Reduce input records=88291  
hduser@hduser-System-Product-Name:~/hadoop$
```

```
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop dfs -ls /home/hduser
Found 3 items
drwxr-xr-x - hduser supergroup          0 2013-12-05 09:21 /home/hduser/gutenberg
drwxr-xr-x - hduser supergroup          0 2013-12-05 09:27 /home/hduser/gutenberg-output
drwxr-xr-x - hduser supergroup          0 2013-12-05 09:13 /home/hduser/tmp
hduser@hduser-System-Product-Name:~/hadoop$ bin/hadoop dfs -ls /home/hduser/gutenberg-output
Found 2 items
drwxr-xr-x - hduser supergroup          0 2013-12-05 09:26 /home/hduser/gutenberg-output/_logs
-rw-r--r--  1 hduser supergroup      766835 2013-12-05 09:27 /home/hduser/gutenberg-output/part-r-00000
hduser@hduser-System-Product-Name:~/hadoop$
```

5.6 หยุดการทำงาน

> *bin/stop-all.sh*

การรัน MapReduce บน ubuntu แบบหลายโหนด*

*แก้ไขจากการทำงานแบบ 1 โหนด

1. แก้ไขไฟล์ /etc/hosts

คอมพิวเตอร์ (โหนด) ติดต่อกันผ่านเครือข่ายอินเทอร์เน็ต ทางที่ง่ายที่สุดคือให้ทุกเครื่องวางอยู่ภายในเครือข่ายเดียวกัน โดยตั้งค่า IP ให้แต่ละเครื่อง แล้วแก้ไขไฟล์ /etc/hosts ทำทุกเครื่อง (master, slave)

[หมายเหตุ] ทุกเครื่องที่รัน Hadoop ต้องมีชื่อผู้ใช้ชื่อเดียวกัน (user)

```
> sudo gedit /etc/hosts
```

เพิ่มรายละเอียดดังต่อไปนี้

IP ชื่อเครื่อง เช่น

```
192.168.1.11 gpuubuntu1*
```

```
192.168.1.12 gpuubuntu2
```

```
...
```

```
...
```

*ระบุเท่ากับจำนวนโหนด บรรทัดละ 1 เครื่อง และให้ระบุเหมือนกันทุกเครื่องในไฟล์เดียวกันนี้

2. ติดตั้ง SSH

เครื่องแม่ (master) ต้องสามารถเชื่อมต่อไปยังบัญชีชื่อผู้ใช้ของตัวเอง (user account) และบัญชีชื่อผู้ใช้เดียวกันนี้ในเครื่องลูก (slave) ผ่านการเข้ารหัส SSH เนื่องจากแก้ไขเพิ่มเติมจากการทำแบบ 1 โหนด เพียงแค่รันคำสั่งต่อไปนี้

ทำที่เครื่อง master

```
> ssh-copy-id -i $HOME/.ssh/id_rsa.pub hduser@IPslave
```

เช่น `ssh-copy-id -i $HOME/.ssh/id_rsa.pub hduser@192.168.1.12`

คำสั่งนี้จะแจ้งให้ใส่รหัสผ่านเข้าสู่ระบบสำหรับผู้ใช้ hduser ของ slave จากนั้นคัดลอกคีย์ SSH เพื่อที่เราจะไม่ต้องใส่รหัสผ่านทุกครั้งเมื่อมีการเข้าถึงบัญชีของ Hadoop โดยจะต้องทำเท่ากับจำนวนเครื่อง slave ถ้ามีทั้งหมด 4 เครื่องจะต้องคัดลอกคีย์ทั้งหมด 4 ครั้ง IP ต่างกันของแต่ละเครื่อง ทดลองการใช้งานด้วยคำสั่ง

```
> ssh gpuubuntu2 (ssh ชื่อเครื่อง)
```

สามารถเข้าถึงบัญชีผู้ใช้ของ Hadoop ที่สร้างไว้ในแต่ละเครื่องได้เลย โดยไม่ต้องใส่รหัสผ่านอีก

```
> exit (ตัดการเชื่อมต่อ ออกจากบัญชีผู้ใช้)
```

3. แก้ไขไฟล์ **core-site.xml** ทุกเครื่องและ **mapred-site.xml** ที่ master

เปลี่ยน localhost เป็น ชื่อเครื่อง master ทั้งหมด

นอกจากนี้สามารถกำหนดจำนวน reduce task ได้ในไฟล์ **mapred-site.xml** ที่ master โดยเติมข้อความต่อไปนี้ระหว่าง **<configuration>**

```
<!-- In: conf/mapred-site.xml -->
<property>
  <name>mapred.reduce.tasks</name>
  <value>2</value>
  <description>set num of reduce task.</description>
</property>
```

4. แก้ไขไฟล์ **mapred-site.xml** ที่ slave

เปลี่ยน localhost เป็น ชื่อเครื่อง slave ของแต่ละเครื่องเอง

5. แก้ไขไฟล์ **hdfs-site.xml** ทุกเครื่อง

เปลี่ยน value ให้เท่ากับจำนวนโน้ต

นอกจากนี้ยังสามารถกำหนดขนาดบล็อกได้ในไฟล์นี้ (default ขนาดบล็อกคือ 64 M) โดยเติมข้อความต่อไปนี้ระหว่าง **<configuration>** ทำที่ master ถ้าขนาดบล็อกใหญ่ขึ้น จะได้จำนวน map task ลดลง

