



รายงานการวิจัย

การส่งเสริมการขายสินค้าของฝากของที่ระลึกด้วยการทำเหมืองข้อมูล
ตามอุปสงค์ของนักท่องเที่ยวของการพยากรณ์อนุกรมเวลา

**Using Data Mining to Increase Gifts and Souvenirs Sales Based
on Tourism Demand of Times Series Forecasting**

กฤษณ์รัฐ พิมพิลา
สาขาวิชาคณิตศาสตร์ คณะครุศาสตร์
มหาวิทยาลัยราชภัฏเพชรบูรณ์

ประจำปีงบประมาณ 2559

สัญญาที่ PCRU_2559_L016

รายงานวิจัยฉบับสมบูรณ์

การส่งเสริมการขายสินค้าของฝากของที่ระลึกด้วยการทำเหมืองข้อมูล
ตามอุปสงค์ของนักท่องเที่ยวของการพยากรณ์อนุกรมเวลา

**Using Data Mining to Increase Gifts and Souvenirs Sales Based on
Tourism Demand of Times Series Forecasting**

กฤษณ์รัฐ พิมพิลา

สาขาวิชาคณิตศาสตร์ คณะครุศาสตร์

มหาวิทยาลัยราชภัฏเพชรบูรณ์

ทุนอุดหนุนโดยมหาวิทยาลัยราชภัฏเพชรบูรณ์

ประจำปีงบประมาณ 2559

(ก)

ชื่องานวิจัย	การส่งเสริมการขายสินค้าของฝากของที่ระลึกด้วยการทำเหมืองข้อมูลตามอุปสงค์ของนักท่องเที่ยวของการพยากรณ์อนุกรมเวลา Using Data Mining to Increase Gifts and Souvenirs Sales Based on Tourism Demand of Times Series Forecasting
ผู้วิจัย	กฤษฏีรัฐ พิมพิลา
ผู้ร่วมวิจัย/ที่ปรึกษา	ภัทรารุณ วงศ์ศักดิ์, น้าฝน เป้าทองคำ
สาขาวิชา	สาขาวิชาคณิตศาสตร์ มหาวิทยาลัยราชภัฏเพชรบูรณ์ 2560

บทคัดย่อ

งานวิจัยนี้มีจุดประสงค์เพื่อ วิเคราะห์หากฎความสัมพันธ์และพยากรณ์ยอดขายของสินค้าของฝากของที่ระลึก ในจังหวัดเพชรบูรณ์โดยการทำเหมืองข้อมูล ด้วยวิธีการวิเคราะห์หากฎความสัมพันธ์โดยใช้วิธี FP-Growth(ชื่อเฉพาะ) และพยากรณ์ยอดขายด้วย Time series(ชื่อเฉพาะ) ข้อมูลที่ใช้มาจากยอดขายผลิตภัณฑ์มะขามแปรรูปย้อนหลังตั้งแต่ 1 ม.ค. 2556 – 31 ธ.ค. 2559 เป็นจำนวนทั้งสิ้น 300,765 รายการผลจากการหากฎความสัมพันธ์ของสินค้าของฝากของที่ระลึก ได้กฎความสัมพันธ์(association rule ชื่อเฉพาะ) ที่ดีที่สุดจำนวน 10 กฎและการพยากรณ์อนุกรมเวลา (Time series) เป็นการนำข้อมูลการขายสินค้าย้อนหลังในแต่ละปีมาหาช่วงเวลาที่ขายสินค้าดีที่สุด โดยแบ่งช่วงการขายออกเป็น 4 ไตรมาส ไตรมาสละ 3 เดือน ผลลัพธ์ที่ได้จากการพยากรณ์อนุกรมเวลามีช่วงเวลาที่ขายสินค้าดีที่สุดอยู่ในไตรมาสที่ 4 และไตรมาสที่ 1 ของแต่ละปี ผลจากการขายสินค้าตามกฎความสัมพันธ์ที่ได้ พร้อมช่วงเวลาขายสินค้าที่ดีที่สุดในแต่ละปีทำให้ยอดขายสินค้าเพิ่มขึ้น ถึง 15%

คำสำคัญ : การทำเหมืองข้อมูล, การจัดกลุ่มข้อมูล

(ข)

กิตติกรรมประกาศ

รายงานการวิจัยฉบับนี้สำเร็จลุล่วงไปได้ด้วยคำแนะนำต่างๆ จากคณาจารย์ในมหาวิทยาลัยราชภัฏเพชรบูรณ์ และความร่วมมือช่วยเหลืออย่างดียิ่งจากบุคคลหลายฝ่าย ที่สละเวลาให้คำแนะนำคำปรึกษา รวมถึงข้อเสนอแนะต่างๆ อันเป็นประโยชน์อย่างยิ่งต่อการดำเนินการวิจัยในครั้งนี้

ผู้วิจัยขอขอบคุณ มหาวิทยาลัยราชภัฏเพชรบูรณ์ ที่ได้จัดสรรงบประมาณเพื่อสนับสนุนการวิจัยในครั้งนี้ งานวิจัยนี้สำเร็จลุล่วงได้ดี ด้วยความอนุเคราะห์ข้อมูลยอดการขายสินค้าของฝากของที่ระลึก จากกลุ่มผู้ประกอบการร้านขายสินค้าของฝากของที่ระลึก ในจังหวัดเพชรบูรณ์ และคำแนะนำจากผู้เชี่ยวชาญจาก ศูนย์เทคโนโลยีอิเล็กทรอนิกส์และคอมพิวเตอร์แห่งชาติ

ขอขอบพระคุณ นักวิจัยจากศูนย์ เทคโนโลยีไมโครอิเล็กทรอนิกส์ ที่ให้ความรู้ด้าน อัลกอริทึมในการเข้าถึงและคัดกรองข้อมูลจากฐานข้อมูลขนาดใหญ่

สุดท้ายนี้ งานวิจัยฉบับนี้จะสำเร็จมิได้ ถ้าขาดผู้ประสานงานที่ดีด้านข้อมูลและด้านวิเคราะห์ข้อมูล จึงขอขอบพระคุณทุกท่านที่มีส่วนเกี่ยวข้องในงานวิจัยฉบับนี้

กฤษณ์รัฐ พิมพิลาและคณะ

15 มีนาคม 2560

(ก)

สารบัญ

	หน้า
บทนำ	1
1.1 บทนำ	1
1.2 วัตถุประสงค์	2
1.3 ขอบเขต	2
1.4 โครงสร้างของเอกสาร	2
2.1 บทนำ	3
2.2 การขุดค้นข้อมูล (Data Mining)	3
2.3 กฎความสัมพันธ์ (Association rule)	7
2.4 งานวิจัยที่เกี่ยวข้อง	36
3.1 บทนำ	40
3.2 การประยุกต์ใช้	42
3.3 การประเมินผล	42
3.4 สรุป	47
4.1 บทนำ	50
4.2 การประยุกต์ใช้	56

สารบัญ (ต่อ)

	หน้า
๕	๖๐
5.1	๖๐
5.2	๖๑
5.3	๖๑
	๖๒
	๖๖
	๖๘
	๖๙

(จ)

สารบัญตาราง

ข้อ	ข้อ
2-1	ข้อ 14
2-2	ข้อ 14
2-3	ข้อ 16
2-4	ข้อ 16
2-5	ข้อ 18
2-6	ข้อ 19
2-7	ข้อ 20
2-8	ข้อ 20
2-9	ข้อ 22
2-10	ข้อ 22
2-11	ข้อ 24
2-12	ข้อ 26
2-13	ข้อ 29
2-14	ข้อ 29
2-15	ข้อ 30
2-16	ข้อ 32

(น)

สารบัญรูป

รูปที่	รูปที่
2-1	รูปที่ Multilayer Feed-Forward Neural Networks..... 8
2-2	รูปที่ Genetic Algorithm..... 10
2-3	รูปที่ Crossover รูปที่ Mutation รูปที่..... 10
3-1	รูปที่..... 37
3-2	รูปที่..... 38
3-3	รูปที่..... 41
3-4	รูปที่..... 41
3-5	รูปที่..... 42
3-6	กราฟเปรียบเทียบเส้นแนวโน้มเฉลี่ยเคลื่อนที่ ของข้อมูลกลุ่มที่ 1 รูปที่..... 30
3-7	รูปที่ Regression..... 46
4-1	รูปที่..... 50
4-2	รูปที่ 2559..... 50
4-3	กราฟแสดงค่าพยากรณ์การขายปี 2559..... 52

บทที่ 1

บทนำ

1.1 ความสำคัญและที่มาของปัญหาที่ทำการวิจัย

ในปัจจุบันจังหวัดเพชรบูรณ์ เป็นจังหวัดที่มีชื่อเสียง ทางการท่องเที่ยว เนื่องจากมี สถานที่ท่องเที่ยวหลายแห่งทั้งสถานที่ท่องเที่ยวทางประวัติศาสตร์และวัฒนธรรม เช่น อุทยานประวัติศาสตร์ศรีเทพ และ สถานที่ท่องเที่ยวทางธรรมชาติที่สวยงาม อาทิเช่น อุทยานแห่งชาติน้ำหนาว เขาค้อ ยอดภูทับเบิก เป็นต้น ทำให้ปริมาณนักท่องเที่ยวที่เข้ามาเยี่ยมชมความงามตามธรรมชาติภายในจังหวัดเพชรบูรณ์ เพิ่มมากขึ้นทุกปี ตามสถิติการสำรวจของ สำนักงานท่องเที่ยวและกีฬาจังหวัดเพชรบูรณ์ พบว่าปริมาณนักท่องเที่ยว มีปริมาณไม่เท่ากันตลอดปี ตามสถิตินักท่องเที่ยวที่เข้ามาเยี่ยมชม ความงามภายในจังหวัดเพชรบูรณ์ จะมีปริมาณมากในช่วงเดือนพฤศจิกายน – ช่วงปลายเดือนกุมภาพันธ์ ตามสถิติ ตั้งแต่ปี 2554 ซึ่งปริมาณนักท่องเที่ยว นั้นสอดคล้องกับ ช่วงฤดูกาล ของผลผลิตมะขามหวานในจังหวัดเพชรบูรณ์ ซึ่งเป็นโอกาสในการขายสินค้าและผลผลิตมะขามหวาน ของเกษตรกร และยังเป็น โอกาสของร้านค้า ที่อยู่ตามสถานที่ท่องเที่ยว ซึ่งมีสินค้าประเภทของที่ระลึกประจำท้องถิ่นหรือสถานที่ท่องเที่ยวนั้นๆ อาทิเช่น พวงกุญแจ เข็มกลัด แมคเน็ต เป็นต้น

แต่ในความเป็นจริงกลับพบปัญหา เนื่องจากปริมาณการขายสินค้ามะขามแปรรูป และ สินค้าที่ระลึก กลับมียอดการขาย ไม่สอดคล้องกับจำนวนนักท่องเที่ยวที่เพิ่มขึ้น เนื่องจาก กลุ่มเกษตรกร และร้านค้า ไม่ได้ทำการวิเคราะห์ความต้องการของผู้บริโภค ไม่ได้วิเคราะห์ ปริมาณการระบายสินค้า ให้เหมาะสมตามช่วงเวลา การจัดเรียงหน้าร้านและการจัดวางสินค้า ไม่เป็นที่น่าสนใจ อีกทั้งไม่มีการส่งเสริมการขายที่สอดคล้องกับความต้องการของนักท่องเที่ยว

ผู้วิจัยจึงมีแนวคิด การแบ่งกลุ่มสินค้าและวิเคราะห์โมเดลหา กฎความสัมพันธ์เพื่อส่งเสริมการขายผลิตภัณฑ์ที่ทำจากมะขามแปรรูปและสินค้าประเภทของฝากของที่ระลึกด้วยการทำเหมืองข้อมูล โดยวิธีวิเคราะห์แบ่งกลุ่มสินค้า จากพฤติกรรมการซื้อขายสินค้าผู้บริโภค พิจารณา จากข้อมูลการขายสินค้าย้อนหลัง จากใบเสร็จประจำวัน ใบสั่งซื้อสินค้า หาความสัมพันธ์ของสินค้าในแต่ละรายการ หากกลุ่มสินค้าที่ถูกซื้อบ่อยๆช่วงเวลา ตามปริมาณนักท่องเที่ยว จากการพยากรณ์

อนุกรมเวลา เพื่อเพิ่มยอดขาย ในรูปแบบโปรโมชัน จัดกลุ่มสินค้า จัดเรียงสินค้าบนชั้นวางให้มีความสอดคล้องง่ายต่อการเลือกซื้อ ตามโมเดลกฎความสัมพันธ์

1.2 วัตถุประสงค์ของโครงการวิจัย

1.2.1 ส่งเสริมการขายสินค้าของฝากและของที่ระลึกด้วยการทำเหมืองข้อมูล ตามการพยากรณ์อนุกรมเวลา

1.2.2 เพื่อถ่ายทอดเทคโนโลยีผ่านการทำโปรโมชัน การจัดเรียงสินค้า ตามกฎความสัมพันธ์ที่หาได้ สู่กลุ่มเป้าหมาย

1.3 ขอบเขตของโครงการวิจัย

1.3.1 ขอบเขตเวลา 10 เดือน เริ่ม 1 ธันวาคม 2558 – 30 กันยายน 2559

1.3.2 ขอบเขตพื้นที่ กลุ่มผู้ขายสินค้าของฝากและของที่ระลึก ในจังหวัดเพชรบูรณ์

1.3.3 ขอบเขตด้านเนื้อหา

1.3.3.1 ศึกษาพฤติกรรมการเลือกซื้อสินค้าของนักท่องเที่ยว ที่เข้ามาท่องเที่ยวในจังหวัดเพชรบูรณ์

1.3.3.2 แบ่งกลุ่มสินค้าและสร้างกฎความสัมพันธ์ของสินค้าแต่ละรายการ โดยการทำเหมืองข้อความ

1.3.3.3 ศึกษาปริมาณนักท่องเที่ยวที่เข้ามาท่องเที่ยวในจังหวัดเพชรบูรณ์ และพยากรณ์ปริมาณนักท่องเที่ยว ตามค่าการแปรผันตามฤดูกาล

1.4 ประโยชน์ที่คาดว่าจะได้รับ

1.4.1 ได้กฎความสัมพันธ์ ของฝากของที่ระลึกที่วางขายตามร้านค้าภายในจังหวัดเพชรบูรณ์

1.4.2 ได้แนวทางการส่งเสริมการ อาทิเช่น โปรโมชันสินค้าตามฤดูกาล วิธีการจัดเรียงสินค้าภายในร้าน

บทที่ 2

เอกสารและงานวิจัยที่เกี่ยวข้อง

การศึกษางานวิจัย การส่งเสริมการขายสินค้าของฝากของที่ระลึกด้วยการทำเหมืองข้อมูลตามอุปสงค์ของนักท่องเที่ยวของการพาณิชย์อนุกรมเวลา ผู้วิจัยได้ทำการศึกษาค้นคว้างานวิจัยที่เกี่ยวข้อง ประกอบด้วยเอกสารและเนื้อหาหลักการแนวคิด วิธีการวิเคราะห์ การทำเหมืองข้อมูล การหาความสัมพันธ์และวิธีการสร้างกฎของความสัมพันธ์ของข้อมูล การพาณิชย์ หาค่าแนวโน้มของการขายสินค้า

