

การรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท
และนิเวศเน็ตเวิร์ก

นายอุดม สถาพรชัยสิทธิ์

วิทยานิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตรปริญญาวิศวกรรมศาสตรมหาบัณฑิต
สาขาวิชาวิศวกรรมคอมพิวเตอร์ ภาควิชาวิศวกรรมคอมพิวเตอร์
คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย
ปีการศึกษา 2549
ลิขสิทธิ์ของจุฬาลงกรณ์มหาวิทยาลัย

THAI OPTICAL CHARACTER RECOGNITION USING MULTI-CLASS PRINCIPAL
COMPONENTS ANALYSIS AND NEURAL NETWORKS

Mr. Udom Sathapornchaiyasit

A Thesis Submitted in Partial Fulfillment of the Requirements
for the Degree of Master of Engineering Program in Computer Engineering

Department of Computer Engineering

Faculty of Engineering

Chulalongkorn University

Academic Year 2006

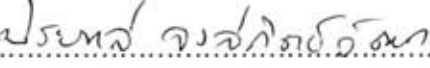
Copyright of Chulalongkorn University

หัวข้อวิทยานิพนธ์	การรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบ
	สำคัญแบบหลายประเภทและนิเวศเน็ตเวิร์ก
โดย	นายอุดม สถาพรชัยสิทธิ์
สาขาวิชา	วิศวกรรมคอมพิวเตอร์
อาจารย์ที่ปรึกษา	รองศาสตราจารย์ ดร.บุญเสริม กิจศิริกุล

คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย อนุมัติให้แนบวิทยานิพนธ์ฉบับนี้
เป็นส่วนหนึ่งของการศึกษาตามหลักสูตรปริญญาโท

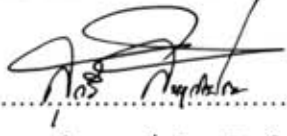

..... คณบดีคณะวิศวกรรมศาสตร์
(ศาสตราจารย์ ดร.ติเรก ลาวัณยศิริ)

คณะกรรมการสอบวิทยานิพนธ์


..... ประธานกรรมการ
(รองศาสตราจารย์ ดร.ประภาส จงสิตย์วัฒนา)


..... อาจารย์ที่ปรึกษา
(รองศาสตราจารย์ ดร.บุญเสริม กิจศิริกุล)


..... กรรมการ
(ผู้ช่วยศาสตราจารย์ ดร.นุชา ลิ้มปิยะกรณ)


..... กรรมการ
(อาจารย์ ดร.สุกรี ลิ้นอุทัย)

นายอุดม สถาพรชัยสิทธิ์ : การรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบ
สำคัญแบบหลายประเภทและนิวรอลเน็ตเวิร์ก (THAI OPTICAL CHARACTER
RECOGNITION USING MULTI-CLASS PRINCIPAL COMPONENT ANALYSIS
AND NEURAL NETWORKS) อ. ที่ปรึกษา: รศ.ดร.บุญเสริม กิจศิริกุล, 60 หน้า.

วิทยานิพนธ์ฉบับนี้มีวัตถุประสงค์เพื่อนำเสนอวิธีการรู้จำตัวอักษรภาษาไทยแบบใหม่
เรียกว่าการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทสำหรับการรู้จำตัวอักษรภาษาไทย
โดยวิธีการนี้มีแนวคิดพื้นฐานมาจากการวิเคราะห์องค์ประกอบสำคัญดั้งเดิม ซึ่งการวิเคราะห์
องค์ประกอบสำคัญดั้งเดิมมีข้อดีที่สามารถลดปริมาณข้อมูลทำให้ข้อมูลที่ได้มีขนาดกะทัดรัดโดย
ใช้กระบวนการแปลงเชิงเส้นและการตัดลดคุณลักษณะที่ไม่สำคัญออก อย่างไรก็ตามการวิเคราะห์
องค์ประกอบสำคัญดั้งเดิมนี้อย่างขาดประสิทธิภาพในการแบ่งแยกข้อมูลที่มีประเภทที่มีจำนวนมาก
ดังเช่นตัวอักษรภาษาไทย ส่วนการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทที่นำเสนอนี้มี
ประสิทธิภาพในการแบ่งแยกข้อมูลโดยการสร้างเซตขององค์ประกอบสำคัญโดยที่แต่ละเซตสร้าง
จากข้อมูลในแต่ละประเภท แต่การสร้างเซตขององค์ประกอบสำคัญหลายเซตนั้นมีจุดอ่อนตรงที่
ต้องใช้ทรัพยากรและเวลาในการคำนวณมากเกินไปซึ่งเป็นการสิ้นเปลืองและเสียเวลาจึงต้องมี
วิธีการในการลดจำนวนองค์ประกอบสำคัญให้น้อยลง ดังนั้นในวิทยานิพนธ์ฉบับนี้จึงนำเสนอ
วิธีการสำหรับการกำหนดจำนวนองค์ประกอบสำคัญที่เหมาะสมสำหรับข้อมูลแต่ละประเภทที่
แตกต่างกัน 4 วิธีด้วยกัน โดยผลการทดลองแสดงให้เห็นว่าวิธีการที่นำเสนอนี้ให้ความถูกต้องใน
การรู้จำตัวอักษรที่สูงขึ้นกว่าวิธีการวิเคราะห์องค์ประกอบสำคัญดั้งเดิม

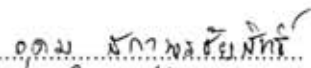
ภาควิชา.... วิศวกรรมคอมพิวเตอร์.....ลายมือชื่อนิสิต.....อุดม.....สาขาวิชา.....
สาขาวิชา....วิศวกรรมคอมพิวเตอร์.....ลายมือชื่ออาจารย์ที่ปรึกษา.....
ปีการศึกษา2549.....

4670662621 : MAJOR COMPUTER ENGINEERING

KEY WORD: PRINCIPLE COMPONENTS ANALYSIS / NEURAL NETWORKS

UDOM SATHAPORNCHAIYASIT : THAI OPTICAL CHARACTER RECOGNITION
USING MULTI-CLASS PRINCIPAL COMPONENT ANALYSIS AND NEURAL
NETWORKS. THESIS ADVISOR : ASSOC. PROF. BOONSERM KIJSIRIKUL,
Ph.D., 60 pp.

The objective of this thesis is to present a novel method, called multi-class principal components analysis (MCPCA), for Thai optical character recognition(Thai OCR). The method is based on the original principal components analysis (PCA) for Thai OCR. The original PCA reduces the original data to smaller size data. It has the advantage of compact representation of data by linear transformation and reduction of unimportant features. However, PCA lacks of discriminative power. The proposed MCPCA possesses strong discriminative power by constructing several sets of principal components, each for one class of data. Each set is the representative for the corresponding class of data. However, constructing several sets of principal components consumes lot of resources and computational time. We propose four methods for determining the appropriate number of principal components for each class. The experimental results show that our proposed methods provide higher accuracy than the original PCA for Thai OCR.

Department.... Computer Engineering.... Student's signature.....
Field of study.... Computer Engineering...Advisor's signature.....
Academic year ...2006.....

กิตติกรรมประกาศ

ขอกราบขอบคุณ รองศาสตราจารย์ ดร.บุญเสริม กิจศิริกุล อาจารย์ที่ปรึกษา ที่ได้ให้การ
อบรมวิชาการ และให้คำชี้แนะในการวิจัย จนทำให้วิทยานิพนธ์ฉบับนี้สำเร็จลุล่วงได้เป็นอย่างดี

ขอขอบคุณคณะกรรมการสอบวิทยานิพนธ์ทุกท่านที่ได้สละเวลาตรวจสอบและให้
ข้อเสนอแนะอันเป็นประโยชน์ในการทำวิจัย

ขอขอบคุณสมาชิกห้องปฏิบัติการอัจฉริยภาพเครื่องจักรและการค้นพบความรู้ (MIND
LAB) ทุกท่านที่ได้ให้คำแนะนำในการแก้ปัญหาในหลายเรื่องในงานวิจัย และส่งเสริมเกื้อกูลให้
งานวิจัยสำเร็จเป็นอย่างดี

ขอกราบขอบพระคุณบิดา ที่ได้ให้การสนับสนุนในด้านต่างๆ แก่ข้าพเจ้าจนประสบ
ความสำเร็จได้ในวันนี้