```
<!-- In: conf/hdfs-site.xml -->
<property>
  <name>dfs.block.size</name>
  <value>134217728</value>
  <description>set dfs block size, 134217728 byte is 128 M</description>
</property>
```

6. กำหนดค่า master และ slave โดยแก้ไขไฟล์ **conf/masters** และ **conf/slaves** ที่ master

6.1 **conf/masters** เติมชื่อเครื่อง master และลบคำว่า localhost

6.2 **conf/slaves** เติมชื่อเครื่องที่ใช้รันงานทั้งหมด 1 บรรทัด 1 เครื่อง และลบคำว่า localhost

7. ลบไฟล์ **conf/masters** และ **conf/slaves** ที่ slave

Running MapReduce

การรันงานแบบหลายโหนดให้รันทุกคำสั่งที่เครื่อง master เท่านั้น

1. ฟอแมตเนมโหนด

> `cd Hadoop`

> `bin/hadoop namenode -format`

2. สตาร์ทงานทุกประเภทบน master

> `bin/start-all.sh`

```
gpuubuntu1@gpuubuntu1: ~/hadoop
14/03/07 13:26:24 INFO namenode.FSNamesystem: supergroup=supergroup
14/03/07 13:26:24 INFO namenode.FSNamesystem: isPermissionEnabled=true
14/03/07 13:26:24 INFO common.Storage: Image file of size 100 saved in 0 seconds.
14/03/07 13:26:24 INFO common.Storage: Storage directory /home/gpuubuntu1/tmp/dfs/name has been successfully formatted.
14/03/07 13:26:24 INFO namenode.NameNode: SHUTDOWN_MSG:
/*****
SHUTDOWN_MSG: Shutting down NameNode at gpuubuntu1/192.168.1.11
*****/
gpuubuntu1@gpuubuntu1:~/hadoop$ bin/start-all.sh
starting namenode, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-namenode-gpuubuntu1.out
gpuubuntu2: starting datanode, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-datanode-gpuubuntu2.out
gpuubuntu3: starting datanode, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-datanode-gpuubuntu3.out
hduser2: starting datanode, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-datanode-hduser2.out
gpuubuntu1: starting datanode, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-datanode-gpuubuntu1.out
gpuubuntu1: starting secondarynamenode, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-secondarynamenode-gpuubuntu1.out
starting jobtracker, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-jobtracker-gpuubuntu1.out
gpuubuntu3: starting tasktracker, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-tasktracker-gpuubuntu3.out
gpuubuntu2: starting tasktracker, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-tasktracker-gpuubuntu2.out
hduser2: starting tasktracker, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-tasktracker-hduser2.out
gpuubuntu1: starting tasktracker, logging to /home/gpuubuntu1/hadoop/bin/../logs/hadoop-gpuubuntu1-tasktracker-gpuubuntu1.out
gpuubuntu1@gpuubuntu1:~/hadoop$ jps
6831 NameNode
7757 Jps
7082 DataNode
7322 SecondaryNameNode
7642 TaskTracker
7401 JobTracker
gpuubuntu1@gpuubuntu1:~/hadoop$
```

งานที่สตาร์ทได้บน master คือ Namenode, Datanode, Jobtracker, Tasktracker และ

Secondarynamenode

หากพบปัญหาสตาร์ทบางงานไม่ได้ให้รันคำสั่งต่อไปนี้แทน bin/start-all.sh (รันที่ **master**)

```
> bin/start-dfs.sh
```

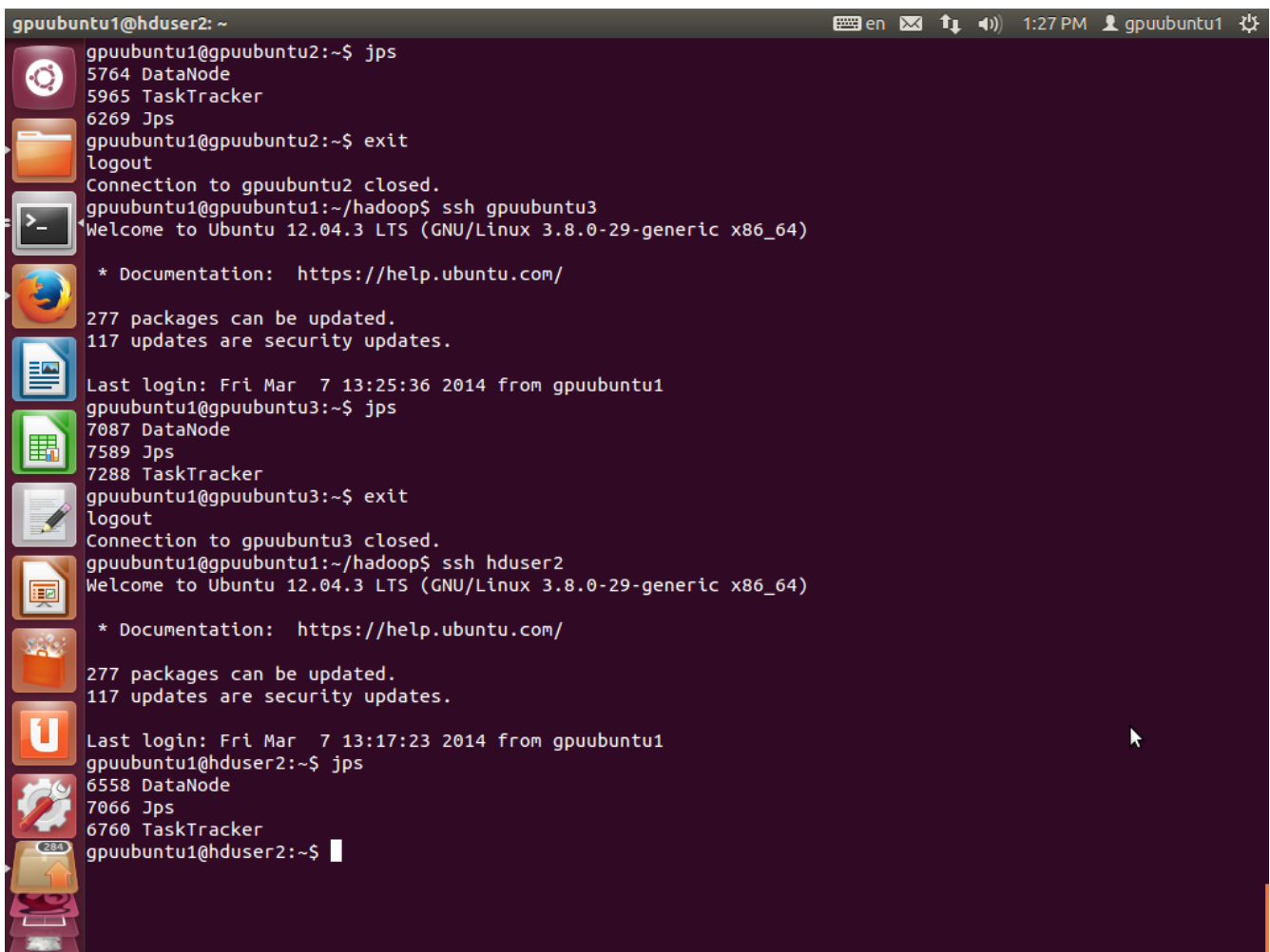
สตาร์ท Namenode, Datanode และ Secondarynamenode บน **master**

สตาร์ท Datanode บน **slave**