- 2.1 การทำเหมืองข้อมูล (Data Mining)
- 2.2 การหากฎความสัมพันธ์ (Association rule)
- 2.3 การหาคุณภาพแบบทดสอบ
- 2.4 งานวิจัยที่เกี่ยวข้อง

2.1 การทำเหมืองข้อมูล(Data Mining)

มีผู้ให้คำจำกัดความของการทำเหมืองข้อมูลไว้หลายคำจำกัดความ ดังนี้

การทำเหมืองข้อมูล หมายถึง การวิเคราะห์ข้อมูล เพื่อหารูปแบบ(patterns) หรือหาความสัมพันธ์(relation) ระหว่างในฐานข้อมูลขนาดใหญ่ เป็นกระบวนการดึงข่าวสารที่น่าสนใจและมีประโยชน์แต่ไม่เคยรู้มาก่อนจากฐานข้อมูลขนาดใหญ่(Han, Kamber,Jian,2012:25)

การทำเหมืองข้อมูล (Data Mining)

การทำเหมืองข้อมูล (Data Mining) คือกระบวนการที่กระทำกับข้อมูลจำนวนมากเพื่อค้นหา รูปแบบและความสัมพันธ์ที่ซ่อนอยู่ในชุดข้อมูลนั้น ในปัจจุบันการทำเหมืองข้อมูลได้ถูกนำไปประยุกต์ใช้ในงานหลายประเภท ทั้งในด้านธุรกิจที่ช่วยในการตัดสินใจของผู้บริหาร ในด้านวิทยาศาสตร์และการแพทย์รวมทั้งในด้านเศรษฐกิจและสังคม

การทำเหมืองข้อมูล (Data Mining) เปรียบเสมือนวิวัฒนาการหนึ่งในการจัดเก็บและตีความหมายข้อมูล จากเดิมที่มีการจัดเก็บข้อมูลอย่างง่าย ๆ มาสู่การจัดเก็บในรูปแบบฐานข้อมูล

สามารถดึงข้อมูลสารสนเทศมาใช้จนถึงการทำเหมืองข้อมูลที่สามารถค้นพบความรู้ที่ซ่อนอยู่ในข้อมูล หรือจะแยกๆ เป็นข้อๆ ได้ดังนี้กระบวนการหรือการเรียงลำดับของการค้นข้อมูลจำนวนมาก และเก็บข้อมูลที่เกี่ยวข้องการนำมาใช้โดยหน่วยงานทางธุรกิจและนักวิเคราะห์ทางการเงินหรือการนำมาใช้งานในด้านวิทยาศาสตร์เพื่อเอาข้อมูลขนาดใหญ่ที่สร้างโดยวิธีการทดลองและการสังเกตการณ์ที่ทันสมัยการสกัดหรือแยกข้อมูลที่เป็นประโยชน์จากข้อมูลขนาดใหญ่หรือฐานข้อมูล การวางแผนทรัพยากรขององค์กรโดยสามารถวิเคราะห์ทางสถิติและตรรกะของข้อมูลขนาดใหญ่ เป็นการมองหารูปแบบที่สามารถช่วยการตัดสินใจได้ (Visit, 2554:4)

วิวัฒนาการของการทำเหมืองข้อมูลปี 1960 Data Collection คือ การนำข้อมูลมาจัดเก็บอย่างเหมาะสมในอุปกรณ์ที่นำเชื่อถือและป้องกันการสูญหายได้เป็นอย่างดีปี 1980 Data Access คือ การนำข้อมูลที่จัดเก็บมาสร้างความสัมพันธ์ต่อกันในข้อมูลเพื่อประโยชน์ในการนำไปวิเคราะห์ และการตัดสินใจอย่างมีประสิทธิภาพปี 1990 Data Warehouse & Decision Support คือ การรวบรวมข้อมูลมาจัดเก็บลงไปในฐานข้อมูลขนาดใหญ่โดยครอบคลุมทุกด้านขององค์กร เพื่อช่วยสนับสนุนการตัดสินใจปี 2000 Data Mining คือ การนำข้อมูลจากฐานข้อมูลมาวิเคราะห์และประมวลผล โดยการสร้างแบบจำลองและความสัมพันธ์ทางสถิติ (อดุลย์ ยิมงาม : 17/06/54)

ทำไมจึงต้องมี Data Mining ข้อมูลที่ถูกเก็บไว้ในฐานข้อมูลหากเก็บไว้เฉยๆ ก็จะไม่เกิดประโยชน์ ดังนั้นจึงต้องมีการสกัดสารสนเทศหรือการคัดเลือกข้อมูลออกมาใช้งานส่วนที่เราต้องการในอดีต เราได้ใช้คนเป็นผู้สืบค้นข้อมูลต่างๆ ในฐานข้อมูลซึ่งผู้สืบค้นจะทำการสร้างเงื่อนไขขึ้นมาตามภูมิปัญญาของผู้สืบค้นในปัจจุบันการวิเคราะห์ข้อมูลจากฐานข้อมูลเดียวอาจไม่ให้ความรู้เพียงพอและลึกซึ้งสำหรับการดำเนินงานภายใต้ภาวะที่มีการแข่งขันสูงและมีการเปลี่ยนแปลงที่รวดเร็วจึงจำเป็นที่จะต้องรวบรวมฐานข้อมูลหลายๆ ฐานข้อมูลเข้าด้วยกัน เรียกว่า “ คลังข้อมูล ” (Data Warehouse) ดังนั้นเราจึงจำเป็นต้องใช้ Data Mining ในการดึงข้อมูลจากฐานข้อมูลที่มีขนาดใหญ่ เพื่อที่จะนำข้อมูลนั้นมาใช้งานให้เกิดประโยชน์สูงสุด (http://std.kku.ac.th : 25/06/54)

ขั้นตอนการทำเหมืองข้อมูล

- 1.Data Cleaning เป็นขั้นตอนสำหรับการคัดข้อมูลที่ไม่เกี่ยวข้องออกไป
- 2.Data Integration เป็นขั้นตอนการรวมข้อมูลที่มีหลายแหล่งให้เป็นข้อมูลชุดเดียวกัน

- 3.Data Selection เป็นขั้นตอนการดึงข้อมูลสำหรับการวิเคราะห์จากแหล่งที่บันทึกไว้
- 4.Data Transformation เป็นขั้นตอนการแปลงข้อมูลให้เหมาะสมสำหรับการใช้งาน
- 5.Data Mining เป็นขั้นตอนการค้นหารูปแบบที่เป็นประโยชน์จากข้อมูลที่มีอยู่
- 6.Pattern Evaluation เป็นขั้นตอนการประเมินรูปแบบที่ได้จากการทำเหมืองข้อมูล
- 7.Knowledge Representation เป็นขั้นตอนการนำเสนอความรู้ที่ค้นพบ โดยใช้เทคนิคในการนำเสนอเพื่อให้เข้าใจ (อศุลย์ ยิมงาม : 17/06/54)

ส่วนประกอบหรือสถาปัตยกรรมของการทำเหมืองข้อมูล

Database, Data Warehouse, World Wide Web และ Other Info Repositories เป็นแหล่งข้อมูลสำหรับการทำเหมืองข้อมูล Database หรือ Data Warehouse Server ทำหน้าที่นำเข้าข้อมูลตามคำขอของผู้ใช้ Knowledge Base ได้แก่ ความรู้เฉพาะด้านในงานที่ทำจะเป็นประโยชน์ต่อการสืบค้น หรือประเมินความน่าสนใจของรูปแบบผลลัพธ์ที่ได้ Data Mining Engine เป็นส่วนประกอบหลักประกอบด้วยโมดูลที่รับผิดชอบงานทำเหมืองข้อมูลประเภทต่างๆ ได้แก่ การหาความสัมพันธ์ การจำแนกประเภท การจัดกลุ่ม

Pattern Evaluation Module ทำงานร่วมกับ Data Mining Engine โดยใช้มาตรวัดความน่าสนใจในการกลั่นกรองรูปแบบผลลัพธ์ที่ได้ เพื่อให้การค้นหามุ่งเน้นเฉพาะรูปแบบที่น่าสนใจ

User Interface ส่วนติดต่อประสานระหว่างผู้ใช้กับระบบการทำเหมืองข้อมูล ช่วยให้ผู้ใช้สามารถระบุงานทำเหมืองข้อมูลที่ต้องการทำ ดูข้อมูลหรือโครงสร้างการจัดเก็บข้อมูล ประเมินผลลัพธ์ที่ได้ (อศุลย์ ยิมงาม : 17/06/54)

ประเภทข้อมูลที่ใช้ทำเหมืองข้อมูล

Relational Database เป็นฐานข้อมูลที่จัดเก็บอยู่ในรูปแบบของตาราง โดยในแต่ละตารางจะประกอบไปด้วยแถวและคอลัมน์ ความสัมพันธ์ของข้อมูลทั้งหมดสามารถแสดงได้โดย Entity Relationship Model

Data Warehouses เป็นการเก็บรวบรวมข้อมูลจากหลายแหล่งมาเก็บไว้ในรูปแบบเดียวกัน และรวบรวมไว้ในที่ๆ เดียวกัน

Transactional Database ประกอบด้วยข้อมูลที่แต่ละทรานแซกชันแทนด้วยเหตุการณ์ในขณะใดขณะหนึ่ง เช่น ใบเสร็จรับเงิน จะเก็บข้อมูลในรูปแบบชื้อลูกค้าและรายการสินค้าที่ลูกค้ารายชื้อ

Advanced Database เป็นฐานข้อมูลที่จัดเก็บในรูปแบบอื่นๆ เช่น ข้อมูลแบบ Object-Oriented ข้อมูลที่เป็น Text File ข้อมูลมัลติมีเดีย ข้อมูลในรูปแบบของ Web (อศุลย์ ยิมงาม : 17/06/54)

ลักษณะเฉพาะของข้อมูลที่ใช้ทำเหมืองข้อมูล

ข้อมูลขนาดใหญ่ เกินกว่าจะพิจารณาความสัมพันธ์ที่ซ่อนอยู่ภายในข้อมูลได้ด้วยตาเปล่า หรือโดยการใช้ Database Management System (DBMS) ในการจัดการฐานข้อมูลข้อมูลที่มาจากหลายแหล่ง โดยอาจรวบรวมมาจากหลายระบบปฏิบัติการหรือหลาย DBMS เช่น Oracle , DB2 , MS SQL , MS Access เป็นต้นข้อมูลที่ไม่มีมีการเปลี่ยนแปลงตลอดช่วงเวลาที่ทำการ Mining หากข้อมูลที่มีอยู่นั้นเป็นข้อมูลที่เปลี่ยนแปลงตลอดเวลาจะต้องแก้ปัญหานี้ก่อน โดยบันทึกฐานข้อมูลนั้นไว้และนำฐานข้อมูลที่บันทึกไว้มาทำ Mining แต่เนื่องจากข้อมูลนั้นมีการเปลี่ยนแปลงอยู่ตลอดเวลา จึงทำให้ผลลัพธ์ที่ได้จากการทำ Mining สมเหตุสมผลในช่วงเวลาหนึ่งเท่านั้น ดังนั้นเพื่อให้ได้ผลลัพธ์ที่มีความถูกต้องเหมาะสมอยู่ตลอดเวลาจึงต้องทำ Mining ใหม่ทุกครั้งในช่วงเวลาที่เหมาะสมข้อมูลที่มิโครงสร้างซับซ้อน เช่น ข้อมูลรูปภาพ ข้อมูลมัลติมีเดีย ข้อมูลเหล่านี้สามารถนำมาทำ Mining ได้เช่นกันแต่ต้องใช้เทคนิคการทำ Data Mining ขั้นสูง (<http://std.kku.ac.th> : 25/06/54)

เทคนิคต่างๆของการทำเหมืองข้อมูล

Association rule Discovery

เป็นเทคนิคหนึ่งของ Data Mining ที่สำคัญ และสามารถนำไปประยุกต์ใช้ได้จริงกับงานต่างๆ หลักการทำงานของวิธีนี้ คือ การค้นหาความสัมพันธ์ของข้อมูลจากข้อมูลขนาดใหญ่ที่มีอยู่เพื่อนำไปใช้ในการวิเคราะห์ หรือทำนายปรากฏการณ์ต่าง ๆ หรือมากจากการวิเคราะห์การซื้อสินค้าของลูกค้าเรียกว่า “ Market Basket Analysis ” ซึ่งประเมินจากข้อมูลในตารางที่รวบรวมไว้ ผลการวิเคราะห์ที่ได้จะเป็นคำตอบของปัญหา ซึ่งการวิเคราะห์แบบนี้เป็นการใช้ “ กฎความสัมพันธ์ ” (Association Rule) เพื่อหาความสัมพันธ์ของข้อมูล

ตัวอย่างการนำเทคนิคนี้ไปประยุกต์ใช้กับงานจริง ได้แก่ ระบบแนะนำหนังสือให้กับลูกค้าแบบอัตโนมัติ ของ SE-ED BOOK ข้อมูลการสั่งซื้อหนังสือของลูกค้า SE-ED BOOK ซึ่งมีขนาดใหญ่มากจึงต้องถูกนำมาประมวลผลเพื่อหาความสัมพันธ์ของข้อมูล คือ เมื่อลูกค้าที่ซื้อหนังสือ 1 เล่ม อาจจะซื้อหนังสือเล่มใดอีกเล่มหนึ่งพร้อมกันด้วยเสมอ ความสัมพันธ์ที่ได้จากกระบวนการนี้สามารถนำไปใช้คาดเดาได้ว่าควรแนะนำหนังสือเล่มใดเพิ่มเติมให้กับลูกค้าที่เพิ่งซื้อหนังสือจากร้านไป

จากการใช้กฎความสัมพันธ์จากการประยุกต์ใช้การทำเหมืองข้อมูลข้างต้น ผู้ประกอบการวิสาหกิจขนาดกลางและขนาดย่อม (SMEs) หรืออาจารย์ นักวิจัย และ นักศึกษาระดับปริญญาโทและเอก ในมหาวิทยาลัยต่างๆ ดังนั้นวิธีการหนึ่งที่จะทำให้เราสามารถวิเคราะห์ข้อมูลเหล่านี้ได้คือการใช้ open source software และ โปรแกรม Rapid Miner ก็เป็น open source software ที่กำลังได้รับความนิยมในปัจจุบันเนื่องจากการได้รับการโหวตว่ามีผู้ใช้งานมากที่สุดจากเว็บไซต์ KDnuggets.com เมื่อปี 2013 และจากการรวบรวมข้อมูลเกี่ยวกับสินค้า OTOP พบปัญหาที่มีผลต่อยอดขายของสินค้า OTOP ด้านหนึ่งคือไม่มีการพัฒนารูปแบบของบรรจุภัณฑ์ผลิตภัณฑ์ ไม่ได้รับรอง มาตรฐาน แบบบรรจุภัณฑ์หรือที่เรียกว่า packaging มีความสำคัญในวงการธุรกิจโดยเฉพาะการเพิ่มยอดขายเนื่องจากการออกแบบบรรจุภัณฑ์เป็นกระบวนการหนึ่งในการตลาด (Marketing) การมีบรรจุภัณฑ์ที่โดนใจผู้บริโภคจะเป็นตัวกลางสำคัญที่จะสื่อสารระหว่างผลิตภัณฑ์กับผู้บริโภค ดังนั้นจากการศึกษามีการนำการทำเหมืองข้อมูล (Data Mining) มาใช้ในการวิเคราะห์ข้อมูลของผู้บริโภคเพื่อส่งเสริมการขายหลากหลายรูปแบบ เช่น การประยุกต์ใช้ในระบบโลจิสติกส์เพื่อช่วยให้ผู้ประกอบการสามารถตัดสินใจวางแผนงานต่างๆ ได้มีประสิทธิภาพมากขึ้น การจัดกลุ่มสินค้าที่เหมาะสมกับการจำหน่ายผ่านทางแท็บเล็ต และเพื่อพัฒนาแบบจำลองที่สามารถสนับสนุนกิจกรรมทางการตลาด ที่สามารถนำไปใช้สนับสนุนกิจกรรมทางการตลาดขององค์กรได้ เช่น การออกไปสเตอร์ประชาสัมพันธ์ หรือการสร้างโปร โมชัน หรือการออกผลิตภัณฑ์ใหม่ที่เหมาะสมกับลูกค้าแต่ละกลุ่ม เพื่อให้เกิดความพึงพอใจสูงสุด และความภักดีต่อองค์กร รวมถึงการกระตุ้นการซื้อผลิตภัณฑ์ของลูกค้า