สารบัญ

หน้า

บทคัดย่อภาษาไทย	ง
บทคัดย่อภาษาอังกฤษ.....	จ
กิตติกรรมประกาศ.....	ฉ
สารบัญ	ช
สารบัญภาพ.....	ฌ
สารบัญตาราง.....	ญ
บทที่ 1 บทนำ.....	1
1.1 ความเป็นมาและความสำคัญของปัญหา	1
1.2 วัตถุประสงค์	2
1.3 ขอบเขตของการวิจัย	2
1.4 ขั้นตอนและวิธีดำเนินงานวิจัย.....	2
1.5 ประโยชน์ที่คาดว่าจะได้รับ	3
1.6 ลำดับการจัดเรียงเนื้อหาในวิทยานิพนธ์	3
1.7 ผลงานที่ตีพิมพ์จากวิทยานิพนธ์.....	3
บทที่ 2 ทฤษฎีและงานวิจัยที่เกี่ยวข้อง	4
2.1 ทฤษฎีที่เกี่ยวข้อง.....	4
2.1.1 การวิเคราะห์องค์ประกอบสำคัญ (Principle components analysis)	4
2.1.2 เกณฑ์การเลือกองค์ประกอบสำคัญ	5
2.1.3 นีวรอลเน็ตเวิร์ก	8
2.2 งานวิจัยที่เกี่ยวข้อง	11
บทที่ 3 การออกแบบและพัฒนาการรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบ สำคัญแบบหลายประเภทและนีวรอลเน็ตเวิร์ก	13

3.1 โครงสร้างของระบบ	13
3.1.1 ขั้นตอนการเรียนรู้.....	13
3.1.2 ขั้นตอนการรู้จำ	17
3.2 เกณฑ์การเลือกองค์ประกอบสำคัญ	18
3.2.1 เกณฑ์กำหนดล่วงหน้า.....	18
3.2.2 เกณฑ์เปอร์เซ็นต์ของความแปรปรวน.....	19
3.2.3 เกณฑ์การทดสอบของเศษหินขอบภูเขา.....	20
3.2.4 เกณฑ์การทดสอบการแบ่งกลุ่ม.....	23
3.3 ข้อกำหนดต่างๆ ของนิรวลเน็ตเวิร์ก	26
บทที่ 4 การทดลอง	28
4.1 ชุดข้อมูลที่ใช้ในการทดลอง	28
4.2 วิธีการทดลอง	29
4.2.1 ขั้นตอนการเรียนรู้.....	29
4.2.2 ขั้นตอนการรู้จำ	30
4.3 ผลการทดลอง.....	31
4.4 วิเคราะห์ผลการทดลอง.....	38
บทที่ 5 สรุปผลการวิจัยและข้อเสนอแนะ.....	41
5.1 สรุปผลการวิจัย	41
5.2 ข้อเสนอแนะ	42
รายการอ้างอิง.....	43
ประวัติผู้เขียนวิทยานิพนธ์	60

สารบัญภาพ

หน้า

รูปที่ 2.1 ตัวอย่างกราฟของการทดสอบแบบกองเศษหินที่ขอบภูเขา	7
รูปที่ 2.2 ตัวอย่างกราฟของการทดสอบแบบกองเศษหินที่ขอบภูเขาที่จุดเปลี่ยนความชันไม่ ชัดเจน	7
รูปที่ 2.3 นิเวศเน็ตเวิร์กแบบแบ็กพรอพพาเกชันที่มีชั้นที่ถูกซ่อน 1 ชั้น	9
รูปที่ 3.1 การหาเมตริกซ์ไฮเกนและเวกเตอร์ค่าเฉลี่ยของข้อมูลชุดสอน.....	14
รูปที่ 3.2 การหาเมตริกซ์การแปลงของการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท	15
รูปที่ 3.3 การแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทและเรียนรู้ด้วยนิเวศเน็ตเวิร์ก	16
รูปที่ 3.4 การแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทและการรู้จำด้วยนิเวศเน็ตเวิร์ก...	17
รูปที่ 3.5 ภาพรวมการเรียนรู้และการรู้จำด้วยนิเวศเน็ตเวิร์ก.....	18

สารบัญตาราง

หน้า

ตารางที่ 3.1 จำนวนข้อมูลนำเข้าและจำนวนโหนดในชั้นซ่อน.....	18
ตารางที่ 3.2 ตัวอย่างค่าของไอเกนและผลรวมสะสมค่าของไอเกนของตัวอักษร ก.....	19
ตารางที่ 3.3 จำนวนข้อมูลนำเข้าและโหนดในชั้นซ่อนของเกณฑ์เปอร์เซ็นต์ของความแปรปรวน	20
ตารางที่ 3.4 จำนวนองค์ประกอบสำคัญของเกณฑ์การทดสอบกองเศษหินขอบภูเขา	20
ตารางที่ 3.5 จำนวนโหนดชั้นซ่อนและข้อมูลนำเข้าของเกณฑ์การทดสอบกองเศษหินขอบภูเขา	22
ตารางที่ 3.6 การจับกลุ่มด้วยการสังเกตตัวอักษรที่คล้ายกัน.....	23
ตารางที่ 3.7 จำนวนโหนดชั้นข้อมูลนำเข้าและจำนวนโหนดใช้ซ่อนจากเกณฑ์การเลือก องค์ประกอบสำคัญแบบต่างๆ.....	27
ตารางที่ 4.1. ผลการทดลองที่ได้จากการใช้วิธีการเลือกข้อมูลที่มีลักษณะสำคัญด้วยวิธีต่างๆ กับชุดทดสอบที่ 1	31
ตารางที่ 4.2 ผลการทดลองที่ได้จากการใช้วิธีการเลือกข้อมูลที่มีลักษณะสำคัญด้วยวิธีต่างๆ กับ ชุดทดสอบที่ 2	32
ตารางที่ 4.3 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิวรอลเน็ตเวิร์กแยกตามตัวอักษรบน ชุด ทดสอบที่ 1	32
ตารางที่ 4.4 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิวรอลเน็ตเวิร์กแยกตามตัวอักษรบน ชุด ทดสอบที่ 2	35
ตารางที่ 4.5 ค่าความถูกต้องของการรู้จำจากการวิเคราะห์องค์ประกอบสำคัญแบบดั้งเดิม เปรียบเทียบกับวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท.....	39
ตารางที่ 4. 6 ข้อดี-ข้อเสียของเกณฑ์การวิเคราะห์องค์ประกอบสำคัญแบบต่างๆ	39

บทที่ 1

บทนำ

1.1 ความเป็นมาและความสำคัญของปัญหา

การใช้งานระบบการรู้จำตัวอักษรจะช่วยย่นระยะเวลาการป้อนตัวอักษรเข้าเครื่องคอมพิวเตอร์และลดการใช้แรงงานมนุษย์ได้ดี ทำให้มีความสะดวกมากยิ่งขึ้น ทั้งนี้ระบบการรู้จำตัวอักษรจะต้องมีความถูกต้องในการรู้จำในอัตราที่สูงเพียงพอจนผู้ใช้ไม่ต้องมาตรวจแก้ผลจากการรู้จำที่ผิดพลาดจนทำให้เสียเวลาการทำงานซึ่งอาจใช้เวลามากกว่ามนุษย์พิมพ์งานเอง

เนื่องจากความต้องการระบบการรู้จำที่มีความถูกต้องในการรู้จำสูง จึงมีการวิจัยและพัฒนาอย่างต่อเนื่อง ในภาษาต่างประเทศมีระบบการรู้จำตัวอักษรที่สามารถใช้งานเชิงพาณิชย์หลายภาษาโดยเฉพาะภาษาอังกฤษมีการใช้งานระบบการรู้จำตัวอักษรอย่างแพร่หลาย ตัวอย่างเช่น งานจัดเก็บเอกสารในรูปแบบอิเล็กทรอนิกส์สำหรับห้องสมุด การแปลงหนังสือในรูปแบบกระดาษให้เป็นหนังสืออิเล็กทรอนิกส์ เป็นต้น