```
> bin/start-mapred.sh
```

สตาร์ท Jobtracker และ Tasktracker บน **master**

สตาร์ท Tasktracker บน **slave**



```
gpuubuntu1@hduser2: ~  
gpuubuntu1@gpuubuntu2:~$ jps  
5764 DataNode  
5965 TaskTracker  
6269 Jps  
gpuubuntu1@gpuubuntu2:~$ exit  
logout  
Connection to gpuubuntu2 closed.  
gpuubuntu1@gpuubuntu1:~/hadoop$ ssh gpuubuntu3  
Welcome to Ubuntu 12.04.3 LTS (GNU/Linux 3.8.0-29-generic x86_64)  
  
* Documentation:  https://help.ubuntu.com/  
  
277 packages can be updated.  
117 updates are security updates.  
  
Last login: Fri Mar  7 13:25:36 2014 from gpuubuntu1  
gpuubuntu1@gpuubuntu3:~$ jps  
7087 DataNode  
7589 Jps  
7288 TaskTracker  
gpuubuntu1@gpuubuntu3:~$ exit  
logout  
Connection to gpuubuntu3 closed.  
gpuubuntu1@gpuubuntu1:~/hadoop$ ssh hduser2  
Welcome to Ubuntu 12.04.3 LTS (GNU/Linux 3.8.0-29-generic x86_64)  
  
* Documentation:  https://help.ubuntu.com/  
  
277 packages can be updated.  
117 updates are security updates.  
  
Last login: Fri Mar  7 13:17:23 2014 from gpuubuntu1  
gpuubuntu1@hduser2:~$ jps  
6558 DataNode  
7066 Jps  
6760 TaskTracker  
gpuubuntu1@hduser2:~$
```

งานที่สตาร์ทได้บน **slave** คือ Datanode และ Tasktracker

3. copy file to HDFS

```
> bin/hadoop dfs -copyFromLocal /home/hduser/Hotel /home/hduser/hdfsHotel
```

4. running MapReduce

```
> hadoop jar ProjectTourism.jar ProjectTourism /home/hduser/hdfsHotel  
/home/hduser/hdfsHotel-output
```

รูปแบบคำสั่งคือ `hadoop jar <jar name.jar> <package.mainclass> <input> <output>`

*ยกตัวอย่างการรันงานต่อข้อความแบบ 4 โหนด

```
gpuubuntu1@gpuubuntu1: ~/hadoop
drwxr-xr-x - gpuubuntu1 supergroup 0 2014-03-07 13:26 /home/gpuubuntu1/tmp
gpuubuntu1@gpuubuntu1:~/hadoop$ bin/hadoop dfs -copyFromLocal /home/gpuubuntu1/Hotel /home/gpuubuntu1/hdfs
Hotel
gpuubuntu1@gpuubuntu1:~/hadoop$ bin/hadoop dfs -ls /home/gpuubuntu1/hdfsHotel
Found 3 items
-rw-r--r-- 4 gpuubuntu1 supergroup 67183 2014-03-07 13:29 /home/gpuubuntu1/hdfsHotel/agoda.txt
-rw-r--r-- 4 gpuubuntu1 supergroup 50693 2014-03-07 13:29 /home/gpuubuntu1/hdfsHotel/atsiam.txt
-rw-r--r-- 4 gpuubuntu1 supergroup 80205 2014-03-07 13:29 /home/gpuubuntu1/hdfsHotel/trip130257.txt
gpuubuntu1@gpuubuntu1:~/hadoop$ hadoop jar ProjectTourism.jar ProjectTourism /home/gpuubuntu1/hdfsHotel /h
ome/gpuubuntu1/hdfsHotel-output
14/03/07 13:29:58 WARN mapred.JobClient: Use GenericOptionsParser for parsing the arguments. Applications
should implement Tool for the same.
14/03/07 13:29:58 INFO input.FileInputFormat: Total input paths to process : 3
14/03/07 13:29:58 INFO mapred.JobClient: Running job: job_201403071326_0001
14/03/07 13:29:59 INFO mapred.JobClient: map 0% reduce 0%
14/03/07 13:30:07 INFO mapred.JobClient: map 66% reduce 0%
14/03/07 13:30:11 INFO mapred.JobClient: map 100% reduce 0%
14/03/07 13:30:20 INFO mapred.JobClient: map 100% reduce 100%
14/03/07 13:30:22 INFO mapred.JobClient: Job complete: job_201403071326_0001
14/03/07 13:30:22 INFO mapred.JobClient: Counters: 17
14/03/07 13:30:22 INFO mapred.JobClient: Job Counters
14/03/07 13:30:22 INFO mapred.JobClient: Launched reduce tasks=1
14/03/07 13:30:22 INFO mapred.JobClient: Launched map tasks=3
14/03/07 13:30:22 INFO mapred.JobClient: Data-local map tasks=3
14/03/07 13:30:22 INFO mapred.JobClient: FileSystemCounters
14/03/07 13:30:22 INFO mapred.JobClient: FILE_BYTES_READ=200270
14/03/07 13:30:22 INFO mapred.JobClient: HDFS_BYTES_READ=198081
14/03/07 13:30:22 INFO mapred.JobClient: FILE_BYTES_WRITTEN=400648
14/03/07 13:30:22 INFO mapred.JobClient: HDFS_BYTES_WRITTEN=192583
14/03/07 13:30:22 INFO mapred.JobClient: Map-Reduce Framework
14/03/07 13:30:22 INFO mapred.JobClient: Reduce input groups=419
14/03/07 13:30:22 INFO mapred.JobClient: Combine output records=0
14/03/07 13:30:22 INFO mapred.JobClient: Map input records=517
14/03/07 13:30:22 INFO mapred.JobClient: Reduce shuffle bytes=200282
14/03/07 13:30:22 INFO mapred.JobClient: Reduce output records=419
14/03/07 13:30:22 INFO mapred.JobClient: Spilled Records=1032
14/03/07 13:30:22 INFO mapred.JobClient: Map output bytes=198305
14/03/07 13:30:22 INFO mapred.JobClient: Combine input records=0
14/03/07 13:30:22 INFO mapred.JobClient: Map output records=516
14/03/07 13:30:22 INFO mapred.JobClient: Reduce input records=516
gpuubuntu1@gpuubuntu1:~/hadoop$
```

5. หยุดการทำงาน