2.2 การหากฎความสัมพันธ์ (Associate Rules)

Associate Rules เป็นเทคนิคที่ใช้ในการค้นหาความสัมพันธ์ที่น่าสนใจ (Interesting associations) ระหว่างลักษณะเฉพาะ (Attributes) หรือ หัวข้อ (Items) ต่างๆที่อยู่ในข้อมูล ดังนั้นผลลัพธ์ (Output) ของ Associate Rules จะแสดงความสัมพันธ์ของ ลักษณะเฉพาะ หรือ หัวข้อได้ตั้งแต่หนึ่งหรือมากกว่า และลักษณะเฉพาะ หรือ หัวข้อที่เป็นผลลัพธ์ของกฎ (Rule) ข้อหนึ่งสามารถเป็นสิ่งนำเข้า (Input) ของกฎอีกข้อหนึ่ง ขั้นตอนวิธี (Algorithm) ที่นิยมประยุกต์ใช้ได้แก่ Apriori Algorithm เสนอโดย Agawal และ Srikant ในปี 1994 เนื่องจาก Association Rules สามารถ

แสดงถึงความสัมพันธ์ระหว่างลักษณะเฉพาะหรือหัวข้อจึงถูกนำมาประยุกต์ใช้ในการตลาดอย่างแพร่หลาย ยกตัวอย่างเช่น Association Rules ค้นพบว่าเมื่อลูกค้าซื้อสินค้า A มักจะซื้อสินค้า B ด้วย โดย 70% ของลูกค้าที่ซื้อสินค้า A จะซื้อสินค้า B และ 40% ของรายการซื้อทั้งหมดลูกค้าซื้อสินค้า A และ B ดังนั้นกฎที่ได้ก็คือ

ลูกค้าซื้อสินค้า A $\Rightarrow$ ลูกค้าซื้อสินค้า B

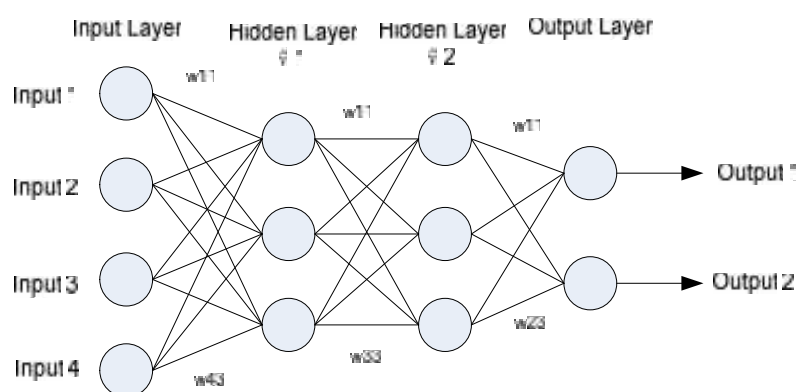
[support = 40%, confidence = 70%]

โดย $\text{Support (A} \Rightarrow \text{B)} = P(\text{A} \cup \text{B})$

$\text{Confidence (A} \Rightarrow \text{B)} = P(\text{B}|\text{A})$

2.2.1 Neural Networks

Neural networks เป็นสมการคณิตศาสตร์ที่นำหลักการทำงานของสมองมาประยุกต์ใช้โดย Han และ Kamber ได้ให้คำนิยามว่าเป็นกลุ่มของข้อมูลนำเข้า (Input) และ ผลลัพธ์ (Output) ที่เชื่อมโยงกัน โดยแต่ละจุดที่เชื่อมโยงกันจะมีน้ำหนัก (Weight) ประจำจุดเชื่อมโยงนั้นๆ เมื่อเกิดการเรียนรู้ น้ำหนักนี้จะถูกปรับเพื่อให้ได้ผลลัพธ์ที่ถูกต้อง รูปภาพที่ 2 แสดง Multilayer Feed-Forward Neural Networks



รูปที่ 2-1 แสดง Multilayer Feed-Forward Neural Networks

ขั้นตอนวิธีการเรียนรู้ (Learning Algorithm) ของ Neural Networks ที่นิยมประยุกต์ใช้ได้แก่ Backpropagation โดยขั้นตอนจะเริ่มจาก นำข้อมูลนำเข้า (Input) และ ผลลัพธ์ (Output) จากข้อมูลที

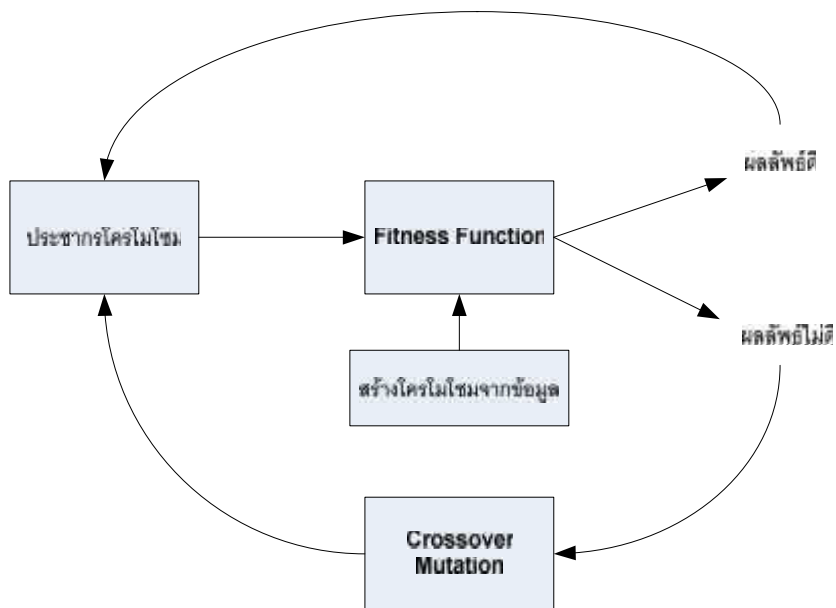
สนใจ มาทำการแบ่งไว้สำหรับการฝึก (Training) และ การทดสอบ (Test) จากนั้นนำข้อมูลที่ใช้ฝึก มาเข้าสู่กระบวนการเรียนรู้ โดยจากรูปที่ 2 Hidden Layer ที่ 1 ประกอบด้วยสาม node ใน node ที่ 1 สามารถเขียนเป็นสามการ ได้ดังนี้

$$\text{Sum}(X1W11, X2W21, X3W31, X4W41) + \text{Bias Function}$$

ในแต่ละ Node ก็จะคำนวณในลักษณะเดียวกัน จากนั้นผลลัพธ์จะถูกส่งต่อไปยัง Hidden Layer ที่ 2 ซึ่งการคำนวณก็จะมีลักษณะเช่นเดียวกัน เมื่อผลลัพธ์ส่งถึง Output Layer ผลลัพธ์จะถูกเปลี่ยนเป็นค่าผลลัพธ์จากการเรียนรู้โดย Activation Function ผลลัพธ์นี้จะถูกเปรียบเทียบกับผลลัพธ์เริ่มต้นเพื่อใช้ในการปรับค่าน้ำหนักในแต่ละ Layer จากนั้นกระบวนการเรียนรู้ก็จะเริ่มใหม่ จนกว่าค่าความถูกต้อง (Accuracy) ของผลลัพธ์จะเป็นที่พอใจ

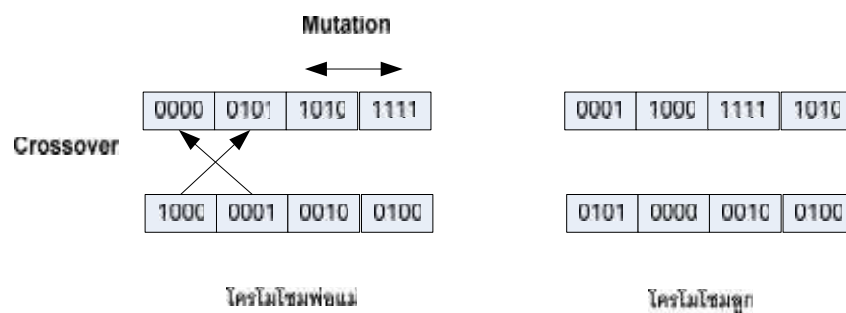
2.2.2 Genetic Algorithms

Genetic Algorithms นิยมนำมาประยุกต์ใช้ในการศึกษาข้อมูลที่มีความซับซ้อนที่ไม่สามารถศึกษาด้วยวิธีการพื้นฐานอย่างเช่นสถิติ เป็นต้น ขั้นตอนวิธีนี้สามารถประยุกต์เพื่อใช้ในการแบ่งกลุ่ม (Clustering), การพยากรณ์ (Predictions), หรือ อาจนำมาสร้างเป็น Associate Rules หลักการของ Genetic Algorithms ถูกนำเสนอครั้งแรกในปี 1986 โดย John Holland ซึ่งเป็นการประยุกต์ของการวิวัฒนาการของชีวภาพ (Biological evolution) โดยโครโมโซม (Chromosomes) ซึ่งเป็นสายของ DNA ที่ต่อกันและในหน่วยย่อยของโครโมโซมซึ่งคือ Genes จะเป็นสิ่งกำหนดลักษณะของสิ่งมีชีวิต เมื่อมีการวิวัฒนาการ Genes ของพ่อแม่จะรวมกันเพื่อไปเป็น Genes ของลูก ในการนำมาใช้ในการทำเหมืองข้อมูลแสดงในรูปภาพที่ 2-2



รูปที่ 2-2 แสดงหลักการของเทคนิค Genetic Algorithm

ข้อมูลจะถูกจัดให้เป็นลักษณะประชากรโครโมโซม (Population Chromosomes) จากนั้นจะนำไปหาผลลัพธ์จากสมการเป้าหมาย (Fitness function) โครโมโซมที่ให้ผลลัพธ์ที่ดีจะถูกเก็บไว้เพื่อใช้หาค่าสมการเป้าหมายอีกครั้ง โดยโครโมโซมที่ให้ค่าผลลัพธ์ที่ไม่ดีจะถูกนำไปทำการ Crossover หรือ Mutation จากนั้นจะนำโครโมโซมทั้งหมดมาหาค่าผลลัพธ์จากสมการเป้าหมายอีกครั้งจากนั้นกระบวนการจะทำซ้ำจนการเงื่อนไขการหยุด (Stopping Criteria) จะถูกทำให้ถูกต้อง รูปภาพที่ 4 แสดงการ Crossover และ Mutation ของโครโมโซม



รูปที่ 2-3 แสดงการ Crossover และ Mutation ของโครโมโซม

Model Construction (Learning)

เป็นขั้นการสร้าง model โดยการเรียนรู้จากข้อมูลที่ได้กำหนดคลาสไว้เรียบร้อยแล้ว (training data) ซึ่ง model ที่ได้อาจแสดงในรูปของ

- 1) โครงสร้างแบบต้นไม้ (Decision Tree)
- 2) นิวรอลเน็ต หรือ นิวรอลเน็ตเวิร์ก (Neural Net)

1) โครงสร้างแบบต้นไม้ของ Decision Tree

เป็นที่นิยมกันมากเนื่องจากเป็นลักษณะที่คนจำนวนมากคุ้นเคย ทำให้เข้าใจได้ง่าย มีลักษณะเหมือนแผนภูมิต้นไม้ โดยที่แต่ละโหนดแสดง attribute แต่ละกิ่งแสดงผลในการทดสอบ และลิฟโหนดแสดงคลาสที่กำหนดไว้

ตัวอย่าง บริษัทขนาดใหญ่แห่งหนึ่ง ทำธุรกิจอสังหาริมทรัพย์มีสำนักงานสาขาอยู่ประมาณ 50 แห่ง แต่ละสาขามีพนักงานประจำ เป็นผู้จัดการและพนักงานขาย พนักงานเหล่านี้แต่ละคนจะดูแลอาคารต่าง ๆ หลายแห่งรวมทั้งลูกค้าจำนวนมาก บริษัทจำเป็นต้องใช้ระบบฐานข้อมูลที่กำหนดความสัมพันธ์ระหว่างองค์ประกอบเหล่านี้ เมื่อรวบรวมข้อมูลแบ่งเป็นตารางพื้นฐานต่าง ๆ เช่น ข้อมูลสำนักงานสาขา (Branch) ข้อมูลพนักงาน (Staff) ข้อมูลทรัพย์สิน (Property) และข้อมูลลูกค้า (Client)

พร้อมทั้งกำหนดความสัมพันธ์ (Relationship) ของข้อมูลเหล่านี้ เช่น ประวัติการเช่าบ้านของลูกค้า (Customer_rental) รายการให้เช่า (Rentals) รายการขายสินทรัพย์ (Sales) เป็นต้น ต่อมาเมื่อมีประจักษ์กรรมการผู้บริหารของบริษัท ส่วนหนึ่งของรายงานจากฐานข้อมูลสรุปว่า “ 40 % ของลูกค้าที่เช่าบ้านนานกว่าสองปี และมีอายุเกิน 25 ปี จะซื้อบ้านเป็นของตนเอง โดยกรณีเช่นนี้เกิดขึ้น 35 % ของลูกค้าผู้เช่าบ้านของบริษัท ” ดังรูป แสดงให้เห็นถึง Decision Tree สำหรับการวิเคราะห์ว่าลูกค้าบ้านเช่าจะมีความสนใจที่จะซื้อบ้านเป็นของตนเองหรือไม่ โดยใช้ปัจจัยในการวิเคราะห์คือระยะเวลาที่ลูกค้าได้เช่าบ้านมา และอายุของลูกค้า

2) Artificial Intelligence: AI (Neural Net)

Artificial Intelligence: AI (Artificial Neural Networks (ANN) (train)

ระบบได้รู้จักที่จะคิดแก้ปัญหาที่กว้างขึ้นได้ ในโครงสร้างของนิวรอนนี้ประกอบด้วยโหนด (node) Input – Output hidden layers input layer , output layer hidden layers layer

Model Evaluation (Accuracy)

1. 데이터셋을 학습 데이터와 테스트 데이터로 나누기 (testing data)
 2. 모델 학습 (model)
 3. 모델 평가

Model Usage (Classification)