การรู้จำตัวอักษรภาษาไทยมีความยากกว่าการรู้จำตัวอักษรภาษาอังกฤษเนื่องจากตัวอักษรมีรูปร่างซับซ้อน ตัวอักษรมีรูปร่างคล้ายคลึงกันหลายตัว และมีจำนวนประเภท (class) มาก ในงานวิจัยนี้เราใช้ตัวอักษรภาษาไทยทั้งหมด 68 ประเภท (พยัญชนะ 44 ประเภท คือ ก, ข, ..., ฮ สระและวรรณยุกต์ 24 ประเภท คือ ะ, า, ... , *) ตัวอย่างของตัวอักษรมีรูปร่างซับซ้อน เช่น ฐ ประกอบด้วยเส้นตรง 5 เส้น ห่วงกลม 2 วง เส้นโค้งแบบอื่นอีก 3 เส้น เป็นต้น ส่วนตัวอักษรที่มีรูปร่างคล้ายกันก็อย่างเช่น ฏ และ ฎ เป็นต้น ซึ่งปัญหาต่างๆ ดังกล่าวทำให้มีความยากในการรู้จำตัวอักษร โดยที่ระบบการรู้จำผ่านมายังมีความถูกต้องในการรู้จำไม่เพียงพอ จึงต้องมีการวิจัยและพัฒนาต่อไปอีกมาก

ด้วยระบบการรู้จำด้วยการใช้นิวรอลเน็ตเวิร์กมีการพัฒนาเป็นอย่างมาก มีการประยุกต์ใช้ในงานในด้านต่างๆ มากมาย งานวิจัยนี้ใช้ระบบการเรียนรู้ด้วยนิวรอลเน็ตเวิร์กและใช้เทคนิคการวิเคราะห์ห้องค้ประกอบสำคัญแบบหลายประเภทซึ่งเป็นการพัฒนาต่อจากงานวิจัยการรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์ห้องค้ประกอบสำคัญแบบดั้งเดิม

อย่างไรก็ดีการวิเคราะห์ห้องค้ประกอบสำคัญดั้งเดิมนี้อย่างมีจุดด้อยในเรื่องของความสามารถของการแบ่งแยก (discriminative power) ทำให้การจำแนกประเภทข้อมูลที่มีลักษณะคล้ายกันทำ

ได้ไม่ดีนัก เช่น ตัวอักษร ฎ และ ฏ เป็นต้น เนื่องจากการวิเคราะห์องค์ประกอบสำคัญไม่ได้ถูกออกแบบเพื่อการจำแนกหรือแบ่งแยกข้อมูลแต่ได้ถูกออกแบบเพื่อการแทนข้อมูลพร้อมๆ กับการลดขนาดของข้อมูลโดยให้คงองค์ประกอบสำคัญไว้เท่านั้น

วิทยานิพนธ์ฉบับนี้มุ่งเน้นการพัฒนาการรู้จำตัวอักษรภาษาไทยโดยใช้วิธีการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท ซึ่งเป็นการปรับปรุงการวิเคราะห์องค์ประกอบสำคัญดั้งเดิมให้มีประสิทธิภาพดีขึ้น วิธีการนี้จะเป็นการสร้างเซตขององค์ประกอบสำคัญของข้อมูลแต่ละประเภท ทำให้เราได้เซตขององค์ประกอบสำคัญที่ต่างๆ กันของข้อมูลแต่ละประเภทจึงเป็นการเพิ่มความสามารถในการแบ่งแยกข้อมูล

1.2 วัตถุประสงค์

เพื่อพัฒนาระบบการรู้จำตัวอักษรภาษาไทยโดยใช้เทคนิคด้านการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทควบคู่กับระบบการรู้จำด้วยนิวรอลเน็ตเวิร์ก

1.3 ขอบเขตของการวิจัย

1. พัฒนาระบบการรู้จำอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท
2. วิเคราะห์และเปรียบเทียบประสิทธิภาพระบบการรู้จำอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทและระบบการรู้จำอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบดั้งเดิม

1.4 ขั้นตอนและวิธีดำเนินงานวิจัย

1. ศึกษาแนวคิดและทฤษฎีการรู้จำอักษรภาษาไทย นิวรอลเน็ตเวิร์ก และการวิเคราะห์องค์ประกอบสำคัญ
2. ศึกษางานวิจัยที่เกี่ยวข้อง
3. ออกแบบอัลกอริทึมและพัฒนาระบบการรู้จำอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท
4. ทดสอบประสิทธิภาพของระบบการรู้จำอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทกับข้อมูลทดสอบ

5. วิเคราะห์ผลเปรียบเทียบระบบการรู้จำอักขรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทกับระบบการรู้จำอักขรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบดั้งเดิมและปรับปรุงประสิทธิภาพ
6. สรุปผลการทดลอง และจัดทำวิทยานิพนธ์

1.5 ประโยชน์ที่คาดว่าจะได้รับ

ได้ระบบการรู้จำตัวอักขรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทและนิวรอลเน็ตเวิร์กซึ่งเป็นระบบการรู้จำแบบใหม่ที่เหมาะสมสำหรับการรู้จำข้อมูลที่มีจำนวนประเภทมากโดยเฉพาะตัวอักขรภาษาไทย

1.6 ลำดับการจัดเรียงเนื้อหาในวิทยานิพนธ์

วิทยานิพนธ์นี้แบ่งเนื้อหาออกเป็น 5 บทดังนี้ บทที่ 1 เป็นบทนำซึ่งกล่าวถึงที่มาและความสำคัญของปัญหา รวมทั้งวัตถุประสงค์ของงานวิจัยชิ้นนี้ บทที่ 2 กล่าวถึงทฤษฎีพื้นฐานและงานวิจัยที่เกี่ยวข้องในงานวิจัยนี้ บทที่ 3 กล่าวถึงรายละเอียดทั้งหมดของระบบการรู้จำอักขรภาษาไทยวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท ซึ่งเป็นวิธีที่นำเสนอในงานวิจัยนี้ บทที่ 4 แสดงรายละเอียดของการทดลองและผลการทดลอง และบทที่ 5 จะเป็นข้อสรุปและข้อเสนอแนะจากการวิจัย

1.7 ผลงานที่ตีพิมพ์จากวิทยานิพนธ์

ส่วนหนึ่งของวิทยานิพนธ์นี้ได้รับการตีพิมพ์เป็นบทความทางวิชาการและนำเสนอในงานประชุมวิชาการ ในหัวข้อ “การรู้จำตัวอักขรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทและนิวรอลเน็ตเวิร์ก” โดยอุดม สถาพรชัยสิทธิ์ และ บุญเสริม กิจศิริกุล ในงานประชุมวิชาการ "The 10th National Computer Science and Engineering Conference (NCSEC'2006)" ซึ่งจัดโดยมหาวิทยาลัยขอนแก่น ณ โรงแรมโซฟิเทลราชาขอนแก่น จังหวัดขอนแก่น ในวันที่ 25-27 ตุลาคม 2549

บทที่ 2

ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

บทนี้จะแบ่งเนื้อหาออกเป็น 2 ส่วนด้วยกัน โดยเนื้อหาในส่วนแรกจะกล่าวถึงทฤษฎีและแนวคิดพื้นฐาน มีเนื้อหาประกอบด้วย การวิเคราะห์องค์ประกอบสำคัญ เกณฑ์การเลือกองค์ประกอบสำคัญ และนิเวศน์เน็ตเวิร์ก เนื้อหาในส่วนที่สองจะกล่าวถึงงานวิจัยที่เกี่ยวข้องกับงานวิจัยนี้

2.1 ทฤษฎีที่เกี่ยวข้อง

2.1.1 การวิเคราะห์องค์ประกอบสำคัญ (Principle components analysis) [1, 2]

การวิเคราะห์องค์ประกอบสำคัญเป็นกระบวนการลดจำนวนมิติข้อมูลเพื่อประโยชน์ในการวิเคราะห์ข้อมูลได้ง่ายและมีประสิทธิภาพมากยิ่งขึ้น โดยใช้เทคนิคการแปลงมิติข้อมูลเชิงเส้น ไปสู่ระบบพิกัดใหม่ โดยที่ข้อมูลที่ฉายเงา (Projection) ลงบนแกนองค์ประกอบสำคัญแรก (First principle component axis) ที่มีค่าความแปรปรวน (variance) มากที่สุด แกนองค์ประกอบสำคัญที่สองและแกนถัดไป มีค่าความแปรปรวนน้อยลงตามลำดับ ซึ่งกระบวนการวิเคราะห์องค์ประกอบสำคัญมีลำดับขั้นตอนดังนี้