```
> bin/stop-all.sh
```

6. ดูผลลัพธ์ผ่านหน้าอินเทอร์เน็ตของเครื่อง master

http://IP ของ Master:50070/ = web UI of the NameNode

ระบุรายละเอียดหน่วยความจำที่ใช้ จำนวนโนหนด ฯลฯ

Hadoop NameNode gpuubuntu1:54310 - Mozilla Firefox

192.168.1.11:50070/dfshealth.jsp

NameNode 'gpuubuntu1:54310'

Started: Fri Mar 07 13:26:34 ICT 2014
Version: 0.20.2, r911707
Compiled: Fri Feb 19 08:07:34 UTC 2010 by chrisdo
Upgrades: There are no upgrades in progress.

[Browse the filesystem](#)
[NameNode Logs](#)

Cluster Summary

7 files and directories, 1 blocks = 8 total. Heap Size is 30 MB / 889 MB (3%)

Configured Capacity	: 1.44 TB
DFS Used	: 96.06 KB
Non DFS Used	: 90.21 GB
DFS Remaining	: 1.35 TB
DFS Used%	: 0 %
DFS Remaining%	: 93.87 %
Live Nodes	: 4
Dead Nodes	: 0

NameNode Storage:

Storage Directory	Type	State
/home/gpuubuntu1/tmp/dfs/name	IMAGE_AND_EDITS	Active

[Hadoop](#), 2014.

สามารถรันคำสั่งผ่าน terminal เพื่อดูรายละเอียดของเนมโนด

> `hadoop dfsadmin -report`

```
gpuubuntu1@gpuubuntu1: ~/hadoop
Connection to hduser2 closed.
gpuubuntu1@gpuubuntu1:~/hadoop$ hadoop dfs-admin -report
Error: Could not find or load main class dfs-admin
gpuubuntu1@gpuubuntu1:~/hadoop$ hadoop dfsadmin -report
Configured Capacity: 1579695218688 (1.44 TB)
Present Capacity: 1482830340156 (1.35 TB)
DFS Remaining: 1482830241792 (1.35 TB)
DFS Used: 98364 (96.06 KB)
DFS Used%: 0%
Under replicated blocks: 0
Blocks with corrupt replicas: 0
Missing blocks: 0

-----
Datanodes available: 4 (4 total, 0 dead)

Name: 192.168.1.14:50010
Decommission Status : Normal
Configured Capacity: 976003944448 (908.97 GB)
DFS Used: 24591 (24.01 KB)
Non DFS Used: 54558961649 (50.81 GB)
DFS Remaining: 921444958208(858.16 GB)
DFS Used%: 0%
DFS Remaining%: 94.41%
Last contact: Fri Mar 07 13:28:54 ICT 2014

Name: 192.168.1.11:50010
Decommission Status : Normal
Configured Capacity: 120110559232 (111.86 GB)
DFS Used: 24591 (24.01 KB)
Non DFS Used: 10639208433 (9.91 GB)
DFS Remaining: 109471326208(101.95 GB)
DFS Used%: 0%
DFS Remaining%: 91.14%
Last contact: Fri Mar 07 13:28:54 ICT 2014