แบบ Model แบบที่เรานำมาทดสอบ (unseen data) แบบที่เรานำมาทดสอบ
แบบ object ใหม่ที่ได้มา หรือ ทำนายค่าออกมาตามที่ต้องการ

2.3 การหาคุณภาพแบบทดสอบ

[illegible]

2.3.1 ความเที่ยงตรง

[illegible]

การตรวจสอบความสอดคล้องของข้อสอบกับเนื้อหาที่สอน (Content Validity)

3. ข้อควรพิจารณา

2.3.1.1. การตรวจสอบความสอดคล้องของข้อสอบกับเนื้อหาที่สอน (content validity) เป็นการตรวจสอบที่ผู้สอนออกแบบทดสอบได้ตรงตามเนื้อหาที่สอน ในการทดสอบความเที่ยงตรงตามเนื้อหา การตรวจสอบความสอดคล้องนี้โดยให้ผู้เชี่ยวชาญในด้านเนื้อหา พิจารณาถึงความสอดคล้องระหว่าง ข้อสอบกับเนื้อหาที่สอน (Index of Item – Objective Congruence : IOC) โดยการคำนวณค่า IOC

$$IOC = \frac{\sum R}{N}$$

เมื่อ

IOC = ค่าความสอดคล้องของข้อสอบกับเนื้อหาที่สอน

$\sum R$ = ผลรวมของค่าความสอดคล้องของข้อสอบกับเนื้อหาที่สอน

N = จำนวนข้อสอบ

การตรวจสอบความสอดคล้องของข้อสอบกับเนื้อหาที่สอน (Content Validity) เป็นการตรวจสอบที่ผู้สอนออกแบบทดสอบได้ตรงตามเนื้อหาที่สอน ในการทดสอบความเที่ยงตรงตามเนื้อหา การตรวจสอบความสอดคล้องนี้โดยให้ผู้เชี่ยวชาญในด้านเนื้อหา พิจารณาถึงความสอดคล้องระหว่าง ข้อสอบกับเนื้อหาที่สอน (Index of Item – Objective Congruence : IOC) โดยการคำนวณค่า IOC

การประเมินความเสี่ยงภัยคุกคามจากภัยคุกคาม IOC การประเมินความเสี่ยงภัยคุกคาม 2-1

การประเมินความเสี่ยงภัยคุกคาม 2-1 การประเมินความเสี่ยงภัยคุกคาม

การประเมินความเสี่ยงภัยคุกคาม	การประเมิน	การประเมินความเสี่ยงภัยคุกคาม		
		1	0	-1
เมื่อเรียนจบเนื้อหาผู้เรียน การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม	1. การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม a. hardware, software, people b. hardware, software c. hardware, people d. hardware, software, internet			

การประเมินความเสี่ยงภัยคุกคาม 2-1 การประเมินความเสี่ยงภัยคุกคาม 5 การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม 2-1

การประเมินความเสี่ยงภัยคุกคาม 2-2 การประเมินความเสี่ยงภัยคุกคาม

การประเมิน	การประเมินความเสี่ยงภัยคุกคาม					การประเมิน	IOC
	การประเมิน[1]	การประเมิน[2]	การประเมิน[3]	การประเมิน[4]	การประเมิน[5]		
1	1	1	0	1	1	4	0.80

การประเมินความเสี่ยงภัยคุกคาม 2-2 การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม 4 การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม 1 การประเมินความเสี่ยงภัยคุกคาม (ΣR) การประเมินความเสี่ยงภัยคุกคาม 4 การประเมินความเสี่ยงภัยคุกคาม IOC การประเมินความเสี่ยงภัยคุกคาม

$$IOC = \frac{4}{5}$$

$$= 0.80$$

การประเมินความเสี่ยงภัยคุกคาม 0.80 การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม 0.8 การประเมินความเสี่ยงภัยคุกคาม 1 การประเมินความเสี่ยงภัยคุกคาม IOC การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม 0.5 การประเมินความเสี่ยงภัยคุกคาม 0.5 ถือว่าข้อสอบข้อนั้นไม่มีความสอดคล้องกับวัตถุประสงค์เชิงพฤติกรรม จะต้องตัดข้อสอบนั้น การประเมินความเสี่ยงภัยคุกคาม การประเมินความเสี่ยงภัยคุกคาม

ตารางที่ 2-3 ข้อมูลการทดสอบที่สร้างขึ้นและเกรดเฉลี่ยของ

ลำดับ	คะแนนการทดสอบที่สร้างขึ้น (X)	เกรดเฉลี่ย (Y)
1	30	3.5
2	28	3.4
3	20	3.0
4	15	2.5
5	18	2.7
6	26	3.0

ตารางที่ 2-3 เป็นคะแนนจากการทำแบบทดสอบที่สร้างขึ้นและเกรดเฉลี่ยของ
นักศึกษา 6 คน ข้อมูลเหล่านี้จะนำมาใช้เพื่อหาความสัมพันธ์ระหว่างคะแนนการทดสอบที่สร้างขึ้นและเกรดเฉลี่ย

ตารางที่ 2-4 ข้อมูลการคำนวณค่าสหสัมพันธ์

ลำดับ	X	Y	X ²	Y ²	XY
1	30	3.50	900	12.25	105.00
2	28	3.40	784	11.56	95.20
3	20	3.00	400	9.00	60.00
4	15	2.50	225	6.25	37.50
5	18	2.70	324	7.29	48.60
6	26	3.00	676	9.00	78.00
Σ	137	18.10	3309	55.35	424.30

จากข้อมูลที่คำนวณได้ในตารางที่ 2-4 ค่าสหสัมพันธ์ r_{xy} สามารถคำนวณได้ดังนี้

$$\begin{aligned}
 r_{xy} &= \frac{2545.80 - 2479.70}{\sqrt{[19854 - 18769][332.10 - 327.61]}} \\
 &= \frac{66.10}{\sqrt{(1085)(4.49)}} \\
 &= \frac{66.10}{69.80} \\
 &= 0.95
 \end{aligned}$$

ค่า r_{xy} อยู่ระหว่าง 0.95 ถึง 1.00 หมายถึง แบบทดสอบมีความเที่ยงตรงเชิงสภาพคำ

2) **predictive validity** (predictive validity) หมายถึง การใช้หลักการเกี่ยวกับการทดสอบความเที่ยงตรงเชิงสภาพ เพียงแต่ถ้าเป็นแบบเชิงพยากรณ์จะใช้ข้อมูลที่เป็นเกรดเฉลี่ยของผู้เรียนที่สำเร็จการศึกษาแล้ว หรือใช้คะแนนที่ได้จากการทดสอบแบบทดสอบที่ต้องการหาความเที่ยงตรงเชิงพยากรณ์ก่อน เมื่อทดสอบสอบแล้วจะต้องรอให้ผู้เรียนกลุ่มนี้ได้คะแนนที่จะใช้เป็นเกณฑ์ในการหาความสัมพันธ์ จึงจะดำเนินการคำนวณหา r_{xy} โดยใช้สูตร $r_{xy} = \frac{\sum (X - \bar{X})(Y - \bar{Y})}{\sqrt{\sum (X - \bar{X})^2 \sum (Y - \bar{Y})^2}}$ ค่า r_{xy} ที่ได้ทดสอบก่อนหน้านี้เพื่อหาค่าความเที่ยงตรงเชิงพยากรณ์ ใช้สูตร $r_{xy} = \frac{\sum (X - \bar{X})(Y - \bar{Y})}{\sqrt{\sum (X - \bar{X})^2 \sum (Y - \bar{Y})^2}}$

construct validity (construct validity) หมายถึง โครงสร้างของแบบทดสอบมาตรฐาน โดยแบบทดสอบที่สร้างขึ้นจะมีมาตรฐานที่วัดลักษณะเดียวกัน เมื่อคำนวณค่าได้แล้วพบว่า ถ้าค่าที่คำนวณได้เข้าใกล้ 1.00 หมายถึง แบบทดสอบมีความเที่ยงตรงเชิงโครงสร้าง

ความเชื่อมั่น หมายถึง ความคงเส้นคงวาของผลการวัดจากการที่นำแบบทดสอบชุดนั้นไปทดสอบกับผู้เรียนไม่ว่าจะทดสอบจำนวนกี่ครั้งคะแนนที่ได้จะไม่แตกต่างกัน ความเชื่อมั่นสามารถคำนวณได้จาก $r_{tt} = \frac{\sum (X - \bar{X})(Y - \bar{Y})}{\sqrt{\sum (X - \bar{X})^2 \sum (Y - \bar{Y})^2}}$

ว่าแบบทดสอบนั้นมีค่าความเชื่อมั่นสูง วิธีการคำนวณหาค่าความเชื่อมั่นสามารถคำนวณหาได้
ดังนี้

วิธีการสอบซ้ำ

การสอบซ้ำ (test – retest) เป็นการนำแบบทดสอบชุดเดิมมาทดสอบซ้ำกับกลุ่มตัวอย่าง
กลุ่มเดิม (stability) อย่างน้อย 2 ครั้ง เพื่อวัด
ความแตกต่างกัน ในการวัดผลจะวัดในเวลาที่แตกต่างกันแล้วนำคะแนนที่ได้ทั้ง 2 ครั้ง
มาหาค่าความเชื่อมั่นได้ดังนี้

ตารางที่ 2-5 ข้อมูลการสอบซ้ำ

ข้อสอบ	คะแนนครั้งที่ 1 (X)	คะแนนครั้งที่ 2 (Y)
1	80	79
2	70	65
3	90	80
4	50	50
5	66	65
6	45	45

ตารางที่ 2-5 แสดงข้อมูลการสอบซ้ำกับกลุ่มตัวอย่าง 2 ครั้ง เพื่อหาค่าความเชื่อมั่น
รวม 100 ข้อสอบ โดยนำคะแนนครั้งที่ 1 (X) และคะแนนครั้งที่ 2 (Y) มาคำนวณหา
ค่าความเชื่อมั่นได้ดังนี้

ตารางที่ 2-6

ตารางที่ 2-6 ข้อมูลดิบของคะแนนสอบวิชาคณิตศาสตร์

ข้อสอบ	X	Y	X ²	Y ²	XY
1	80	79	6400	6241	6320
2	70	65	4900	4225	4550
3	90	80	8100	6400	7200
4	50	50	2500	2500	2500
5	66	65	4356	4225	4290
6	45	45	2025	2025	2025
Σ	401	384	28281	25616	26885

การคำนวณค่าสัมประสิทธิ์สหสัมพันธ์ของเพียร์สัน (2-6) โดยใช้ข้อมูลดิบของคะแนนสอบวิชาคณิตศาสตร์

$$\begin{aligned}
 r_{XY} &= \frac{161310 - 153984}{\sqrt{[169686 - 160801][153696 - 147456]}} \\
 &= \frac{7326}{\sqrt{(8885)(6240)}} \\
 &= \frac{7326}{7445.96} \\
 &= 0.98
 \end{aligned}$$

การคำนวณค่าสัมประสิทธิ์สหสัมพันธ์ของเพียร์สัน (2-6) โดยใช้ข้อมูลดิบของคะแนนสอบวิชาคณิตศาสตร์

มากเนื่องจาก ค่าที่คำนวณได้มีค่าเข้าใกล้ 1 มาก

วิธีการใช้แบบทดสอบคู่ขนาน

แบบทดสอบคู่ขนาน (parallel form) เป็นแบบทดสอบที่มีลักษณะคล้ายคลึงกันมากกับแบบทดสอบเดิม แต่มีข้อแตกต่างที่สำคัญคือ แบบทดสอบคู่ขนานจะมีจำนวนข้อสอบที่เท่ากัน แต่มีเนื้อหาที่แตกต่างจากแบบทดสอบเดิม โดยแบบทดสอบคู่ขนานจะประกอบด้วยข้อสอบที่เหมือนกัน 2 ข้อ และข้อสอบที่แตกต่างจากแบบทดสอบเดิม 2 ข้อ แบบทดสอบคู่ขนานนี้สามารถใช้เพื่อวัดความน่าเชื่อถือของแบบทดสอบได้ โดยแบบทดสอบคู่ขนานจะประกอบด้วยข้อสอบที่เหมือนกัน 2 ข้อ และข้อสอบที่แตกต่างจากแบบทดสอบเดิม 2 ข้อ

ตารางที่ 2-7 ข้อมูลตัวอย่างของคะแนนครั้งที่ 1 และ 2

ลำดับ	คะแนนครั้งที่ 1 (X)	คะแนนครั้งที่ 2 (Y)
1	60	70
2	50	40
3	40	60
4	70	70
5	90	80
6	80	75

ถ้าเรานำข้อมูลในตารางที่ 2-7 เมื่อนำไปคำนวณค่าต่าง ๆ จะได้ข้อมูลดังตารางที่ 2-8

ตารางที่ 2-8 ข้อมูลตัวอย่างของคะแนนครั้งที่ 1 และ 2

ลำดับ	X	X ²	Y	Y ²	XY
1	60	3600	70	4900	4200
2	50	2500	40	1600	2000
3	40	1600	60	3600	2400
4	70	4900	70	4900	4900
5	90	8100	80	6400	7200
6	80	6400	75	5625	6000
Σ	390	27100	395	27025	26700

ถ้าเรานำข้อมูลในตารางที่ 2-8 ไปคำนวณค่าต่าง ๆ จะได้ข้อมูลดังตารางที่ 2-9

$$\begin{aligned}
 r_{tt} &= \frac{160200-154050}{\sqrt{[162600-152100][162150-156025]}} \\
 &= \frac{6150}{\sqrt{[10500][6125]}} \\
 &= \frac{6150}{8019.50} \\
 &= 0.77
 \end{aligned}$$

ค่าสัมประสิทธิ์สหสัมพันธ์ r_{tt} เท่ากับ 0.77 แสดงให้เห็นว่ามีความสัมพันธ์ที่ค่อนข้างสูงระหว่าง
การวัดซ้ำครั้งที่ 1 และการวัดซ้ำครั้งที่ 2

วิธีการแบ่งครึ่งแบบทดสอบ

วิธีแบ่งครึ่งแบบทดสอบ (split – half) เป็นการนำข้อสอบมาแบ่งเป็น 2 ส่วนเท่าๆกัน
และนำผลรวมของข้อสอบทั้งสองส่วนมาหารกันเพื่อหาความสัมพันธ์ระหว่างข้อสอบทั้งสองส่วน
ซึ่งจะได้ค่าความสัมพันธ์ (r1) แล้วนำค่าความสัมพันธ์ที่ได้ไปคำนวณหาความเชื่อมั่นทั้งฉบับโดยใช้สูตร
ของ Spearman – Brown ดังนี้

$$r_t = \frac{2r_{\frac{1}{2}}}{1 + r_{\frac{1}{2}}}$$

เมื่อ

r_t คือ สัมประสิทธิ์ความเชื่อมั่นแบบทดสอบทั้งฉบับ

$r_{\frac{1}{2}}$ คือ สัมประสิทธิ์ความเชื่อมั่นแบบทดสอบครึ่งหนึ่ง

การคำนวณหาความเชื่อมั่นแบบทดสอบทั้งฉบับโดยใช้สูตรของ Spearman – Brown 2-9

例題 2-9 二変量データの相関係数の計算

順位	X	Y
1	10	10
2	8	8
3	7	5
4	6	8
5	10	8
6	8	9

例題 2-9 のデータから、 $n=6$ の二変量データの相関係数を計算する。20 点満点の試験問題として出題された。

例題 2-10 二変量データの相関係数の計算

順位	X	X ²	Y	Y ²	XY
1	10	100	10	100	100
2	8	64	8	64	64
3	7	49	5	25	35
4	6	36	8	64	48
5	10	100	8	64	80
6	8	64	9	81	72
Σ	49	413	48	398	399