2.1.1.1 การหาเมตริกซ์ค่าเฉลี่ยและเมตริกซ์โคแวกเรียน

การหาเวกเตอร์ค่าเฉลี่ย (mean vector) – M_x และเมตริกซ์โคแวกเรียนซ์ (covariance matrix) – C_x เมื่อภาพของตัวอักษรที่ต้องการรู้จำมีความกว้าง a_1 จุดและความยาว a_2 จุด เราแทนแต่ละภาพด้วยเวกเตอร์ X_i มีขนาดเท่ากับ N ($N = a_1 \times a_2$) โดยที่ N คือจำนวนจุดภาพทั้งหมด (ในงานวิจัยนี้ เวกเตอร์ X_i มีขนาดเท่ากับ $32 \times 32 = 1,024$) เวกเตอร์ข้อมูล X_i จะถูกเรียงต่อกันเป็นแถวตั้งแต่ X_1 จนถึง X_T หรือ $X = \{X_1, X_2, \dots, X_T\}$ การหาเวกเตอร์ค่าเฉลี่ย M_x และเมตริกซ์โคแวกเรียนซ์ C_x ที่มีขนาด $N \times N$ คำนวณได้ตามสมการ (2.1) และ (2.2) ตามลำดับ

$$M_x = \frac{1}{T} \sum_{i=1}^T X_i \quad (2.1)$$

$$C_x = \frac{1}{T} \sum_{i=1}^T X_i X_i^T - M_x M_x^T \quad (2.2)$$

โดยที่ T คือ จำนวนเวกเตอร์ของภาพตัวอย่างทั้งหมดที่ใช้ในการเรียนรู้

2.1.1.2 การหาเมตริกซ์ไอเกนเวกเตอร์

จากเมตริกซ์โคแวนเรียนซ์ C_x ที่ได้ นำไปหาเมตริกซ์ไอเกนเวกเตอร์ (eigenvectors matrix) ขนาด $N \times N$ ซึ่งการคำนวณเมตริกซ์ไอเกนเวกเตอร์นั้น สามารถทำได้โดยการหาค่าไอเกน (eigen value - λ_i) และไอเกนเวกเตอร์ (eigen vector - e_i) ตามสมการ (2.3)

$$C_x e_i = \lambda_i e_i \quad \text{โดยที่ } i = 1, 2, \dots, N \quad (2.3)$$

ซึ่งเมตริกซ์ค่าของไอเกน คือ

$$e = \{e_1, e_2, \dots, e_k\} \quad \text{โดยที่ } N \text{ เป็นจำนวนของไอเกนเวกเตอร์} \quad (2.4)$$

2.1.1.3 การแปลงแบบเค-แอล (K-L Transform)

เนื่องจากเมตริกซ์โคแวนเรียนซ์ ขนาด 1024×1024 มีสมาชิกเป็นจำนวนจริงทั้งหมด และเป็นเมตริกซ์สมมาตร (symmetric) จะสามารถหาไอเกนเวกเตอร์ทั้ง 1024 ตัวได้เสมอ เมื่อนำไอเกนเวกเตอร์ ที่ได้จากเมตริกซ์โคแวนเรียนซ์ มาสร้างเป็นเมตริกซ์ A โดยที่แต่ละแถวของเมตริกซ์ A จะประกอบด้วยสมาชิกที่ได้จากไอเกนเวกเตอร์ โดยให้เรียงไอเกนเวกเตอร์ ซึ่งมีค่าของไอเกน (eigenvalue) มากที่สุดอยู่แถวแรก และเรียงไอเกนเวกเตอร์ที่มีค่ารองลงมาอยู่แถวถัดไป ตามลำดับ จะสามารถทำการแปลงแบบเค-แอลได้ตามสมการ (2.5)

$$y = A(x - m_x) \quad (2.5)$$

เนื่องจากการใช้เมตริกซ์ A ขนาด 1024×1024 เป็นข้อมูลขนาดใหญ่มาก จึงต้องใช้หน่วยความจำจำนวนมาก ทำให้สิ้นเปลืองทรัพยากรของระบบการรู้จำมากเกินไป ดังนั้นจึงทำการสร้างเมตริกซ์ของการแปลง A_k ซึ่งมีขนาด $K \times N$ และเวกเตอร์ Y ที่ได้จะมีขนาด K เนื่องจากค่าไอเกนมีคุณสมบัติที่ว่า ค่าของมันจะลดลงเมื่อลำดับมีค่าสูงขึ้น ดังนั้นเราจะใช้เพียงแค่ค่าไอเกนแรกๆ ที่มีค่าสูงและจะทิ้งตัวหลังๆ ที่มีค่าน้อยๆ ซึ่งถือว่าไม่มีความสำคัญมากนัก ด้วยวิธีการนี้ทำให้เราสามารถลดขนาดข้อมูลจาก $N \times N$ มาเป็น $K \times N$ ได้ โดยที่ $K \leq N$ โดยจำนวนของ K ในงานวิจัยนี้จะขึ้นกับเกณฑ์การเลือกองค์ประกอบสำคัญ โดยจะกล่าวในหัวข้อถัดไป

2.1.2 เกณฑ์การเลือกองค์ประกอบสำคัญ

เกณฑ์การเลือกจำนวนองค์ประกอบสำคัญของข้อมูลแต่ละประเภทจะนำไปใช้สำหรับเป็นข้อมูลนำเข้าเพื่อการเรียนของนิวรอลเน็ตเวิร์กต่อไป เราได้ประยุกต์ใช้วิธีของ Factor Analysis [3]

ซึ่งเป็นเทคนิคสำหรับคัดเลือกข้อมูลจากระบบการวิเคราะห์องค์ประกอบสำคัญดั้งเดิมและมีหลายวิธีด้วยกัน โดยการพิจารณาไอเกนเวกเตอร์และค่าของไอเกน [4] ซึ่งบางวิธีเป็นวิธีที่ใช้เฉพาะสำหรับปัญหาบางปัญหา ในวิทยานิพนธ์นี้จะใช้เกณฑ์การเลือกที่แตกต่างกัน 4 วิธี ซึ่งทำให้มีจำนวนข้อมูลสำหรับการเรียนรู้ที่แตกต่างกัน รายละเอียดของเกณฑ์การเลือกองค์ประกอบสำคัญมีดังนี้

1. เกณฑ์กำหนดล่วงหน้า [2] (A priori criterion)

เป็นกระบวนการที่ง่ายที่สุดในการเลือกองค์ประกอบสำคัญของแต่ละประเภทโดยเลือกองค์ประกอบที่มีค่าไอเกนจากค่ามากที่สุดที่เรียงกันไปหาน้อยเป็นองค์ประกอบสำคัญของประเภท n จำนวน P_n ตัวเท่ากันทุกๆ ประเภท โดย $P_n = k$ ตามที่ได้กำหนดไว้ล่วงหน้า ซึ่งจะทำให้ได้จำนวนข้อมูลนำเข้าทั้งหมดดังสมการ (2.6)

$$\sum_{n=1}^N P_n = ck \quad (2.6)$$

โดยที่ c คือ จำนวนประเภท

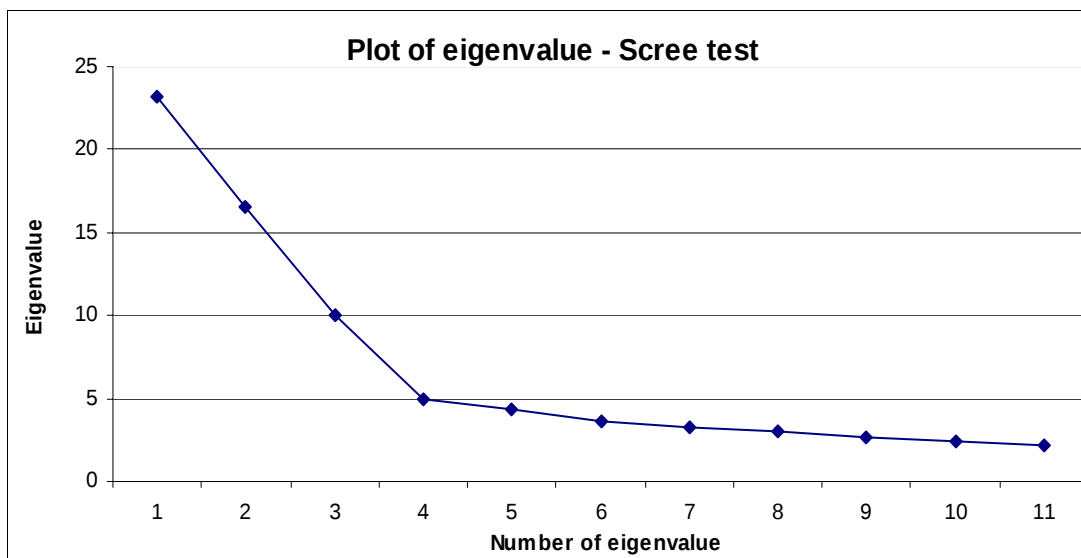
k คือ ค่าคงที่ของจำนวนองค์ประกอบที่ได้กำหนดไว้ล่วงหน้า