Name: 192.168.1.13:50010
Decommission Status : Normal
Configured Capacity: 243845963776 (227.1 GB)
```

gpuubuntu1@gpuubuntu1: ~/hadoop

en 1:29 PM gpuubuntu1

Non DFS Used: 54558961649 (50.81 GB)
DFS Remaining: 921444958208(858.16 GB)
DFS Used%: 0%
DFS Remaining%: 94.41%
Last contact: Fri Mar 07 13:28:54 ICT 2014

Name: 192.168.1.11:50010
Decommission Status : Normal
Configured Capacity: 120110559232 (111.86 GB)
DFS Used: 24591 (24.01 KB)
Non DFS Used: 10639208433 (9.91 GB)
DFS Remaining: 109471326208(101.95 GB)
DFS Used%: 0%
DFS Remaining%: 91.14%
Last contact: Fri Mar 07 13:28:54 ICT 2014

Name: 192.168.1.13:50010
Decommission Status : Normal
Configured Capacity: 243845963776 (227.1 GB)
DFS Used: 24591 (24.01 KB)
Non DFS Used: 15968636913 (14.87 GB)
DFS Remaining: 227877302272(212.23 GB)
DFS Used%: 0%
DFS Remaining%: 93.45%
Last contact: Fri Mar 07 13:28:54 ICT 2014

Name: 192.168.1.12:50010
Decommission Status : Normal
Configured Capacity: 239734751232 (223.27 GB)
DFS Used: 24591 (24.01 KB)
Non DFS Used: 15698071537 (14.62 GB)
DFS Remaining: 224036655104(208.65 GB)
DFS Used%: 0%
DFS Remaining%: 93.45%
Last contact: Fri Mar 07 13:28:54 ICT 2014

gpuubuntu1@gpuubuntu1:~/hadoop\$

http://IP ของ Master:50030/ = web UI of the JobTracker

ระบุชื่อเครื่องที่รัน ชื่องาน เพอร์เซ็นต์การทำงาน ฯลฯ

gpuubuntu1 Hadoop Map/Reduce Administration - Mozilla Firefox

192.168.1.11:50030/jobtracker.jsp

Cluster Summary (Heap Size is 30 MB/889 MB)

Maps	Reduces	Total Submissions	Nodes	Map Task Capacity	Reduce Task Capacity	Avg. Tasks/Node	Blacklisted Nodes
0	0	1	1	2	2	4.00	0

Scheduling Information

Queue Name	Scheduling Information
default	N/A

Filter (Jobid, Priority, User, Name)
Example: 'user:smith 3200' will filter by 'smith' only in the user field and '3200' in all fields

Running Jobs

[none](#)

Completed Jobs

Jobid	Priority	User	Name	Map % Complete	Map Total	Maps Completed	Reduce % Complete	Reduce Total	Reduces Completed	Job Scheduling Information
job_201403071326_0001	NORMAL	gpuubuntu1	projecttourism	100.00%	3	3	100.00%	1	1	NA

Failed Jobs

[none](#)

เมื่อคลิกเข้าไปที่ JobId จะระบุรายละเอียดงานเพิ่มขึ้นอีก

Hadoop job_201403071326_0001 on gpuubuntu1 - Mozilla Firefox

192.168.1.11:50030/jobdetails.jsp?jobid=job_201403071326_0001&refresh=0

Hadoop job_201403071326_0001 on gpuubuntu1

User: gpuubuntu1
Job Name: projecttourism
Job File: hdfs://gpuubuntu1:54310/home/gpuubuntu1/tmp/mapred/system/job_201403071326_0001/job.xml
Job Setup: [Successful](#)
Status: Succeeded
Started at: Fri Mar 07 13:29:58 ICT 2014
Finished at: Fri Mar 07 13:30:21 ICT 2014
Finished in: 22sec
Job Cleanup: [Successful](#)

Kind	% Complete	Num Tasks	Pending	Running	Complete	Killed	Failed/Killed Task Attempts
map	100.00%	3	0	0	3	0	0 / 0
reduce	100.00%	1	0	0	1	0	0 / 0

	Counter	Map	Reduce	Total
Job Counters	Launched reduce tasks	0	0	1
	Launched map tasks	0	0	3
	Data-local map tasks	0	0	3
FileSystemCounters	FILE_BYTES_READ	0	200,270	200,270
	HDFS_BYTES_READ	198,081	0	198,081
	FILE_BYTES_WRITTEN	200,378	200,270	400,648
	HDFS_BYTES_WRITTEN	0	192,583	192,583
	Reduce input groups	0	419	419
	Combine output records	0	0	0
	Map input records	517	0	517

http://IP ของ Master:50060/ = web UI of the TaskTracker

ระบบงานที่กำลังทำ

tracker_gpuubuntu1:ip6-localhost/127.0.0.1:50380 Task Tracker Status - Mozilla Firefox

192.168.1.11:50060/tasktracker.jsp

Most Visited ▾ Getting Started

 **hadoop**

Version: 0.20.2, r911707
Compiled: Fri Feb 19 08:07:34 UTC 2010 by chrisdo

Running tasks

Task Attempts	Status	Progress	Errors
attempt_201403071326_0001_m_000002_0	RUNNING	0.00%	
attempt_201403071326_0001_r_000000_0	RUNNING	0.00%	

Non-Running Tasks

Task Attempts	Status
attempt_201403071326_0001_m_000001_0	SUCCEEDED
attempt_201403071326_0001_m_000000_0	SUCCEEDED

Tasks from Running Jobs

Task Attempts	Status	Progress	Errors
attempt_201403071326_0001_m_000002_0	RUNNING	0.00%	