例題 2-10 のデータから、相関係数を計算する。

$$\begin{aligned}
 r_{1t} &= \frac{2394 - 2352}{\sqrt{[2478 - 2401][2388 - 2304]}} \\
 &= \frac{42}{\sqrt{[77][84]}} \\
 &= \frac{42}{80.42} \\
 &= 0.52
 \end{aligned}$$

$$r_{\frac{1}{2}} = 0.52$$

ฉบับ เมื่อนำค่าที่คำนวณได้นี้ไปคำนวณหาความเชื่อมั่นแบบทดสอบทั้งฉบับ

$$r_t = \frac{2 \times 0.52}{1 + 0.52} = 0.68$$

0.68 แสดงว่าแบบทดสอบมีความเชื่อมั่นในระดับปานกลาง เนื่องจากค่าที่ได้มีค่าใกล้เคียง 1

(Kuder – Richardson : KR) 0.68

1. KR-20

ง่ายในลักษณะกระจาย สูตรที่ใช้ในการหาวิธีแบบนี้

$$r_t = \frac{n}{n-1} \left\{ 1 - \frac{\sum p_q}{S_t^2} \right\}$$

$$S_t^2 = \frac{N \sum x^2 - (\sum x)^2}{N^2}$$

r_t ค่าความเชื่อมั่นแบบทดสอบ

n จำนวนข้อสอบ

$$p = \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6}$$

$$q = \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6}$$

$$s_t^2 = \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6}$$

$$N = 6 \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6} \quad \frac{1}{6}$$

จากสูตรสามารถอธิบายวิธีการคำนวณหาค่าความเชื่อมั่นได้ โดยใช้ข้อมูลที่แสดงตามตารางที่ 2-11

ตารางที่ 2-11 ข้อมูลการคำนวณหาความเชื่อมั่น

ค่า ความเชื่อมั่น	1	2	3	4	5	ค่า ความเชื่อมั่น (X)	X ²
1	1	1	1	1	1	5	25
2	1	1	1	1	1	5	25
3	1	1	1	1	1	5	25
4	0	0	0	1	1	2	4
5	1	0	0	1	0	2	4
6	0	0	0	1	0	1	1
Σ	4	3	3	6	4	20	84
p	0.67	0.50	0.50	1.00	0.67	$\bar{x} = 3.33$	
q	0.33	0.50	0.50	0.00	0.33		
pq	0.22	0.25	0.25	0.00	0.22	Σpq = 0.94	

จากตารางที่ 2-11 ข้อมูลการคำนวณหาความเชื่อมั่น สามารถคำนวณหาความเชื่อมั่นได้ โดยใช้สูตรต่อไปนี้

$$s_t^2 = \frac{6(84) - (20 \times 20)}{(6 \times 6)}$$

$$s_t^2 = 2.89$$

$$r_t = \frac{5}{5-1} \left\{ 1 - \frac{0.94}{2.89} \right\}$$

$$r_t = 1.25 * 0.67$$

$$r_t = 0.84$$

ค่าความเชื่อมั่น 0.84 หมายถึง แบบทดสอบชุดนี้มีความเชื่อมั่นสูง เนื่องจากค่าความเชื่อมั่น
 อยู่ระหว่าง 0.80 ถึง 1.00 ค่าความเชื่อมั่น 0.6 ถึง 1.0

2. KR-21 เป็นค่าความเชื่อมั่นแบบทดสอบที่ใช้ในการคำนวณมีรูปแบบดังนี้
 ของข้อสอบแต่ละข้อมีค่าใกล้เคียงกัน สูตรที่ใช้ในการคำนวณมีรูปแบบดังนี้

$$r_t = \frac{n}{n-1} \left\{ 1 - \frac{\bar{X}(n-\bar{X})}{n s_t^2} \right\}$$

ซึ่ง

r_t คือ สัมประสิทธิ์ของความเชื่อมั่นของแบบทดสอบทั้งฉบับ

n คือ จำนวนข้อสอบ

$\bar{X}$ คือ ค่าเฉลี่ยของข้อสอบ

s_t^2 คือ ค่าความแปรปรวนของข้อสอบ

N คือ จำนวนผู้สอบ

ตัวอย่างการคำนวณค่าความเชื่อมั่นแบบทดสอบ KR-21

$$r_t = \frac{5}{5-1} \left\{ 1 - \frac{3.33(5-3.33)}{5*(2.89)} \right\}$$

$$= 0.77$$

ค่าความเชื่อมั่น 0.77 อยู่ระหว่าง 0.70 ถึง 0.80 ค่าความเชื่อมั่น 0.6 ถึง 1.0
 ค่าความเชื่อมั่น 1.0 อยู่ระหว่าง 0.80 ถึง 1.00 ค่าความเชื่อมั่น 0.6 ถึง 1.0

วิธีการหาสัมประสิทธิ์แอลฟา

สัมประสิทธิ์แอลฟา (α - Coefficient) คือ สัมประสิทธิ์ความเชื่อมั่นของแบบทดสอบ ค่าความเชื่อมั่นที่คำนวณได้จากสูตรครอนบาช (Cronbach) แบบทดสอบโดยแบบทดสอบค่าคะแนนที่ได้ อาจจะเป็นค่าอะไรก็ได้ที่มีค่ามากกว่า 1

$$\alpha = \frac{n}{n-1} \left\{ 1 - \frac{\sum s_i^2}{s_t^2} \right\}$$

เมื่อ

α คือ สัมประสิทธิ์แอลฟา

n คือ จำนวนข้อสอบ

s_i^2 คือ ค่าความแปรปรวนของข้อสอบแต่ละข้อ

s_t^2 คือ ค่าความแปรปรวนของแบบทดสอบ

ค่าความเชื่อมั่นของแบบทดสอบ โดยวิธีการหาสัมประสิทธิ์แอลฟา

ตาราง 2-12 แสดงค่าความเชื่อมั่นของแบบทดสอบ

ข้อสอบ / ข้อสอบ	1	2	3	4	5	ค่าเฉลี่ย (X)	X ²
1	7	10	9	10	8	44	1936
2	8	8	8	8	8	401	1600
3	9	10	9	7	9	44	1936
4	7	8	7	6	7	35	1225
5	8	7	6	7	6	34	1156
6	6	6	6	8	7	33	1089

ΣX	45	49	45	46	45	230	8942
$(\Sigma x)^2$	2025	2401	2025	2116	2025		
$\Sigma(x^2)$	343	413	347	362	343		
s_i^2	0.92	2.14	1.58	1.55	0.92	7.11	

ข้อ 2-12 มี 5 ข้อ ข้อละ 10
เมื่อคำนวณค่าต่าง ๆ แล้วสามารถนำมาคำนวณหาค่าความแปรปรวนแบบทดสอบทั้งฉบับและสัม
ประสิทธิ์

$$S_t^2 = \frac{(6 \cdot 8942) - (230 \cdot 230)}{(6 \cdot 6)}$$

$$S_t^2 = 20.89$$

$$a = \frac{5}{5-1} \left\{ 1 - \frac{7.11}{20.89} \right\}$$

$$= 0.82$$

ค่าสัมประสิทธิ์ความแปรปรวน 0.82
ข้อ 1
ข้อ 2

ข้อ 2 กลุ่มในที่นี้คือ ผู้เรียนกลุ่มเก่งและผู้เรียนกลุ่มอ่อน หรือกลุ่มที่ชอบและกลุ่มที่ไม่ชอบ ค่า
-1 ถึง 1

0.40

0.30 – 0.39

0.20 – 0.29

0.20

จะต้องตัดข้อสอบข้อนั้นทิ้งไป

當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。

$$D = PH - PL$$

當 $PH = 1/3$

$PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。

$PL = 1/3$ 時， $PH = 1/3$ ， $D = 1/3$ 。

當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。

當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。

$$PH = 3/3$$

$$= 1$$

$$PL = 1/3$$

$$= 0.33$$

$$D = 1 - 0.33$$

$$= 0.67$$

當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。

當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。

當 $PH = 1/3$	1	2	3	4	5
當 $PH = 1/3$	3	2	3	3	3
當 $PH = 1/3$	1	1	0	0	1
當 D	0.67	0.33	1	1	0.67

當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。當 $PH = 1/3$ 時， $PL = 1/3$ ， $D = 1/3$ 。

การหาค่าสัมประสิทธิ์สหสัมพันธ์จุด-อันดับ (point biserial correlation) เป็นการหาค่าสัมประสิทธิ์สหสัมพันธ์ระหว่างตัวแปรที่เป็นอันดับ (ordinal) กับตัวแปรที่เป็นจุด (dichotomous) ค่าสัมประสิทธิ์นี้จะอยู่ในช่วง 0 ถึง 1 ถ้าค่าสัมประสิทธิ์เป็น 1 แสดงว่ามีความสัมพันธ์กันอย่างสมบูรณ์ ถ้าเป็น 0 แสดงว่าไม่มีความสัมพันธ์กัน

$$r_{p.bis} = \frac{\bar{X}_p - \bar{X}_f}{S_t} \cdot \sqrt{pq}$$

เมื่อ

$r_{p.bis}$ คือ ค่าสัมประสิทธิ์สหสัมพันธ์จุด-อันดับ

$\bar{X}_p$ คือ ค่าเฉลี่ยอันดับของคนที่ตอบข้อถูก

$\bar{X}_f$ คือ ค่าเฉลี่ยอันดับของคนที่ตอบข้อผิด

p คือ สัดส่วนผู้เรียนที่ทำข้อสอบข้อนั้นได้

q คือ สัดส่วนผู้เรียนที่ทำข้อสอบข้อนั้นผิด $1-p$

S_t คือ ค่าเบี่ยงเบนมาตรฐานของอันดับ

$$S_t = \sqrt{\frac{N \sum X^2 - (\sum X)^2}{N(N-1)}}$$

$$\bar{X}_p = \frac{\sum xy}{n_p}$$

$$\bar{X}_f = \frac{\sum x - \sum xy}{n_f}$$

เมื่อ

N คือ จำนวนข้อสอบทั้งหมด

X คือ จำนวนข้อสอบที่ตอบถูก

n_f คือ จำนวนข้อสอบที่ตอบผิด

n_p คือ จำนวนข้อสอบที่ตอบถูก

การหาค่าสัมประสิทธิ์สหสัมพันธ์จุด-อันดับสามารถใช้หาค่าสัมประสิทธิ์สหสัมพันธ์อันดับได้

Table 2-16: Data for the regression analysis

Order	Price (X)	Quality (Y)	XY	X ²
1	30	1	30	900
2	20	1	20	400
3	25	1	25	625
4	15	0	0	225
5	14	0	0	196
6	18	1	18	324
Σ	122	4	93	2670

Table 2-16: Data for the regression analysis. The data shows the relationship between Price (X) and Quality (Y). The mean of Price is 20.33, the mean of Quality is 0.67, and the correlation coefficient is 0.34.

$$\bar{X}_p = \frac{93}{4} = 23.25$$

$$\bar{X}_f = \frac{122-93}{2} = 14.5$$

$$p = \frac{4}{6} = 0.67$$

$$q = \frac{2}{6} = 0.34$$

$$St = \sqrt{\frac{6(2670) - (122)^2}{6(6-1)}}$$

$$= \sqrt{\frac{16020 - 14884}{30}}$$

$$= 6.15$$

$$\text{rp.bis} = \frac{23.25 - 14.50}{6.15} \sqrt{0.67 * 0.34}$$

$$= 1.42 * 0.48$$

$$= 0.68$$

การคำนวณค่า rp.bis นี้ได้มาจากการนำค่าความยากง่ายของข้อสอบแต่ละข้อมาลบด้วยค่าเฉลี่ยของความยากง่ายของข้อสอบทั้งหมด แล้วนำผลที่ได้มาหารด้วยค่าความยากง่ายของข้อสอบทั้งหมด แล้วนำผลที่ได้มาคูณด้วยค่าความยากง่ายของข้อสอบทั้งหมด

ความยากง่าย

การคำนวณค่าความยากง่ายของข้อสอบแต่ละข้อทำได้โดยการนำค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบแต่ละข้อมาหารด้วยจำนวนข้อสอบทั้งหมด แล้วนำผลที่ได้มาคูณด้วยค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบทั้งหมด แล้วนำผลที่ได้มาคูณด้วยค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบทั้งหมด

นั่นต้องนำไปทดสอบหลาย ๆ ครั้ง และปรับปรุงจนได้ค่าความยากง่ายใกล้เคียงกับ 50%

การคำนวณค่าความยากง่ายของข้อสอบแต่ละข้อทำได้โดยการนำค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบแต่ละข้อมาหารด้วยจำนวนข้อสอบทั้งหมด แล้วนำผลที่ได้มาคูณด้วยค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบทั้งหมด แล้วนำผลที่ได้มาคูณด้วยค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบทั้งหมด

การคำนวณค่าความยากง่ายของข้อสอบแต่ละข้อทำได้โดยการนำค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบแต่ละข้อมาหารด้วยจำนวนข้อสอบทั้งหมด แล้วนำผลที่ได้มาคูณด้วยค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบทั้งหมด แล้วนำผลที่ได้มาคูณด้วยค่าความถี่ของคำตอบที่ถูกต้องของข้อสอบทั้งหมด

$$P = \frac{R}{N}$$

เมื่อ

P คือ ค่าความยากง่าย

R คือ จำนวนคำตอบที่ถูกต้อง

N คือ จำนวนข้อสอบ

การสุ่มตัวอย่างแบบง่าย (Simple Random Sampling) จำนวน 30 คน ผลการสุ่มได้ดังนี้

$$p = \frac{20}{30} \\ = 0.67$$

การสุ่มตัวอย่างแบบง่าย (Simple Random Sampling) จำนวน 30 คน ผลการสุ่มได้ดังนี้ 0.67

ความเป็นปรนัย

การสุ่มตัวอย่างแบบง่าย (Simple Random Sampling) จำนวน 30 คน ผลการสุ่มได้ดังนี้

1. การสุ่มตัวอย่างแบบง่าย (Simple Random Sampling) จำนวน 30 คน ผลการสุ่มได้ดังนี้

2. การสุ่มตัวอย่างแบบง่าย (Simple Random Sampling) จำนวน 30 คน ผลการสุ่มได้ดังนี้

3. ความเข้มแข็งในการแปลความหมายของคะแนน หมายถึงเมื่อทุกคนได้เห็นคะแนนที่ได้

การสุ่มตัวอย่างแบบง่าย (Simple Random Sampling) จำนวน 30 คน ผลการสุ่มได้ดังนี้

บทสรุป

การสุ่มตัวอย่างแบบง่าย (Simple Random Sampling) จำนวน 30 คน ผลการสุ่มได้ดังนี้

2.4 งานวิจัยที่เกี่ยวข้อง

สมนึกและคณะ (2548) ได้ทำการศึกษาผลของอัตราการป้อนของข้าวเปลือกและความดันลูกยางต่อประสิทธิภาพของการกะเทาะข้าวเปลือก โดยพบว่าเมื่อกำหนดให้อัตราป้อนคงที่ เมื่อความดันลูกยางเพิ่มขึ้น เปอร์เซ็นการกะเทาะและเปอร์เซ็นต์ข้าวกลิ้งหักและร้าวเพิ่มขึ้น เมื่อกำหนดให้