2. เกณฑ์เปอร์เซ็นต์ของความแปรปรวน [3] (Percentage of variance)

เกณฑ์เปอร์เซ็นต์ของความแปรปรวน ใช้การพิจารณาจากผลรวมสะสมค่าของค่าไอเกนตั้งแต่ตัวแรกที่มีค่ามากที่สุดเป็นต้นไปจนถึงตัวที่ผลรวมค่าไอเกนสะสมยังน้อยกว่าค่าขีดแบ่ง L สำหรับผลรวมสะสมของค่าไอเกนของแต่ละประเภท ซึ่งเป็นเลขจำนวนจริงมีค่าระหว่าง 0 ถึง 1 หรือ $L \in \mathcal{R}; L = \{0, 1\}$ โดยที่จะเลือกเฉพาะองค์ประกอบสำคัญตั้งแต่ตัวที่ 1 จนถึง n โดยที่องค์ประกอบสำคัญตัวที่ n มีผลรวมสะสมค่าของไอเกนน้อยกว่าหรือเท่ากับ L และองค์ประกอบตั้งแต่ตัวที่ $n+1$ จนถึงตัวสุดท้าย จะไม่ถูกเลือกใช้

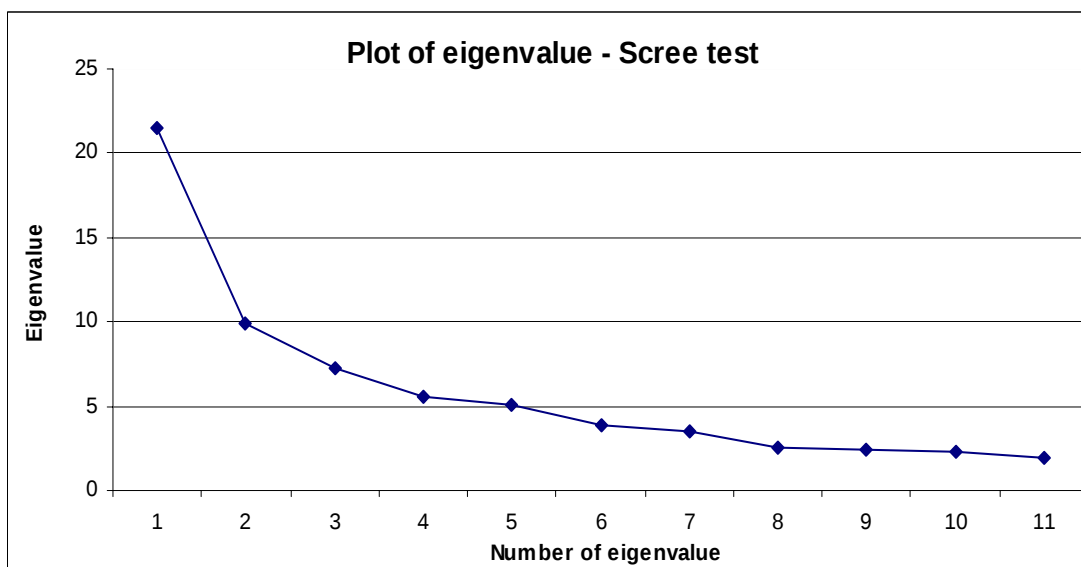
3. เกณฑ์การทดสอบกองเศษหินขอบภูเขา [5] (Scree test criterion)

เกณฑ์การทดสอบกองเศษหินขอบภูเขา เป็นกระบวนการหาองค์ประกอบสำคัญโดยอาศัยการพล็อตกราฟเชิงเส้นด้วยค่าไอเกนของไอเกนเวกเตอร์แต่ละตัวโดยเรียงลำดับจากมากไปน้อย จำนวนองค์ประกอบสำคัญของข้อมูลแต่ละประเภทจะถูกเลือกตั้งแต่ตัวแรกจนถึงตัวที่จุดเปลี่ยนกราฟเป็นความชันน้อยๆ หรือมีค่าใกล้เคียงศูนย์ดังตัวอย่างรูปที่ 2.1 เมื่อสังเกตกราฟจะพบว่าที่จุดที่ค่าไอเกนตัวที่ 4 ความชันเปลี่ยนไปมากดังนั้นเราจึงเลือกข้อมูลนำเข้าตั้งแต่ตัวที่ 1 ถึงตัวที่ 4 หรือ $P_n = 4$ เป็นต้น



รูปที่ 2.1 ตัวอย่างกราฟของการทดสอบแบบกองเศษหินที่ขอบภูเขา

แต่ในบางกรณีข้อมูลมีลักษณะที่เมื่อพล็อตกราฟค่าไอเกนแล้วมีจุดเปลี่ยนความชันที่ไม่ชัดเจนดังเช่นรูปที่ 2.2 เราอาจใช้การประมาณการโดยเลือกจำนวนข้อมูลให้ใกล้เคียงก็ได้ จากรูปที่ 2.2 เมื่อสังเกตในกราฟเราอาจประมาณเลือกค่าไอเกนตัวที่ 1 จนถึงตัวที่ 4 5 หรือ 6 ก็ได้ ดังนั้นเราอาจจะประมาณการเลือกข้อมูลนำเข้าตั้งแต่ตัวที่ 1 ถึงตัวที่ 5 เป็นต้น



รูปที่ 2.2 ตัวอย่างกราฟของการทดสอบแบบกองเศษหินที่ขอบภูเขาที่จุดเปลี่ยนความชันไม่ชัดเจน

4. เกณฑ์การทดสอบการแบ่งกลุ่ม

เกณฑ์การทดสอบการแบ่งกลุ่มนี้ใช้การพิจารณา และสังเกตตัวอักษรภาษาไทยต่างๆ ที่มีลักษณะคล้ายกัน โดยที่จะจัดกลุ่มที่ตัวอักษรคล้ายกันให้เป็นกลุ่มเดียวกัน จากข้อสมมติฐานว่าตัวอักษรหลายๆ ประเภทที่มีลักษณะคล้ายกันยังมีจำนวนตัวอักษรที่คล้ายกันมากขึ้นจะทำให้มีความยากในการแบ่งกลุ่มมากกว่าตัวอักษรที่มีเอกลักษณ์หรือไม่คล้ายกับตัวอักษรอื่น ดังนั้นกลุ่มที่มีจำนวนตัวอักษรคล้ายกันจำนวนมากควรมีจำนวนข้อมูลนำเข้ามากกว่า ในวิทยานิพนธ์นี้เรากำหนดให้

$$P_n = (V_n \times 2) + k \quad (2.7)$$

โดยที่ P_n คือ จำนวนองค์ประกอบสำคัญของข้อมูลประเภท n

V_n คือ จำนวนตัวอักษรที่คล้ายกันในกลุ่ม

k คือ จำนวนเต็มใดๆ

โดยเกณฑ์ทั้ง 4 วิธีที่กล่าวมาจะสามารถคำนวณจำนวนข้อมูลนำเข้าสู่นิวรอลเน็ตเวิร์กได้ดังสมการ