attempt_201403071326_0001_m_000001_0	SUCCEEDED	100.00%	
attempt_201403071326_0001_r_000000_0	RUNNING	0.00%	
attempt_201403071326_0001_m_000000_0	SUCCEEDED	100.00%	

Local Logs

[Log](#) directory

ผลการทดลองแสดงเวลาการทำงาน

เวลาการทำงานวัดจากการรัน map บวกกับการรัน reduce

1. รัน wordcount 1 โหนดของไฟล์ขนาด 4 GB จำนวน 1 ไฟล์

เครื่องที่ใช้รัน

- core = 4
- memory = 7.7 GiB
- disk = 976.0 GB

1 Node	Time				
Block Size (M)	64	128	256	512	1024
Task	Map = 61	Map = 31	Map = 16	Map = 8	Map = 4
Reduce = 1	5 mins, 21 sec	5 mins, 5 sec	4 mins, 45 sec	4 mins, 26 sec	4 mins, 23 sec
Reduce = 2	5 mins, 19 sec	5 mins, 4 sec	4 mins, 33 sec	4 mins, 15 sec	4 mins, 14 sec

2. รัน wordcount 4 โหนดของไฟล์ขนาด 4 GB จำนวน 1 ไฟล์

เครื่องที่เป็น master

- core = 4
- memory = 7.7 GiB
- disk = 976.0 GB

4 Node	Time				
Block Size (M)	64	128	256	512	1024
Task	Map = 61	Map = 31	Map = 16	Map = 8	Map = 4
Reduce = 1	5 mins, 22 sec	4 mins, 45 sec	4 mins, 42 sec	4 mins, 16 sec	4 mins, 15 sec
Reduce = 2	5 mins, 10 sec	4 mins, 38 sec	4 mins, 27 sec	4 mins, 12 sec	4 mins, 14 sec

3. รัน dbpedia (concat string) 1 โหนดของไฟล์ขนาด 4 GB จำนวน 1 ไฟล์

เครื่องที่ใช้รัน

- core = 4
- memory = 7.7 GiB
- disk = 976.0 GB

1 Node	Time				
Block Size (M)	64	128	256	512	1024
Task	Map = 61	Map = 31	Map = 16	Map = 8	Map = 4
Reduce = 1	7 mins, 35 sec	7 mins, 31 sec	6 mins, 40 sec	6 mins, 2 sec	6 mins, 48 sec
Reduce = 2	7 mins, 10 sec	7 mins, 23 sec	6 mins, 29 sec	5 mins, 40 sec	6 mins, 12 sec

4. รัน dbpedia (concat string) 4 โหนดของไฟล์ขนาด 4 GB จำนวน 1 ไฟล์

เครื่องที่ใช้รัน

- core = 4
- memory = 7.7 GiB
- disk = 976.0 GB

4 Node	Time				
Block Size (M)	64	128	256	512	1024
Task	Map = 61	Map = 31	Map = 16	Map = 8	Map = 4
Reduce = 1	7 mins, 9 sec	7 mins, 10 sec	6 mins, 49 sec	6 mins, 35 sec	6 mins, 31 sec
Reduce = 2	6 mins, 50 sec	7 mins, 30 sec	6 mins, 12 sec	5 mins, 22 sec	6 mins, 2 sec

ปัจจัยที่มีผลต่อเวลาการทำงาน

1. จำนวน data disk
จำนวนเพิ่มขึ้น เวลาการทำงานลดลง
2. การบีบอัดแฮดพุทของ map
เมื่อบีบอัดแฮดพุท เวลาการทำงานลดลง
3. ขนาดบล็อกข้อมูล
ขนาดบล็อกมาก เวลาการทำงานลดลง
4. จำนวน map
จำนวนงานน้อย เวลาการทำงานลดลง
5. จำนวน reduce
จำนวนงานมาก เวลาการทำงานลดลง

การปรับพารามิเตอร์ใน Hadoop

hdfs-site.xml

พารามิเตอร์	คำอธิบาย
dfs.block.size	ขนาดบล็อกข้อมูล (byte)
dfs.name.dir	ไดเรกทอรีเก็บข้อมูล namenode
dfs.data.dir	ไดเรกทอรีเก็บข้อมูล datanode
dfs.datanode.handler.count	จำนวน server threads ของ datanode
dfs.namenode.handler.count	จำนวน server threads ของ namenode
dfs.replication	แบบจำลองจำนวนบล็อก
dfs.replication.max	แบบจำลองจำนวนบล็อกสูงสุด
dfs.namenode.replication.min	แบบจำลองจำนวนบล็อกต่ำสุด
dfs.client.block.write.retries	จำนวนการลองเขียนบล็อกที่ datanode
dfs.namenode.max.objects	จำนวนสูงสุดของไฟล์ ไดเรกทอรี และบล็อก
dfs.namenode.replication.interval	ระยะเวลาที่ namenode คำนวณงานให้ datanode (second)

mapred-site.xml

พารามิเตอร์	คำอธิบาย
mapred.compress.map.output	บีบอัดเอาต์พุตของ map ก่อนส่งผ่านเครือข่าย
mapred.map.tasks.speculative.execution	กำหนดงาน map รันแบบขนาน
mapred.reduce.tasks.speculative.execution	กำหนดงาน reduce รันแบบขนาน
mapred.tasktracker.map.tasks.maximum	จำนวนงานสูงสุดของ map ที่รันโดย TaskTracker
mapred.tasktracker.reduce.tasks.maximum	จำนวนงานสูงสุดของ reduce ที่รันโดย TaskTracker
io.sort.mb	ขนาดหน่วยความจำบัฟเฟอร์ของ map ที่ใช้เรียงเอาต์พุต (megabytes)
mapred.job.reuse.jvm.num.tasks	จำนวนงานที่รันผ่าน JVM
mapred.reduce.parallel.copies	จำนวนเรดที่ไค้คัดลอกผลของ map ไปสู่การ reduce
mapred.min.split.size	ขนาดต่ำสุดในการแบ่งอินพุตของ map
mapred.max.split.size	ขนาดสูงสุดในการแบ่งอินพุตของ map
io.sort.spill.percent	Spill เนื้อหาไปสู่ดิสก์
mapred.local.dir	ไดเรกทอรีเก็บข้อมูลการทำงาน
mapred.job.shuffle.input.buffer.percent	เปอร์เซ็นต์หน่วยความจำที่ใช้เรียงลำดับเอาต์พุตของ map ระหว่างการ shuffle
mapred.job.shuffle.merge.percent	เปอร์เซ็นต์หน่วยความจำที่ใช้จัดเก็บเอาต์พุตของ map
mapred.task.timeout	ระยะเวลาในการอ่าน เขียน และแก้ไขข้อมูล (millisecond)