พัศกรและคณะ (2548) ได้ทำการศึกษาผลของอุณหภูมิการเก็บรักษาและความชื้นของข้าวเปลือกที่มีผลต่อการเปลี่ยนแปลงคุณภาพการสี โดยข้าวเปลือกที่เก็บไว้ที่อุณหภูมิต่ำกว่าอุณหภูมิห้องจะมีคุณภาพของการสีข้าวต่ำกว่าข้าวเปลือกที่เก็บไว้ที่อุณหภูมิห้องการประยุกต์ใช้การ (Applying Data Mining)Chen, Chiu, Chang (2005) ศึกษาการประยุกต์ใช้การวิเคราะห์ข้อมูลเพื่อทำนายพฤติกรรมการซื้อสินค้าของลูกค้า โดยศึกษาข้อมูลการซื้อสินค้าของลูกค้าในช่วงเวลาหนึ่งปี และนำมาใช้ในการวิเคราะห์โดยเริ่มต้นแบ่งกลุ่มของลูกค้าโดยลูกค้าที่มาซื้อสินค้าบ่อยๆและซื้อสินค้า

Apriori algorithm Associate Rules ได้ 111 และ 78 กฎตามลำดับ จากนั้นเปรียบเทียบกฎระหว่างช่วงเวลาพบการนำไปใช้ในการปรับกลยุทธ์การตลาดเพื่อสอดคล้องกับการเปลี่ยนแปลง

Liao, Chen, Wu (2007) ศึกษาการประยุกต์ใช้ Apriori alorithm K-Means Data

mining เทคนิคสำหรับการขยายประเภทของสินค้าและยี่ห้อในร้านค้าปลีกของไต้หวันชื่อ Carrefour

การขยายประเภทของสินค้าเป็นการเพิ่มความหลากหลายให้กับยี่ห้อซึ่งวัตถุประสงค์ก็เพื่อแสวงหา รายย่อยเป็นสิ่งจำเป็นที่ต้องระบุ การขายปลีกจึงเป็นแหล่งข้อมูลที่จะได้มาซึ่ง Demand Chain รายย่อยเป็นสิ่งจำเป็นที่ต้องระบุ การขายปลีกจึงเป็นแหล่งข้อมูลที่จะได้มาซึ่ง Demand Chain Lia Associate Rules Clustering

Chen Lin (2007) สินค้าและจัดสรรพื้นที่บนชั้นแสดงสินค้า สินค้าที่อยู่บนชั้นวางสินค้าที่เหมาะสมและอยู่ใน สินค้าและจัดสรรพื้นที่บนชั้นแสดงสินค้า สินค้าที่อยู่บนชั้นวางสินค้าที่เหมาะสมและอยู่ใน Associate Rules Associate Rules Associate Rules Multi-level Associate Rules

Liao, Hsieh, Huang (2007)

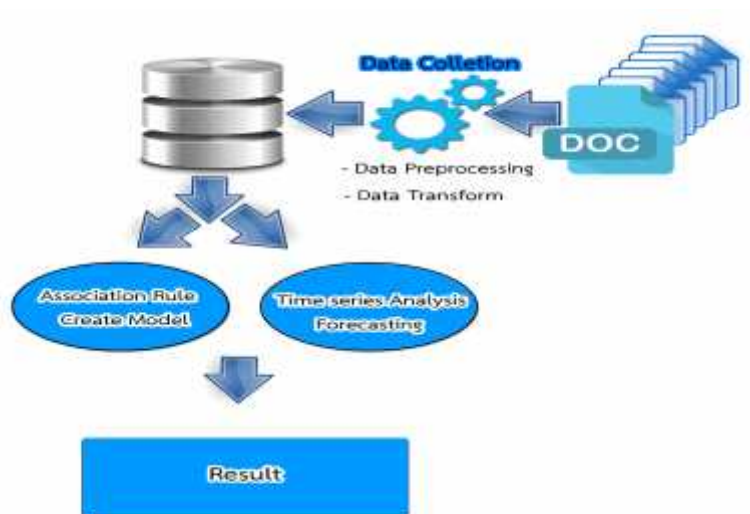
บทที่ 3

วิธีดำเนินงานวิจัย

ในบทนี้จะอธิบายถึงขั้นตอนและวิธีดำเนินงานในการหาความสัมพันธ์ของผลิตภัณฑ์ของฝากของที่ระลึกที่วางขายอยู่ตามสถานที่ท่องเที่ยวภายในจังหวัดเพชรบูรณ์ โดยมีการศึกษาเกี่ยวกับการทำเหมืองข้อมูล การหาความสัมพันธ์ของสินค้า และการวิเคราะห์รูปแบบการจัดทำโปรโมชั่นเพื่อออกวางขายตามฤดูกาลท่องเที่ยวให้สัมพันธ์กับปริมาณนักท่องเที่ยวในจังหวัดเพชรบูรณ์ตามหัวข้อต่อไปนี้

3.1 การศึกษาปัญหาและความเป็นไปได้

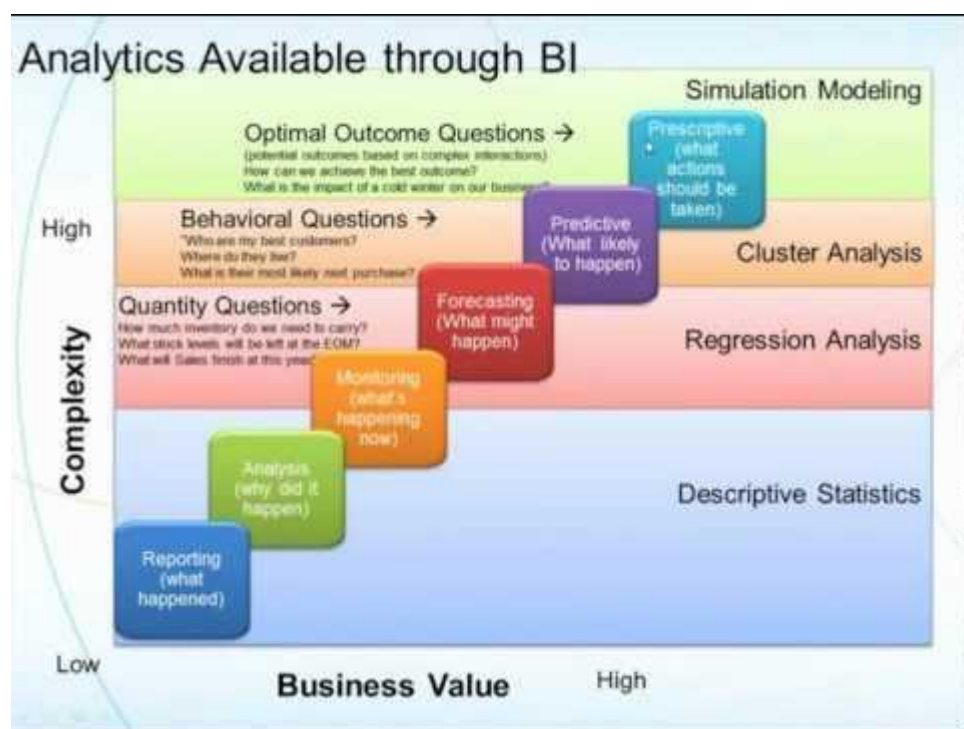
ขั้นตอนการศึกษาปัญหาและความเป็นไปได้ ผู้วิจัย



ภาพที่ 3-1 กรอบแนวคิด

จากภาพที่ 3-1 จากภาพที่ 1 เก็บข้อมูลจากใบเสร็จใบส่งของนำมาแปลงให้อยู่ในรูปแบบของตาราง โดยบันทึกการขายซื้อสินค้าในแต่ละรายการตามจำนวนการซื้อ ราคาขายต่อหน่วย เพื่อใช้ในการวิเคราะห์หาความสัมพันธ์ด้วย วิธี FP-Growth จากโปรแกรม Rapid Miner ซึ่งโปรแกรม

Rapid Miner จะหาค่า Support และค่า Confidence ที่ดีที่สุดของยอดขายสินค้าทั้งหมดเรียงลำดับจากมากไปหาน้อย ผลลัพธ์ที่ได้จากการหาความสัมพันธ์นี้ ผู้วิจัยได้เก็บความสัมพันธ์ทั้งหมดกว่า 5,000 รายการแต่เลือกเก็บความสัมพันธ์ที่ดีที่สุดเพื่อใช้ในการส่งเส (time series) Microsoft Excel 4 (time series) 10



ภาพที่ 3-2

3-2 การวิเคราะห์ข้อมูลเชิงลึก (Data Mining) เป็นการวิเคราะห์ข้อมูลเพื่อค้นหาความสัมพันธ์ที่ซ่อนอยู่หรือแนวโน้มในอนาคต ซึ่งสามารถช่วยในการตัดสินใจทางธุรกิจได้ ตัวอย่างเช่น การวิเคราะห์ข้อมูลการขายสามารถช่วยในการตัดสินใจว่าควรเพิ่มหรือลดการผลิตสินค้าชนิดใดบ้าง หรือการวิเคราะห์ข้อมูลลูกค้าสามารถช่วยในการตัดสินใจว่าควรเพิ่มหรือลดการบริการลูกค้าชนิดใดบ้าง

ตารางที่ 3-2 ข้อมูลราคาอาบน้ำของลูกค้า

ID	Price(Bath)
1	10
2	20
3	50
4	100

ตัวอย่างข้อมูล 3-2 ข้อมูลราคาอาบน้ำของลูกค้า

ตารางที่ 3-3 ข้อมูลจำนวนลูกค้า

ID	Count(ลูกค้า)
1	1
2	2
3	จำนวนลูกค้า 2 คน

ตัวอย่างข้อมูล 3-3 ข้อมูลจำนวนลูกค้า

3.3.1 การเตรียมข้อมูลด้วย Rapid Miner

การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

1. การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

1. การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

1. การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

1. การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

1. การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

1. การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

1. การเตรียมข้อมูลด้วย Rapid Miner มีขั้นตอนดังนี้

Id	outlook	windy	play
1	sunny	TRUE	No
2	sunny	TRUE	Yes
3	rainy	FALSE	Yes
4	sunny	TRUE	No

numeric nominal binominal

ภาพที่ 3-3 ประเภทของข้อมูลตามลักษณะการวัด

- Nominal ข้อมูลประเภท category (ประเภท) มีค่าเพียง 2 ค่า
- Binominal ข้อมูลประเภท category (ประเภท) มีค่า 2 ค่า
- numeric ข้อมูลตัวเลข
- Text ข้อมูลข้อความ

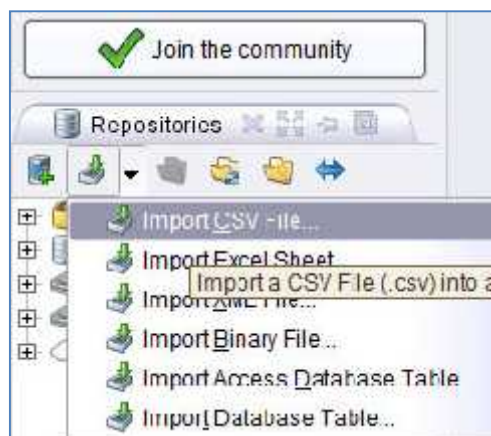
Reciept_id	height	weight	gender
1	170	65	male
2	194	106	male
3	168	61	female
4	170	65	female
5	176	74	male
6	188	98	male
7	185	87	male
8	170	58	female

ID

attribute

label

ภาพที่ 3-4 ประเภทของข้อมูลตามลักษณะการวัด



ภาพที่ 3-5 เมนูนำเข้าข้อมูลของ Rapid Miner

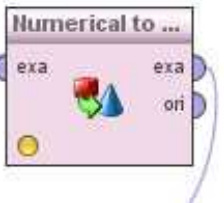
ภาพที่ 3-5 แสดงเมนูนำเข้าข้อมูลของ Rapid Miner ซึ่งผู้ใช้สามารถเลือกนำเข้าข้อมูลจากแหล่งต่าง ๆ ได้ เช่น CSV, Excel, Access Database Table, และ Database Table. การนำเข้าข้อมูลเหล่านี้จะช่วยให้ผู้ใช้สามารถวิเคราะห์ข้อมูลได้อย่างมีประสิทธิภาพมากขึ้น.

การนำเข้าข้อมูล (Attribute) เป็นขั้นตอนที่สำคัญในการเตรียมข้อมูลสำหรับการวิเคราะห์ข้อมูล. ผู้ใช้สามารถเลือกนำเข้าข้อมูลจากแหล่งต่าง ๆ ได้ เช่น CSV, Excel, Access Database Table, และ Database Table. การนำเข้าข้อมูลเหล่านี้จะช่วยให้ผู้ใช้สามารถวิเคราะห์ข้อมูลได้อย่างมีประสิทธิภาพมากขึ้น.

Repository เป็นที่เก็บข้อมูลที่ใช้ในการวิเคราะห์ข้อมูล. ผู้ใช้สามารถนำเข้าข้อมูลจากแหล่งต่าง ๆ ได้ เช่น CSV, Excel, Access Database Table, และ Database Table. การนำเข้าข้อมูลเหล่านี้จะช่วยให้ผู้ใช้สามารถวิเคราะห์ข้อมูลได้อย่างมีประสิทธิภาพมากขึ้น.

Repository เป็นที่เก็บข้อมูลที่ใช้ในการวิเคราะห์ข้อมูล. ผู้ใช้สามารถนำเข้าข้อมูลจากแหล่งต่าง ๆ ได้ เช่น CSV, Excel, Access Database Table, และ Database Table. การนำเข้าข้อมูลเหล่านี้จะช่วยให้ผู้ใช้สามารถวิเคราะห์ข้อมูลได้อย่างมีประสิทธิภาพมากขึ้น.