$$I = \sum_{n=1}^{68} P_n \quad (2.8)$$

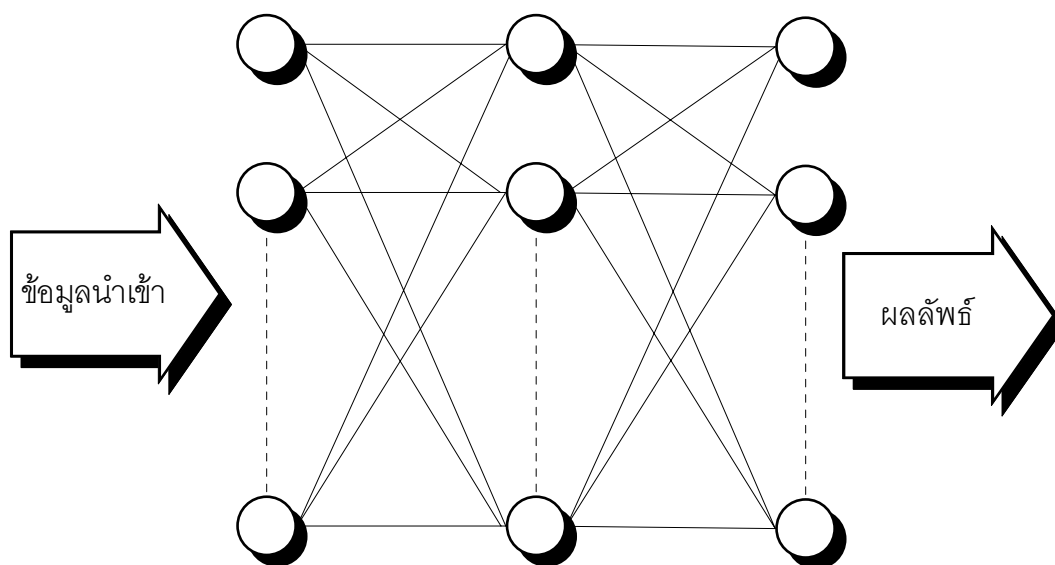
โดยที่ I คือ จำนวนข้อมูลนำเข้าสู่นิวรอลเน็ตเวิร์ก

P_n คือ จำนวนองค์ประกอบสำคัญของข้อมูลประเภท n

เนื่องจากจำนวนประเภทมี 68 ประเภทและจำนวนองค์ประกอบสำคัญที่มากที่สุดของข้อมูลแต่ละตัวอักษรเป็น 1024 ตัว จำนวนข้อมูลนำเข้า I ที่เป็นไปได้มีค่าน้อยที่สุดคือ 68 (68 คูณ 1) และมากที่สุดคือ 69632 (68 คูณ 1024)

2.1.3 นิวรอลเน็ตเวิร์ก [6]

เราใช้นิวรอลเน็ตเวิร์กแบบแบ็กพรอพagation ในการรู้จำรูปแบบ (Pattern Recognition) [7] นิวรอลเน็ตเวิร์กมีชั้นซ่อน (hidden layer) ซึ่งอยู่ระหว่างชั้นทางเข้า และชั้นทางออก 1 ชั้นขึ้นไป รูปที่ 2.3 แสดงนิวรอลเน็ตเวิร์กที่มีชั้นซ่อน 1 ชั้น



รูปที่ 2.3 นิวรอลเน็ตเวิร์กแบบแบ็กพรอพากะชันที่มีชั้นที่ถูกซ่อน 1 ชั้น

ขั้นตอนการสอนนิวรอลเน็ตเวิร์กแบบแบ็กพรอพากะชันมี 3 ขั้นตอน ได้แก่

1. การป้อนรูปแบบที่ต้องการสอนไปข้างหน้า (feed-forward) จากชั้นทางเข้าไปยังชั้นทางออก
2. การคำนวณและกระจายป้อนกลับ (back propagation) ของค่าความผิดพลาดของผลลัพธ์ (output pattern) กับรูปแบบผลลัพธ์ที่ควรจะเป็น (target pattern)
3. การปรับค่าน้ำหนัก (weights)

การป้อนเวกเตอร์รูปแบบที่ต้องการสอนไปข้างหน้า

กำหนดค่าน้ำหนักและค่าไบแอสเริ่มต้น โดยใช้การสุ่มตัวเลขจำนวนจริงที่มีค่าน้อยๆ เช่น ระหว่าง -0.05 ถึง 0.05

แต่ละโหนด ในแต่ละชั้นยกเว้นชั้นทางเข้า จะให้ค่าผลลัพธ์ $a_i^{(l)}$ จากค่า $u_i^{(l)}$ ที่ป้อนสู่โหนด i ในแต่ละชั้น l ตามสมการ (2.9)

$$a_i^{(l)} = f(u_i^{(l)}) \quad (2.9)$$

โดยที่ $f(x)$ เป็นฟังก์ชันการกระตุ้น (activation function) โดยใช้ฟังก์ชันซิกมอยด์ (sigmoid function) ซึ่งมีคำนวณได้ตามสมการ (2.10)

$$f(x) = \frac{1}{1 + e^{-\frac{x}{\sigma}}} \quad (2.10)$$

ข้อมูลนำเข้าของค่าที่ป้อนเข้าสู่โหนด i ในแต่ละชั้น l แสดงได้ตามสมการ

$$u_i^{(l)} = \sum_{j=1}^{N^{(l-1)}} w_{ij}^{(l)} a_j^{l-1} + \theta_i^{(l)} \quad (2.11)$$

โดยที่ $\theta_i^{(l)}$ เป็นค่าไบแอสของโหนด i ในชั้นที่ l

การคำนวณและกระจายป้อนกลับของค่าความผิดพลาดของผลลัพธ์หาได้ตามสมการ

$$\delta_i^{(l)}(m) = -\frac{\partial E}{\partial a_i^{(l)}(m)} \quad (2.12)$$

เมื่อ $\delta_i^{(l)}(m)$ เป็นค่าความผิดพลาดของผลลัพธ์โหนดที่ i ในชั้นที่ l เมื่อป้อนด้วยเวกเตอร์รูปแบบสำหรับสอนเวกเตอร์ที่ m และ E เป็นผลต่างกำลังสองที่น้อยที่สุด (least mean square)

$$E = \frac{1}{2} \sum_{m=1}^M \sum_{n=1}^N [t_i^{(l)}(m) - a_i^{(l)}(m)]^2 \quad (2.13)$$

โดยที่ M คือ จำนวนข้อมูลนำเข้าที่ใช้สอน

N คือ มิติของผลลัพธ์ (จำนวนสมาชิกของเวกเตอร์ผลลัพธ์)

$t_i^{(l)}(m)$ คือ ค่าเป้าหมายที่โหนดที่ i ในชั้นที่ l

$a_i^{(l)}(m)$ คือ ค่าที่ได้จริงขณะนั้นเมื่อป้อนด้วยเวกเตอร์รูปแบบสำหรับการสอน

การหาค่าความผิดพลาดในแต่ละชั้นของนิเวศน์เน็ตเวิร์กโดยแบ่งเป็นชั้นผลลัพธ์และชั้นที่ไม่ใช่ผลลัพธ์ ซึ่งคำนวณได้ดังนี้

$$\delta_i^{(0)}(m) = t_i^{(0)}(m) - a_i^{(0)}(m) \quad (2.14)$$

$$\delta_i^{(l)}(m) = \sum_{j=1}^{N^{(l+1)}} \delta_j^{(l+1)} f'(u_j^{l+1}(m)) w_{ji}^{l+1}(m) \quad (2.15)$$

การปรับค่าน้ำหนัก จะทำการปรับให้ค่าน้ำหนักและค่าไบแอสมีผลต่างกำลังสองของเวกเตอร์ผลลัพธ์ที่ได้กับเวกเตอร์ผลลัพธ์เป้าหมายมีค่าน้อยที่สุด หรือใกล้เคียง 0 ดังสมการ

$$w_{ij}^{(l)}(m+1) = w_{ij}^{(l)}(m) + \Delta w_{ij}^{(l)}(m) \quad (2.16)$$

โดย $w_{ij}^{(l)}$ คือ ค่าน้ำหนักจากโหนดที่ i ในชั้นที่ l ที่เชื่อมต่อกับโหนดที่ j

การหาค่า $\Delta w_{ij}^{(l)}$ สามารถโหนดได้ดังสมการ

$$\Delta w_{ij}^{(l)}(m) = \eta \delta_i^{(l)}(m) f'(u_j^{l+1}(m)) a_{ji}^{(l-1)}(m) \quad (2.17)$$

2.2 งานวิจัยที่เกี่ยวข้อง

งานวิจัยนี้มีความเกี่ยวข้องกับงานวิจัยด้านการรู้จำตัวอักษรพิมพ์ภาษาไทยโดยได้มีงานวิจัยซึ่งใช้การเลือกคุณลักษณะ (feature extraction) และการรู้จำรูปแบบ (pattern recognition) แตกต่างกันไป โดยงานวิจัยนี้เกี่ยวข้องกับงานวิจัยของธเนศ [8] มาก เนื่องจากเป็นการปรับปรุงการรู้จำตัวอักษรพิมพ์ภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญและนิรขอลเน็ตเวิร์กแบบดั้งเดิมให้มีประสิทธิภาพทางด้านการจำแนกข้อมูลที่มีประเภทมากยิ่งขึ้น งานวิจัยที่เกี่ยวข้องต่างๆมีดังนี้