io.sort.record.percent	เปอร์เซ็นต์ของ io.sort.mb
mapred.job.tracker.handler.count	จำนวน server threads ของ JobTracker
mapred.max.tracker.failures	จำนวนงานที่ล้มเหลวบน TaskTracker
mapred.map.tasks	จำนวน map
mapred.reduce.tasks	จำนวน reduce
mapred.job.reduce.input.buffer.percent	เปอร์เซ็นต์หน่วยความจำที่ใช้เก็บเอาต์พุตของ map

mapred..child.java.opts	Java opts ของ TaskTracker
mapred.job.tracker	host และ port ที่ติดตามงานของ MapReduce

ภาคผนวก ข คำค้นในการทดสอบประสิทธิภาพด้านเวลา

ตารางที่ ค่าคั่นที่ใช้ในการทดลองเพื่อวัดเวลาในการสับคั่นขั้นสูง

ลำดับที่	กำหนดเงื่อนไขการคั่น	จำนวน ผลลัพธ์	เวลาในการคิวรี (s)
1	Cha-Am, Day Spa	20	0.063963136
2	Cha-Am, Day Spa, Facial Mask, Detox Massage	0	0.068681728
3	Hua Hin, Day Spa, Clairvoyant	1	0.061341696
4	Hua Hin, Hotel & Resort Spa, Clairvoyant, Cleansing	0	0.055574528
5	Cha-Am, Hua Hin, Clairvoyant, Aroma Massage	26	0.061341696
6	Cha-Am, Day Spa, Clairvoyant, Aroma Massage, Body Massage	9	0.066060288
7	Cha-Am, Hua Hin, Clairvoyant, Aroma Massage, Collagen Massage, Free Service	2	0.06815744
8	Cha-Am, Hua Hin, Gold, Face Massage	1	0.06815744
9	Cha-Am, Hua Hin, Aroma Massage, Body Massage, Collagen Massage,<500	18	0.07340032
10	Cha-Am, Hua Hin, Facial Mask, Herbal Steam, Manicure, Skin Scrub, Spa Hand & Foot	34	0.065011712
11	Cha-Am, Hua Hin, Hotel & Resort Spa, Day Spa, Cleansing, Facial Mask, Foot Massage, WIFI	0	0.061865984
12	Cha-Am, Hua Hin, Day Spa, Aroma Massage, Hot Oil Massage, Neck Shoulder Back Massage, <5000	25	0.072351744
13	Hua Hin, Day Spa, Head Massage	2	0.065011712
14	Hua Hin, Head Massage, Hot Oil Massage, Neck Shoulder Back Massage, Nature Massage	19	0.06291456
15	Hua Hin, Hotel & Resort Spa, Clairvoyant, Cleansing, Slimming Massage, Free Service, WIFI, Single Room	0	0.069206016
16	Hua Hin, Free Service, WIFI, Gym, Sauna	16	0.054525952

ลำดับที่	กำหนดเงื่อนไขการค้น	จำนวน ผลลัพธ์	เวลาในการคิวรี (s)
17	Hua Hin, Skin Scrub, Waxing, Body Massage, Detox Massage	20	0.061865984
18	Hua Hin, Slimming Massage, WIFI	0	0.061865984
19	Hua Hin, Spa Hand & face, Treatment Face, Waxing, Aroma Massage	22	0.053477376
20	Cha-Am, Hotel & Resort, Herbal Steam, Treatment Face, Face Massage, Head Massage, General Room, Yoga Room, <2000	0	0.067108864
21	Hua Hin, Destination, Treatment Face, Face Massage, WIFI	0	0.066060288
22	Hua Hin, Day Spa, Sliver, Collagen Massage	0	0.056623104
23	Hua Hin, Aroma Massage, Coconut Oil Massage, Collagen Massage	19	0.05767168
24	Cha-Am, Hua Hin, Facial Mask, Treatment Face, Hot Oil Massage,<500	5	0.056623104
25	Cha-Am, Hua Hin, Head Massage, <500	0	0.054525952
26	Cha-Am, Body Massage, Foot Massage	18	0.049283072
27	Cha-Am, Sliver	0	0.065011712
28	Cha-Am, Hua Hin, Day Spa, Detox Massage, Hot Oil Massage, Hot Stone Massage	0	0.07340032
29	Hua Hin, Destination, Herbal Steam, Head Massage	0	0.050331648
30	Hua Hin, Cleansing, Facial Mask, Herbal Steam, Manicure, Waxing	16	0.058720256
31	Hua Hin, Aroma Massage, Detox Massage, Face Massage, Nature Massage, Slimming Massage	20	0.06291456
32	Hua Hin	156	0.04194304
33	Hua Hin, Destination	3	0.04194304
34	Hua Hin, Hotel & Resort Spa	31	0.048234496
35	Hua Hin, Hotel& Resort Spa, Gold, Sliver	6	0.054525952

ลำดับที่	กำหนดเงื่อนไขการค้น	จำนวนผลลัพธ์	เวลาในการคิวรี (s)
36	Cha-Am, Hua Hin, Facial Mask, Free Service, <500	0	0.06291456
37	Hua Hin, Manicure, Skin Scrub, Collagen Massage, Detox Massage	22	0.050331648
38	Hua Hin, Aroma Massage, Body Massage, Coconut Oil Massage, Collagen Massage, Detox Massage, Face Massage, Foot Massage, Head Massage, Hot Oil Massage	29	0.058720256
39	Hua Hin, Body Massage, Detox Massage, Single Room, Yoga Room	2	0.060817408
40	Hua Hin, Body Massage, Face Massage, Free Service, WIFI, Gym, Sauna	2	0.06291456
41	Cha-Am, Skin Scrub, Aroma Massage	13	0.054525952
42	Hua Hin, Cleansing, Facial Mask, Waxing, Detox Massage, Head Massage, Hot oil, Massage, Nature Massage	14	0.069206016