ตารางที่ 3-4 ขั้นตอนการประมวลผลข้อมูลเพื่อหาความสัมพันธ์

โอเปอเรเตอร์	คำอธิบาย
	Retrieve ข้อมูลจาก Repository Dataset ที่ใช้ Rapid Miner
	Numerical to Binominal Numerical Repository category True False
	FP-Growth category FP-Tree support FP-Tree
	Create Association Rules Support Confidence

				ค่าเฉลี่ยแต่ละไตรมาส		ค่าเฉลี่ยแต่ละไตรมาส		ค่าเฉลี่ยแต่ละไตรมาส	
				ไตรมาส		ไตรมาส		ไตรมาส	
year	Quarter	มูลค่าการขาย	ไตรมาส	ไตรมาส	ไตรมาส	ไตรมาส	ไตรมาส	ไตรมาส	ไตรมาส
1	2013	1	237						
2		2	131						
3		3	123						
4		4	146						
5	2014	1	231						
6		2	144						
7		3	181						
8		4	299						
9	2015	1							
10		2							
11		3							
12		4							

ภาพที่ 3-7 การคำนวณค่าเฉลี่ยแต่ละไตรมาส

เมื่อได้คำนวณค่าเฉลี่ยแต่ละไตรมาสเรียบร้อยแล้ว การคำนวณค่าเฉลี่ยแต่ละไตรมาส (linear regression) สามารถทำได้โดยใช้ Microsoft Excel

ตารางที่ 3-9 ข้อมูลค่าเฉลี่ยแต่ละไตรมาส

ปี	ไตรมาส	ค่าพยากรณ์มูลค่าการขาย
2559	1	5940.60
	2	3690.72
	3	5360.69
	4	7900.55

ตารางที่ 3-9 แสดงข้อมูลค่าเฉลี่ยแต่ละไตรมาส ซึ่งได้จากการคำนวณค่าเฉลี่ยแต่ละไตรมาส (linear regression) โดยใช้ Microsoft Excel

บทที่ 4

ผลและวิจารณ์ผลการวิจัย

ในการวิจัยครั้งนี้ ผู้วิจัยได้วิเคราะห์และสรุปข้อมูล โดยแบ่งการนำเสนอผลการวิเคราะห์ข้อมูลตามลำดับ ดังนี้

4.1 ผลการสร้างกฎความสัมพันธ์และการพยากรณ์อนุกรมเวลา

4.2 ผลการประเมินความพึงพอใจ

4.1 ผลการสร้างกฎความสัมพันธ์และการพยากรณ์อนุกรมเวลา

ตารางที่ 4-1 ผลลัพธ์กฎความสัมพันธ์

กฎความสัมพันธ์	ผลลัพธ์กฎความสัมพันธ์
กฎที่ 1 pro_001 + [pro_005+pro_006]	มะขามคลุกน้ำตาล ราคา 99 บาท ขายคู่กับ มะขามคลุกบ๊วย ราคา 29 บาทและมะขามจี๊ดจ๊าดราคา 29 บาท
กฎที่ 2 1 pro_001 + [pro_004+pro_009]	มะขามคลุกน้ำตาล ราคา 99 บาท ขายคู่กับ มะขามคลุกน้ำตาล ราคา 29 บาทและพวงกุญแจรูปฝักมะขามราคา 99 บาท
กฎที่ 3 1 pro_001 + [pro_004+pro_008]	มะขามคลุกน้ำตาล ราคา 99 บาท ขายคู่กับ มะขามคลุกน้ำตาล ราคา 29 บาทและมะขามอบแห้งราคา 99 บาท
กฎที่ 4 1 pro_001 + [pro_009+pro_008]	มะขามคลุกน้ำตาล ราคา 99 บาท ขายคู่กับ พวงกุญแจรูปฝัก มะขามราคา 99 บาทและมะขามอบแห้งราคา 99 บาท
กฎที่ 5 1 pro_001 + [pro_003+pro_005+pro_11]	มะขามคลุกน้ำตาล ราคา 99 บาท ขายคู่กับ มะขามคลุกบ๊วย ราคา 99 บาทและมะขามจี๊ดจ๊าดราคา 29 บาทและน้ำเปล่า ราคา 10 บาท

ตารางที่ 4-1 แสดงกฎความสัมพันธ์ของสินค้าขายดีที่สุดคือ มะขามคลุกน้ำตาล ราคา 99 บาท ซึ่งขายคู่กับสินค้าชนิดอื่นๆ

TIME SERIES METHODOLOGY

การพยากรณ์อนุกรมเวลา ข้อมูลที่ใช้ในการพยากรณ์อนุกรมเวลา ใช้ข้อมูลจากใบเสร็จใบสั่งของ ย้อนหลังตั้งแต่ 1 ม.ค. 2556- 31 ธ.ค. 2558 จากข้อมูลในใบเสร็จและใบสั่งของ รายการจำนวนสินค้า ซึ่งผู้วิจัยได้ทำการวิเคราะห์ข้อมูลด้วยการพยากรณ์อนุกรมเวลาเพียงหาช่วงเวลาที่มียอดขายสูงสุดของแต่ละปีดังนี้

ตารางที่ 4-2 มูลค่าการขายสินค้าเฉลี่ยต่อวัน

ปีที่เก็บข้อมูลย้อนหลัง	ไตรมาส	มูลค่าการขาย (บาท)
2556	1	437
	2	234
	3	323
	4	546
2557	1	531
	2	344
	3	484
	4	599
2558	1	735
	2	433
	3	511
	4	870

จากตาราง ที่ 4-2 แสดงรายการขายสินค้าย้อนหลังตั้งแต่ปี 2556 - 2558 ในแต่ละกลุ่มจะถูกแบ่งเป็น 4 ไตรมาสในแต่ละปี ซึ่งในแต่ละไตรมาสจะแบ่งตามจำนวนเดือนเท่ากันในแต่ละไตรมาส และบันทึกมูลค่าการขายสินค้าในแต่ละเดือน โดย การพยากรณ์อนุกรมเวลา มีสมการในการพยากรณ์ ซึ่งการพยากรณ์อนุกรมเวลาใช้หลักการพยากรณ์ดังนี้

ตารางที่ 4-3 ค่า Moving average ของสินค้า

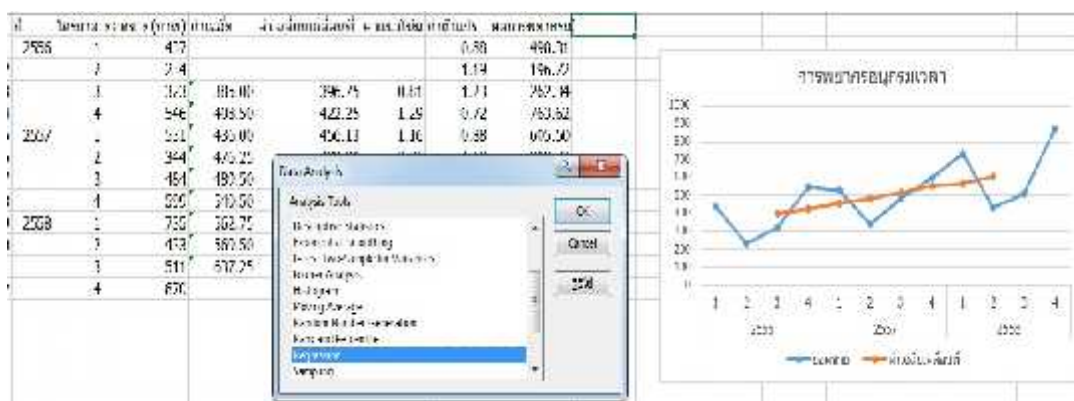
ดัชนีฤดูกาล	ปี	ไตรมาส	ยอดขาย(บาท)	ค่าเฉลี่ย	ค่าเฉลี่ยเคลื่อนที่
1	2556	1	437		
2		2	234		
3		3	323	385.00	396.75
4		4	546	408.50	422.25
5	2557	1	531	436.00	456.13
6		2	344	476.25	482.88
7		3	484	489.50	515.00
8		4	599	540.50	551.63
9	2558	1	735	562.75	566.13
10		2	433	569.50	603.38
11		3	511	637.25	
12		4	870		

จากตารางที่ 4-3 แสดงค่า Moving average จากการนำค่าเฉลี่ยของยอดขายแต่ละไตรมาส มาหาค่า จุดกึ่งกลางแนวโน้มเฉลี่ยเคลื่อนที่ เพื่อทดสอบหาค่า ดัชนีการพยากรณ์ต่อฤดูกาล เพื่อทำการพยากรณ์ยอดขายไตรมาสต่อไป ของสินค้า



ภาพที่ 4-1 กราฟแสดงค่าแนวโน้มจากการหาค่าเฉลี่ยเคลื่อนที่

จากภาพที่ 4-1 แสดงให้เห็นว่าข้อมูลที่ใช้ในการพยากรณ์กรรมเวลานั้นมีลักษณะการเคลื่อนที่ของข้อมูลในรูปแบบใด หลังจากนั้นนำค่าที่ได้จากการหาค่าแนวโน้มเฉลี่ยไปใช้กับสูตรการคำนวณค่าเฉลี่ยเคลื่อนที่ (simple linear regression analysis) โดยใช้โปรแกรม Microsoft Excel ในการคำนวณ



ภาพที่ 4-2 การใช้สมการถดถอยเพื่อพยากรณ์ยอดขายปี 2559

ข้อ 4-2 เมื่อได้คำนวณหาค่าจุดกึ่งกลางค่าแนวโน้มเฉลี่ยเคลื่อนที่จากการคำนวณ
 ข้อ 4-1 แล้วให้นำค่าจุดกึ่งกลางค่าแนวโน้มเฉลี่ยเคลื่อนที่ค่า 2559 ไปคำนวณค่า
 (linear regression) โดยใช้ Microsoft Excel ดังตัวอย่างต่อไปนี้

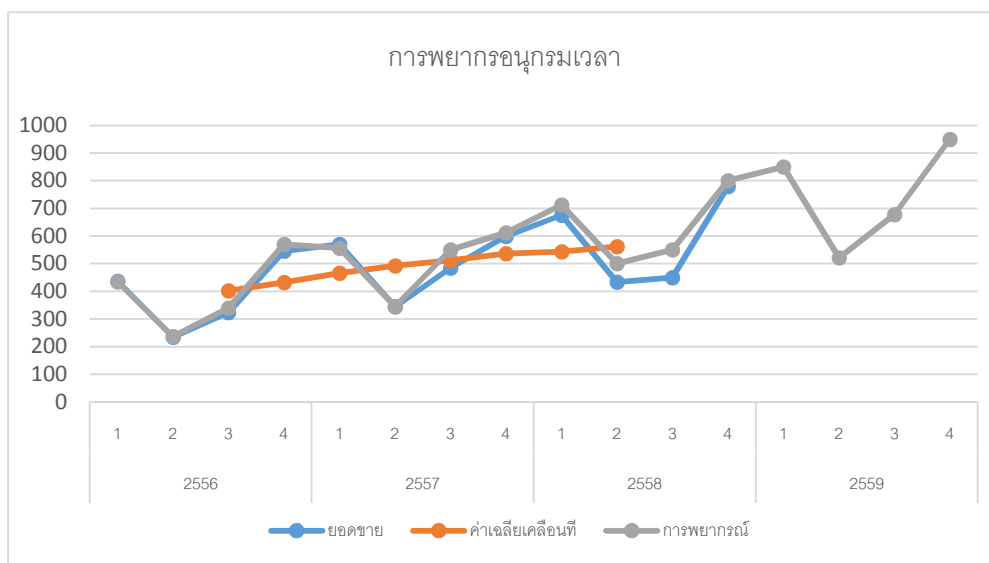
ข้อ 4-4 ตารางค่าการขายสินค้าปี 2559

ปี	ไตรมาส	ค่าการขายสินค้า
2559	1	850.00
	2	521.00
	3	678.00
	4	950.00

จากตารางค่าการขายสินค้าปี 2559 แสดงให้เห็นว่ายอดขายสินค้าในไตรมาสที่ 1 และไตรมาสที่ 4 จากกราฟที่แสดงยังทำนายยอดขายสินค้าในปี 2559 ว่ามียอดขายที่สูงขึ้น

ข้อ 4-5 ตารางค่าการขายสินค้าปี 2558-2559

t	ปี	ไตรมาส	ค่าการขายสินค้า (บาท)	ค่าการขายสินค้าปี 2558	ค่าการขายสินค้าปี 2559	ค่าการขายสินค้าปี 2560	ค่าการขายสินค้าปี 2561	ค่าการขายสินค้าปี 2562	ค่าการขายสินค้าปี 2563	ค่าการขายสินค้าปี 2564
1	2556	1	437				0.87	499.76	358.51	435.00
2		2	234				1.19	196.63	390.50	237.00
3		3	323	385.00	401.63	0.80	1.23	262.02	422.49	340.00
4		4	546	418.25	432.00	1.26	0.73	743.40	454.49	570.00
5	2557	1	570	445.75	465.88	1.22	0.87	651.86	486.48	556.00
6		2	344	486.00	492.63	0.70	1.19	289.06	518.48	345.00
7		3	484	499.25	512.38	0.94	1.23	392.63	550.47	550.00
8		4	599	525.50	536.63	1.12	0.73	815.56	582.47	612.00
9	2558	1	675	547.75	543.50	1.24	0.87	771.93	614.46	712.00
10		2	433	539.25	561.88	0.77	1.19	363.85	646.46	501.00
11		3	450	584.50			1.23	365.04	678.45	550.00
12		4	780				0.73	1061.99	710.45	800.00
13	2559	1					0.87		742.44	850.00
14		2					1.19		774.44	521.00
15		3					1.23		806.43	678.00
16		4					0.73		838.43	950.00



ภาพที่ 4-3 กราฟแสดงค่าพยากรณ์การขายปี 2559

ภาพที่ 4-3 จะเห็นว่า เส้นกราฟที่ได้จากการขายสินค้า มีการเคลื่อนที่ขึ้นลงตามไตรมาสในแต่ละปี และเส้นกราฟที่แสดงค่าแนวโน้มเฉลี่ยเคลื่อนที่แสดง ค่าแปรผันของฤดูกาลในแต่ละปี นั้นหมายความว่า การอนุกรมเวลาของการขายสินค้า มีการเคลื่อนที่แปรผันตามฤดูกาล

2559

2559

4.2 ผลการวิเคราะห์ข้อมูลเชิงพรรณนา

ผลการวิเคราะห์ข้อมูลทั่วไปเกี่ยวกับ สถานะ เพศ อายุ และอาชีพ ของผู้ตอบ
แบบสอบถาม

4-6 ผลการวิเคราะห์ข้อมูลเชิงพรรณนาเกี่ยวกับ สถานะ

สถานะ	จำนวนผู้ตอบแบบสอบถาม	ร้อยละ
โสด	80	32.33
สมรส	170	67.66
รวม	250	100

4-6 ผลการวิเคราะห์ข้อมูลเชิงพรรณนาเกี่ยวกับ สถานะ

โสด 32.33 เปอร์เซ็นต์ สมรส 67.66

4-7 ผลการวิเคราะห์ข้อมูลเชิงพรรณนาเกี่ยวกับ เพศ

เพศ	จำนวนผู้ตอบแบบสอบถาม	ร้อยละ
ชาย	85	34.14
หญิง	165	66.66
รวม	250	100

4-7 ผลการวิเคราะห์ข้อมูลเชิงพรรณนาเกี่ยวกับ เพศ

ชาย 34.14 เปอร์เซ็นต์

4-8 ผลการวิเคราะห์ข้อมูลเชิงพรรณนาเกี่ยวกับ อายุ

อายุ	จำนวนผู้ตอบแบบสอบถาม	ร้อยละ
20-25 ปี	56	22.66
26-30 ปี	54	20
31-35 ปี	77	30.83
36-40 ปี	33	13.33
41-45 ปี	30	12.33