ธเนศ [8] นำเสนอการรู้จำตัวอักษรพิมพ์ภาษาไทย โดยใช้เทคนิคด้านการวิเคราะห์องค์ประกอบสำคัญและนิรขอลเน็ตเวิร์ก เป็นงานวิจัยเกี่ยวกับการรู้จำตัวอักษรภาษาไทย โดยการหาคุณลักษณะ (feature) ด้วยเทคนิคการวิเคราะห์องค์ประกอบสำคัญ โดยมีจุดมุ่งหมายการลดจำนวนข้อมูลนำเข้าโดยคงคุณลักษณะที่สำคัญไว้มากที่สุด เพื่อให้ระบบการรู้จำด้วยนิรขอลเน็ตเวิร์กทำได้มีประสิทธิภาพ ไม่ใช่ทรัพยากรทั้งด้านหน่วยความจำ และการประมวลผลมากเกินไป อีกทั้งยังคงความถูกต้องในการรู้จำที่ดีอีกด้วย

สุรพันธ์ [9] นำเสนอวิธีการรู้จำตัวอักษรลายมือเขียนภาษาไทยโดยการนำหัวของตัวอักษรมาพิจารณาเพื่อทำการจำแนกประเภทของตัวอักษรออกเป็นประเภทย่อยๆ โดยให้ตัวอักษรที่มีหัวอยู่บริเวณเดียวกันอยู่ประเภทเดียวกัน จากนั้นก็จะพิจารณาเปรียบเทียบตัวอักษรที่อยู่ในประเภทของตนเองแตกต่างไป ทั้งนี้ขึ้นกับลักษณะเด่นของตัวอักษรที่อยู่ในประเภทนั้นๆ ซึ่งทำให้ประเภทของตัวอักษรที่ต้องเปรียบเทียบมีจำนวนลดลง ทำให้เปรียบเทียบได้อย่างรวดเร็ว

พิพัฒน์ และมนลดา [10] นำเสนอวิธีการรู้จำตัวอักษรไทยหลายรูปแบบโดยใช้วิธีไดนามิกโปรแกรมมิ่ง โดยการพิจารณาเส้นแสดงของของอักษรโดยนำรหัสทิศทางแบบลูกโซ่ของฟรีแมนกับความแตกต่างของทิศทางของเส้นแสดงของอักษรมาใช้ในการตัดแบ่งเส้นแสดงของของอักษรออกเป็นส่วนโค้งเว้าและส่วนโค้งนูน เพื่อนำไปใช้ในการเปรียบเทียบแบบไดนามิกโปรแกรมมิ่ง โดยการคำนวณค่าความคล้ายคลึง (Similarity) ของส่วนโค้งเว้าและส่วนโค้งนูนที่ได้กับส่วนโค้งเว้าและส่วนโค้งนูนของตัวอักษรต้นแบบ

สนธยา [11] นำเสนอเรื่องการศึกษาการรู้จำตัวอักษรพิมพ์ภาษาไทยโดยวิธีชินแทกติก ซึ่งเป็นการพิจารณาโครงสร้างของตัวอักษรโดยการเปลี่ยนแปลงเส้นแสดงของของตัวอักษรให้อยู่ในรูปรหัสทิศทางแบบลูกโซ่ของฟรีแมน และเปลี่ยนรหัสลูกโซ่เป็นรหัสเวกเตอร์เส้นตรงและวงกลมเพื่อนำมาจัดเก็บเป็นรูปของประโยคที่ประกอบด้วยพรีมิตีฟ (Primitive) ในลักษณะของโครงสร้างแบบต้นไม้ และอาศัยวิธีการวิเคราะห์ประโยคที่ได้จากการเปลี่ยนเส้นแสดงของของตัวอักษร

เปรียบเทียบกับประโยคของอักษรต้นแบบ โดยเลือกเปรียบเทียบเฉพาะตัวอักษรที่เป็นตัวอักษรอยู่ในระดับเดียวกันเท่านั้น (โดยการระบุเส้นบอกระดับ) สำหรับตัวอักษรที่แตกต่างกันไม่มากจะถูกนำไปเปรียบเทียบคุณลักษณะ (feature) อีกครั้งหนึ่งโดยการเก็บลักษณะพิเศษของแต่ละตัวอักษรไว้ หากผลการรู้จำในขั้นต้นไม่อยู่ในเกณฑ์ที่ยอมรับได้เวกเตอร์ของตัวอักษรจะถูกนำมา ปรับปรุงเพื่อตัดส่วนเกินออก หรือเชื่อมเวกเตอร์ที่อยู่ใกล้เคียงกันเข้าด้วยกันแล้วจึงนำมาทำการรู้จำ โดยวิธีเดิมอีกจนกว่าผลการรู้จำจะอยู่ในเกณฑ์ที่ยอมรับได้หรือไม่สามารถทำการปรับปรุงเวกเตอร์ได้อีก

เดชา [12] นำเสนอเรื่องการรู้จำตัวอักษรพิมพ์ภาษาไทยโดยใช้เทคนิคแบบฟัซซีโลจิก และวิธีซินแทกติก โดยทำการปรับปรุงวิธีการซินแทกติกของ สนธยา [11] โดยการนำเทคนิคแบบฟัซซีโลจิกเข้ามาใช้เมื่อการใช้วิธีการทางซินแทกติกไม่สามารถรู้จำตัวอักษร รวมทั้งปรับปรุงวิธีการทำตัวอักษรให้บาง โดยการใช้วิธีทำตัวอักษรให้บาง แบบเอสพีทีเอ (SPTA, Save Point Thinning Algorithm)

อภิญา [13] นำเสนอเรื่องการประยุกต์ใช้การโปรแกรมตรรกะเชิงอุปนัยในการรู้จำตัวอักษรพิมพ์ภาษาไทย โดยนำการเรียนรู้โดยการอุปนัยโดยใช้การโปรแกรมตรรกะเชิงอุปนัย (Inductive Logic Programming, ILP) โดยใช้ความรู้ภูมิหลัง (background knowledge) ในการสร้างสมมติฐานใหม่ที่สอดคล้องกับตัวอย่างที่ได้รับ ซึ่งเทคนิคขั้นต้นจะเป็นการพิจารณาโครงสร้างของตัวอักษรโดยทำการเปลี่ยนขอบของตัวอักษรเป็นรหัสทิศทางแบบลูกโซ่ของฟรีแมน ทำการเปลี่ยนรหัสทิศทางแบบลูกโซ่ของฟรีแมนเป็นเวกเตอร์เส้นตรง และเวกเตอร์วงกลม แล้วทำการเปลี่ยนเวกเตอร์เป็นหน่วยสร้างพื้นฐาน นำการโปรแกรมตรรกะเชิงอุปนัยมาทำการเรียนรู้ลักษณะของหน่วยสร้างพื้นฐานที่ได้จากตัวอักษรต้นแบบ เช่น ระดับของตัว ขนาดของตัวอักษร ลักษณะส่วนหัวของตัวอักษร ลักษณะส่วนปลายของตัวอักษร เป็นต้น

สุขวสา [14] นำเสนอเรื่องการรู้จำตัวอักษรพิมพ์ภาษาไทยโดยการใช้กลุ่มก๊อนนิวรอลเน็ตเวิร์ก ซึ่งใช้หลักการของกลุ่มก๊อนของตัวแยกแยะ (Ensemble of Classifiers) ในการรู้จำตัวอักษรไทย โดยใช้ตัวแยกแยะหลายตัว ตัวแยกแยะที่ดีจะต้องมีความถูกต้องสูงและมีความหลากหลายในการตอบคำตอบที่ผิด ถ้าความผิดพลาดจากตัวแยกแยะแต่ละตัวเป็นอิสระจากกัน จะทำให้ความผิดพลาดของกลุ่มก๊อนลดลงเป็นศูนย์ได้

บทที่ 3

การออกแบบและพัฒนาการรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทและนิรอลเน็ตเวิร์ก

บทนี้จะกล่าวถึงรายละเอียดการออกแบบและพัฒนาการรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทและนิรอลเน็ตเวิร์ก