43	Cha-Am, Free Service, Gym, Locker, Yoga Room	1	0.054525952
44	Hua Hin, Head Massage, Skin Scrub, < 1000	12	0.056623104
45	Cha-Am, Day Spa, Treatment Face	0	0.054525952
46	Hua Hin, Body Massage, <1000	0	0.058720256
47	Hua Hin, Head Massage, Foot Massage	27	0.048234496
48	Hua Hin, Herbal Steam, Manicure, Skin Scrub, Aroma Massage, Body Massage, Coconut Oil Massage, Collagen Massage, Detox Massage, Face Massage	26	0.067108864
49	Hua Hin, Free Service, WIFI, Gym, Sauna, <2000	3	0.06291456
50	Cha-Am, Hotel & Resort, Silver	0	0.046137344
51	Hua Hin, Hotel & Resort, Gold, Nature Massage	0	0.171835392

ลำดับที่	กำหนดเงื่อนไขการค้น	จำนวน ผลลัพธ์	เวลาในการคิวรี (s)
52	Cha-Am, Hua Hin, Gold, Hot Oil Massage, Hot Stone Massage, Single Room	1	0.194248704
53	Hua Hin, Detox Massage, Face Massage, General Room	2	0.074448896
54	Hua Hin, Hotel & Resort Spa, Sauna	11	0.064487424
55	Cha-Am, Hotel& Resort Spa, Detox Massage, Yoga Room	0	0.077856768
56	Cha-Am, Hua Hin, Day Spa, Manicure, <500	10	0.135266304
57	Cha-Am, Hua Hin, Day Spa, Spa Hand & Foot, Aroma Massage	26	0.062390272
58	Hua Hin, Destination Spa, Single Room	0	0.07077888
59	Hua Hin, Detox Massage, Free Service	0	0.062390272
60	Hua Hin, Facial Mask, Detox Massage, Nature Massage	7	0.058195968
61	Hua Hin, Day Spa, Neck Shoulder Back Massage, Nature Massage	16	0.06291456
62	Hua Hin, Day Spa, Manicure, Spa Hand & Foot	12	0.058195968
63	Hua Hin, Day Spa, Treatment Face, Collagen Massage	7	0.060817408
64	Hua Hin, Day Spa, Hot Oil Massage, Neck Shoulder Back Massage, <2000	13	0.062390272
65	Cha-Am, Day Spa, Manicure, Foot Massage	16	0.052953088
66	Cha-Am, Hua Hin, Cleansing, Sauna, Yoga Room	0	0.058195968
67	Hua Hin, Hotel & resort Spa, Gold	4	0.055574528
68	Hua Hin, Hotel & resort, Hot Stone Massage	1	0.054001664
69	Hua Hin, Detox Massage, Nature Massage, Single Room	1	0.060817408
70	Hua Hin, Face Massage, Foot Massage, Head Massage, Hot Oil Massage, Neck Shoulder Back Massage, <1000	22	0.061341696

ลำดับที่	กำหนดเงื่อนไขการค้น	จำนวน ผลลัพธ์	เวลาในการคิวรี (s)
71	Hua Hin, Aroma Massage	19	0.04980736
72	Hua Hin, Collagen Massage, Hot oil massage	2	0.04980736
73	Hua Hin, Foot Massage	26	0.04980736
74	Cha- Am, Hua Hin, Gym	6	0.050331648
75	Cha-Am, Body Massage, WIFI	0	0.057147392
76	Hua Hin, Skin Scrub, Body Massage	20	0.044040192
77	Cha-Am, Hua Hin, Face Massage, Foot Massage, <1000	31	0.054525952
78	Cha-Am, Hua Hin, Slimming Massage, General room	0	0.055574528
79	Hua Hin, Day Spa, Body Massage, Neck Shoulder Back Massage	15	0.053477376
80	Hua Hin, Skin Scrub, Spa Hand & Foot, Treatment Face, Waxing	19	0.0524288
81	Hua Hin, Sauna, WIFI	12	0.044040193
82	Hua Hin, Aroma Massage, Face Massage, Foot Massage, Head Massage	29	0.051380224
83	Cha-Am, Hua Hin, Cleansing, Facial Mask, Face Massage	7	0.055574528
84	Cha-Am, Body Massage, Coconut Massage, <1000	0	0.054525952
85	Hua Hin, Aroma Massage, Body Massage, Neck Shoulder Back Massage, Nature massage	23	0.050331648
86	Hua Hin, Waxing	7	0.061865984
87	Cha-Am, Hua Hin	179	0.042991616
88	Hua Hin, Herbal Steam, Coconut Oil Massage	2	0.050331648
89	Hua Hin, Treatment Face, Single Room	2	0.051380224
90	Cha-Am, Hua Hin, Manicure, General Room	13	0.0524288
91	Hua Hin, Facial Mask, General Room	1	0.05767168

ลำดับที่	กำหนดเงื่อนไขการค้น	จำนวน ผลลัพธ์	เวลาในการคิวรี (s)
92	Cha-Am, Hua Hin, Skin Scrub, Treatment Face, General room	21	0.051380224
93	Hua Hin, Day Spa, Hotel & Resort Spa, Detox Massage	2	0.122683392
94	Hua Hin, Collagen Massage, Free Service	0	0.050331648
95	Cha- Am, Hua Hin, Foot Massage, Slimming Massage	45	0.050331648
96	Hua Hin, Hot Oil Massage, WIFI	0	0.050331648
97	Hua Hin, Hotel & Resort, Body Massage	1	0.045088768
98	Hua Hin, Coconut Oil Massage, Collagen Massage, Detox Massage	3	0.045088768
99	Hua Hin, Face Massage, <1000	0	0.056623104
100	Cha-Am, Hua Hin, Body Massage, Collagen Massage	5	0.053477376
	เวลาเฉลี่ย	13	0.061618