20-25 21.66 26-30 20 31-35 30.83 36-40 13.83 41-45 12.33

4-9

	45	18.5
	35	14.16
	57	22.83
	113	45.16
	10	0.04

4-9 18.5 14.16 22.83 45.16 0.04

表 4-11 不同处理组在不同处理时间下对土壤微生物多样性的影响

处理组	处理前		处理后		t	p
	$\bar{x}$	S.D	$\bar{x}$	S.D		
1. 对照处理组	4.10	0.41	4.38	0.31	1.08	0.28
2. 处理组	4.56	0.43	4.65	0.33	1.71	0.90
3. 处理组	4.01	0.42	3.78	0.62	-1.58	0.11
4. 处理组	3.98	0.44	3.68	0.46	0.79	0.43
5. 处理组	4.46	0.41	4.39	0.43	1.04	0.25
6. 处理组	4.01	0.44	3.89	0.46	0.89	0.67
7. 处理组	4.14	0.43	4.05	0.41	-0.89	0.74
8. 处理组	3.97	0.45	4.12	0.47	0.98	0.75
总计	4.17	0.46	4.17	0.45	0.82	0.46

表 4-11 不同处理组在不同处理时间下对土壤微生物多样性的影响

不同处理组在不同处理时间下对土壤微生物多样性的影响

บทที่ 5

อภิปรายผลและข้อเสนอแนะ

5.1 ผลการวิเคราะห์ แบ่งกลุ่มข้อมูล สินค้า สามารถสรุปได้ดังนี้

ผลจากการหาความสัมพันธ์พบว่า เมื่อร้านค้าขายของฝากของที่ระลึกในแหล่งท่องเที่ยวของจังหวัดเพชรบูรณ์จัดโปรโมชั่นส่งเสริมการขายตาม กฎความสัมพันธ์โดย จัดวางสินค้าตามกฎความสัมพันธ์ที่ดีที่สุด 5 กฎความสัมพันธ์ พบว่า ผู้บริโภคมักซื้อสินค้าที่มีราคาต่ำกว่า 100 ร้อยบาทมากที่สุด จากสินค้าที่มีค่า Support สูงที่สุด และเมื่อจัดวางสินค้าตาม ตามกฎความสัมพันธ์ พบว่า ผู้บริโภค และจะซื้อสินค้าที่รวมกันแล้วไม่เกิน 500 บาทมากที่สุด ในที่นี้พบว่าเมื่อนำสินค้าที่มีราคาสูงกว่า 500 บาทนำมาจัดรายการลดราคา จะทำให้สินค้านั้นมียอดขายเพิ่มขึ้น 10% และเมื่อมีการนำสินค้า ราคาไม่เกิน 100 บาทมาขายคู่กับสินค้าที่ราคา ไม่เกิน 50 บาทจะทำให้สินค้านั้นมียอดขายเพิ่มขึ้น 15% และเมื่อถึงช่วงไตรมาสที่ 1 และที่ 4 ซึ่งผลลัพธ์ของการพยากรณ์พบว่าจะมีมูลค่าการขายสูงสุดของปี เมื่อร้านค้านำสินค้าที่นำมาลดราคาไม่เกิน 20% ของราคาขาย จะมียอดขายเพิ่มขึ้น 250% ของไตรมาส 1 และ 4 ในแต่ละปี ตามการพยากรณ์อนุกรมเวลา

สำหรับผู้วิจัยได้วิเคราะห์กฎความสัมพันธ์ของสินค้าประเภทของฝากของที่ระลึก ทั้งนี้ข้อมูลจาก ใบเสร็จ ใบสั่งของ ยังสามารถแบ่งกลุ่มสินค้า จากราคาขายต่อหน่วย รูปทรงและการจัดวางสินค้าภายในร้านค้า เพื่อเป็นทางเลือกในการนำเสนอ โปรโมชั่น และการส่งเสริมการที่หลากหลายมากขึ้น

สำหรับผู้วิจัยได้วิเคราะห์กฎความสัมพันธ์ของสินค้าประเภทของฝากของที่ระลึก ทั้งนี้ข้อมูลจาก ใบเสร็จ ใบสั่งของ ยังสามารถแบ่งกลุ่มสินค้า จากราคาขายต่อหน่วย รูปทรงและการจัดวางสินค้าภายในร้านค้า เพื่อเป็นทางเลือกในการนำเสนอ โปรโมชั่น และการส่งเสริมการที่หลากหลายมากขึ้น

5.2 ข้อเสนอแนะงานวิจัย

5.2.1 ควรมีการแบ่งกลุ่มยอดการขาย ตามสถานที่การขาย เช่นแหล่งชุมชน แหล่งสถานที่ท่องเที่ยว หรือ แหล่งของฝาก เป็นต้น

5.2.2 เนื่องจากจังหวัดเพชรบูรณ์เป็นสถานที่ท่องเที่ยวจึงมีนักท่องเที่ยวจำนวนมากแต่ไม่เท่ากันตลอดทั้งปี ควรมีการจัดสรรแนวทางในการบริหารจัดการแหล่งท่องเที่ยวให้สามารถมีปริมาณนักท่องเที่ยวที่สูงตลอดทั้งปี

5.2.3 ควรให้ความรู้แก่ ผู้ประกอบการ โรงแรม ทั้งร้านค้า ร้านอาหาร ร้านกาแฟ นำภูมิความสัมพันธที่ไ้จากการวิจัยไปใช้เป็นแนวทางในการส่งเสริมการขายของผู้ประกอบการนั้นๆ เนื่องจากเข้าถึงข้อมูลข่าวสารและเทคโนโลยี เกี่ยวกับงานวิจัยให้มากขึ้น เนื่องจากมีผู้ทำการวิจัย แต่ขาดการนำไปใช้จริงจำนวนมาก

5.3 ข้อเสนอแนะงานวิจัยครั้งต่อไป

ควรมีการวิเคราะห์ข้อมูลการขายแบบไตรมาสในรายปีที่มีการพยากรณ์ในรูปแบบของทำนายแบบไม่มีดัชนีฤดูกาล เพื่อ หาช่วงเวลาการขายดีที่สุดของในแต่ละปีเพื่อเสนอโปรโมชั่นที่ตรงใจผู้บริโภค

บรรณานุกรม

- กฤษณพร สุริยะบรรเทิงม และดร.กมลเกียรติ เรืองกมลลา. การสร้างแบบจำลองผลิตภัณฑ์ประกันภัยให้กับผู้สูงอายุกลุ่มบัญชีออมทรัพย์โดยการทำเหมืองข้อมูล. วิทยาลัยนวัตกรรมการมหาวิทยาลัยธรรมศาสตร์, 2555.
- กัลยา วานิชย์บัญชา. การวิเคราะห์ข้อมูลหลายตัวแปร. พิมพ์ครั้งที่ 4. กรุงเทพมหานคร . ภาควิชาสถิติคณะพาณิชยศาสตร์และการบัญชี จุฬาลงกรณ์มหาวิทยาลัย, 2552.
- โกเมศ อัมพวัน. วิธีการหาความสัมพันธ์แบบใหม่โดยต้นไม้แสดงรายการความถี่. วิทยานิพนธ์วิทยาศาตรมหาบัณฑิต สาขาวิชาวิทยาศาสตร์คอมพิวเตอร์ คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย, 2552.
- ริม โอทกานนท์. การตลาดบัตร KTC [ออนไลน์]. เข้าถึงได้จาก : http://inside.cm.mahidol.ac.th/mkt/index.php?option=com_content&view=article&id=144:-ktc&catid=1:mk-articles&Itemid=11.
- มหาวิทยาลัยเกษตรศาสตร์ วิทยาเขตกำแพงแสน. ระบบสารสนเทศเพื่อการจัดการ ภูมิศึกษา : ระบบประมวลผลภาพใบสั่งของกรรมการขนส่งในกรุงนิวยอร์ก
- ชฎารัตน์ พิพัฒน์นันท์. การวิเคราะห์ข้อมูลทางเศรษฐกิจและสังคมของครัวเรือนด้วยวิธีแบ่งกลุ่มและหาความสัมพันธ์ สำหรับการหาเหมืองข้อมูล. คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยหอการค้าไทย, 2552.
- คำย เชิงจี. ทฤษฎีทดสอบและวัดผลการศึกษา. ภาควิชาประเมินผลและวิจัยการศึกษา: คณะศึกษาศาสตร์ มหาวิทยาลัยเชียงใหม่, 2552.
- ณรงค์ศักดิ์ คงทิม และจิรัฐา ภูบุญชอบ. การประยุกต์ใช้เอฟพี-กโรทกับงานแนะแนวการศึกษาต่อในระดับอุดมศึกษา. CIT2011 & UniNOMS. 21 (กันยายน 2550). 4-15.
- น้ำทิพย์ สมศรีไทย. การพัฒนาห้องสมุดประชาชนโดยใช้ภูมิความรู้ผ่านระบบอิเล็กทรอนิกส์. ภาควิชาเทคโนโลยีสารสนเทศ คณะวิทยาศาสตร์และเทคโนโลยี วิทยาลัยราชพฤกษ์, 2554

ปราณี มณีรัตน์. การสร้างโมเดลการจัดการระบบนักศึกษาสัมพันธ์โดยการทำเหมืองข้อมูล.

สาขาวิทยาการคอมพิวเตอร์:คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยศรีปทุม, 2554
 ปริญญ์ สมักรการ และกมล เกียรติเรืองกมล. การสร้างแบบจำลองรายการธุรกรรมผิดปกติของ
 บัญชีเงินฝากออมทรัพย์ผู้สูงวัยโดยการทำเหมืองข้อมูล กรณีศึกษาธนาคารพาณิชย์แห่งหนึ่ง.
 วิทยาลัยนวัตกรรมการบริหาร มหาวิทยาลัยธรรมศาสตร์, 2552.

พิจิตรา จอมศรี. การทำนายเนื้อหาของเว็บโดยเทคนิคเหมืองข้อมูล. วิทยานิพนธ์วิทยาศาตร
 มหบัณฑิต สาขาวิชาวิทยาการคอมพิวเตอร์ ภาควิชาคอมพิวเตอร์ บัณฑิตวิทยาลัย มหาวิทยาลัย
 ศิลปากร, 2552

พีรภาวี เกียรติเฉลิมคุณ. การพัฒนาแบบจำลองสำหรับการพยากรณ์ยอดขายด้วยการทำเหมือง
 ข้อมูล.กรณีศึกษาบริษัท ที สไตล์โปรดักส์ จำกัด. วิทยานิพนธ์วิทยาศาตรมหาบัณฑิต สาขาวิชา
 บริหารเทคโนโลยี วิทยาลัยนวัตกรรมการบริหาร มหาวิทยาลัยธรรมศาสตร์, 2553.

ฟูไคละห์ ลือมอญ. ขั้นตอนวิธีสำหรับการค้นหาข้อมูลที่ปรากฏร่วมกันบ่อยโดยรองรับรายการ
 ข้อมูลที่คล้ายคลึงกัน. วิทยานิพนธ์วิทยาศาตรมหาบัณฑิต สาขาวิทยาการคอมพิวเตอร์
 มหาวิทยาลัยสงขลานครินทร์, 2553.

ภาวิณี เหล่าพิพัฒน์ไพบุลย์และพาชิตชนัด ศิริพานิช. การจัดกลุ่มพฤติกรรมการใช้บัตรเดบิตและ
 บัตรกดเงินสดของคนทำงานในกรุงเทพมหานคร. การประชุมเสนอผลงานระดับ
 บัณฑิตศึกษา

มหาวิทยาลัยสุโขทัยธรรมมาธิราช ครั้งที่ 2, 2555.

สมโภชน์ ศรีสมุทร. การวิเคราะห์การจัดกลุ่มโรงเรียนตามมาตรฐานการศึกษาเพื่อการประเมิน
 คุณภาพภายนอกระดับการศึกษาขั้นพื้นฐาน ของโรงเรียนในสามจังหวัดชายแดนภาคใต้.
 วิทยานิพนธ์ปริญญาวิทยาศาตรมหาบัณฑิต สาขาวิชาการวัดผลและวิจัยการศึกษา
 มหาวิทยาลัยสงขลานครินทร์, 2553.

อดุลย์ ยิมงาม. การทำเหมืองข้อมูล Data Mining[ออนไลน์]. เอกสารการเรียนการสอนรายวิชา
 801405

Data Mining Chapter 1Introduction. มิถุนายน, 2554www.open-miner.com . Introduction

to Data Mining [ออนไลน์]. เข้าถึงได้จาก : <http://opeminer.com/2009/11/03/introduction-datamining>.

อรอุมา นองเนื่อง และณัฐวิ อุดกฤษฎ์. ระบบช่วยวิเคราะห์บริการทางการเงินเพื่อกลุ่มลูกค้านิติบุคคลกรณีศึกษาธนาคารกสิกรไทย. วิทยานิพนธ์วิทยาศาตรมหาบัณฑิต ภาควิชาการจัดการเทคโนโลยีสารสนเทศ คณะเทคโนโลยีสารสนเทศ มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ, 2550.

อรชุน ฟองประไพ และรศ.ดร.ประสาร บุญเสริม. การพยากรณ์อุปสงค์น้ำมันปาล์มดิบในประเทศไทย.วารสารการวิจัยทางธุรกิจ และการบริหาร ปีที่ 2 ฉบับที่ 1 มกราคม – มิถุนายน 2557

อำนาจ มณีศรีวงศ์กุล.2551.Cluster Analysis . วารสารวิจัย มหาวิทยาลัยขอนแก่น.

อัจฉราพร สีหิรัญวงศ์และรณชัย คงสกนธ์. แบบวัด Hamilton Rating Scale for Depression : การวิเคราะห์การรวมกลุ่ม. วารสารสมาคมจิตแพทย์แห่งประเทศไทย, 2551.

เอกสิทธิ์ พัทธวงศ์ศักดิ์ .การวิเคราะห์ข้อมูลด้วยเทคนิคคตาต้า ไมน์นิง เบืองตัน.พิมพ์ครั้งที่ 2.

กรุงเทพฯ: เอเชีย ดิจิตอลการพิมพ์ จำกัด

Aiman Moyaid Said,Dr. P D D. Dominic,Dr. Azween B Abdullah,2014. A Comparative Study of FP-growth Variations.2015.

B.Santhosh Kumar.K.V.Rukman.Implementation of Web Usage Mining UsingAPRIORI and FP-Growth Algorithms.2013.

Jiawei Han, Micheline Kamber and Jian Pei.Data Mining: Concepts and Techniques, 3rd ed.2015

Joseph F. Hair, Jr. , William C. Black , Barry J. Babin , Rolph E. Anderson . Multivariate Data Analysis . A GLOBAL PERSPECTIVE .2012

Johnson , R.A. and Wichern , D.W. Applied Multivariate Statistical Analysis 3rd ed. Englewood Cliffs : Prentice-Hall .1999.

Norusis , M.J . SPSSx Advance Statistics Guide . New York : McGraw-Hill.2010.

Onuma Nongnuang, Nattavee Utakrit,” Analysis System for Corporate Customer:

CaseStudy of Product Cash Management Group for Kasikorn Bank” The 6TH National

Conference on Computing and Information Technology, NCCIT2010-154

Preyanush Samakkan¹ and Kamol Keatruangkamala²,” Modeling that unusual saving deposits transaction of the elderly using data mining process: case study of one commercial bank”2012 ,pp40-69Education Study”, CIT2011 & UniNOMS2011

Zolt'an Prekopcs'ak,G'abor Makrai,Tam'as Henk,Csaba,” Radoop Analyzing Big Data with RapidMiner and Hadoop”,2009,pp 1-69

ภาคผนวก

ภาคผนวก ก

ภาพประกอบงานวิจัย