3.1 โครงสร้างของระบบ

การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทมีวัตถุประสงค์เพื่อเพิ่มประสิทธิภาพในการแบ่งแยกข้อมูลที่อยู่คนละประเภทให้ออกจากกันได้ดียิ่งขึ้น ข้อมูลที่เป็นตัวอักษรเดียวกัน (เช่น ข้อมูลภาพ 'ก' หลายๆ ภาพ) จะถูกจัดอยู่ในประเภทเดียวกัน ในงานวิจัยนี้เราแบ่งประเภทข้อมูลทั้งหมด ออกเป็น 68 ประเภทด้วยกัน ทำการวิเคราะห์องค์ประกอบสำคัญกับทุกๆ ประเภท และนำข้อมูลมารวมกัน จากข้อมูลที่ได้ทำการเรียนรู้ และรู้จำด้วยนิรอลเน็ตเวิร์กต่อไป

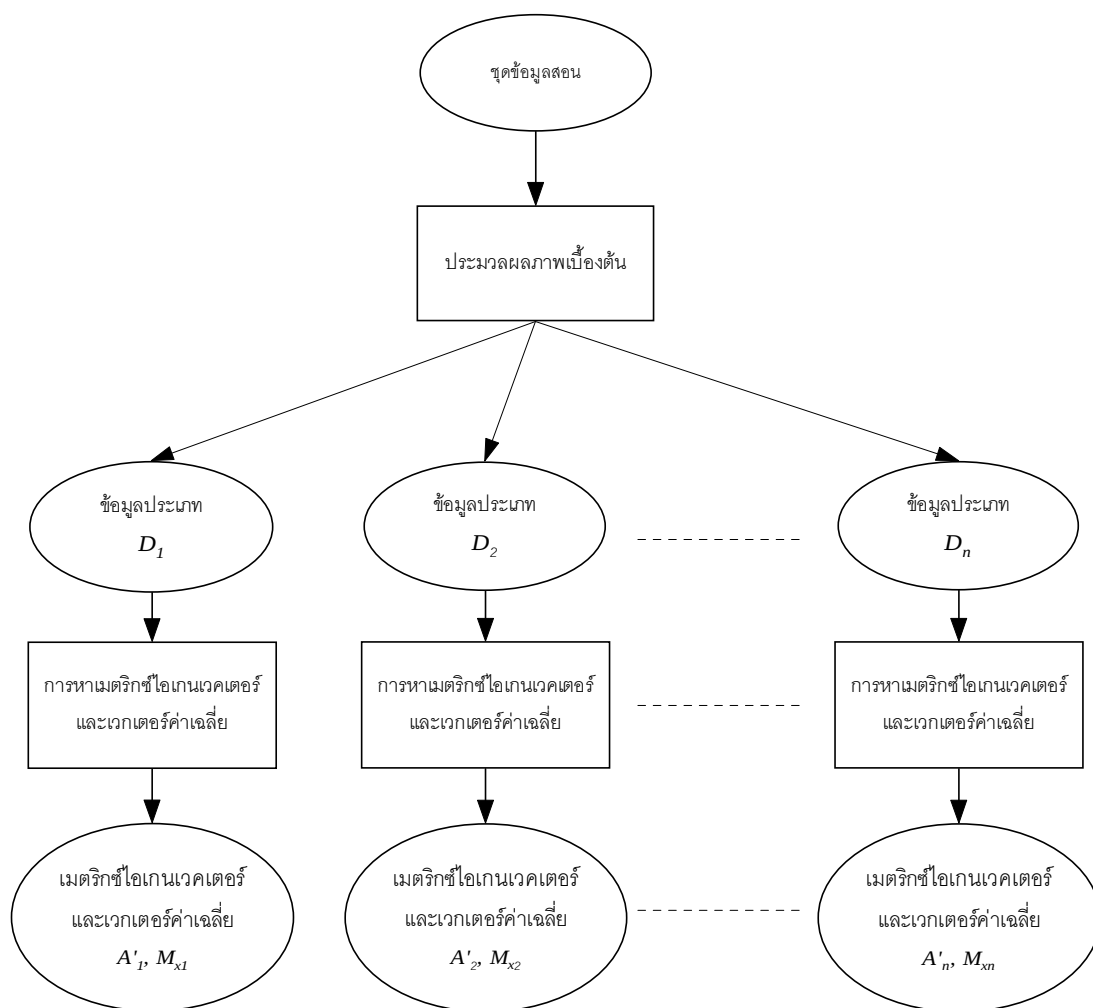
กระบวนการรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทและนิรอลเน็ตเวิร์กจะแบ่งเป็น 2 ขั้นตอน คือ ขั้นตอนการเรียนรู้ และขั้นตอนการรู้จำ มีรายละเอียดดังนี้

3.1.1 ขั้นตอนการเรียนรู้

ขั้นตอนการเรียนรู้แบ่งเป็น 2 ขั้นตอนย่อย คือ 1) การหาเมตริกซ์การแปลงและเวกเตอร์เฉลี่ย 2) ขั้นตอนการแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทและเรียนรู้ด้วยนิรอลเน็ตเวิร์ก มีรายละเอียดดังนี้

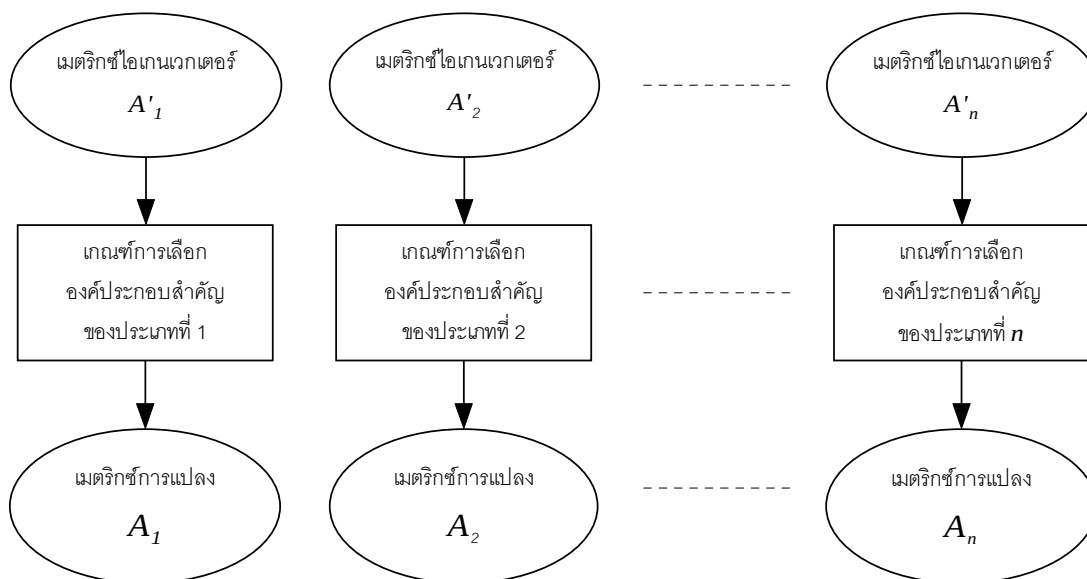
3.1.1.1 การหาเมตริกซ์การแปลงและเวกเตอร์เฉลี่ย

จากข้อมูลภาพชุดสอนเราทำการวิเคราะห์องค์ประกอบสำคัญด้วยวิธีการแปลงแบบเค-แอล โดยแยกทำกับชุดข้อมูลที่ละประเภทโดยใช้ข้อมูลชุดสอนที่ได้ทำการประมวลผลข้อมูลเบื้องต้นแล้ว ข้อมูลแต่ละประเภทจะรวมเป็นเมตริกซ์ข้อมูลของประเภทนั้นๆ จากนั้นทำการหาเมตริกซ์ไอเกนเวกเตอร์ (A'_i) เมตริกซ์ค่าไอเกน และเวกเตอร์ค่าเฉลี่ย (M_{xi}) 1 ชุดสำหรับข้อมูลแต่ละประเภท ซึ่งรวมทั้งสิ้นจำนวน 68 ชุด โดยข้อมูลดังกล่าวจะถูกนำไปใช้ทั้งในขั้นตอนการเรียนรู้และรู้จำ การหาเมตริกซ์การแปลง และเวกเตอร์ค่าเฉลี่ย มีขั้นตอนแสดงดังรูปที่ 3.1



รูปที่ 3.1 การหาเมตริกซ์ไอเกนและเวกเตอร์ค่าเฉลี่ยของข้อมูลชุดสอน

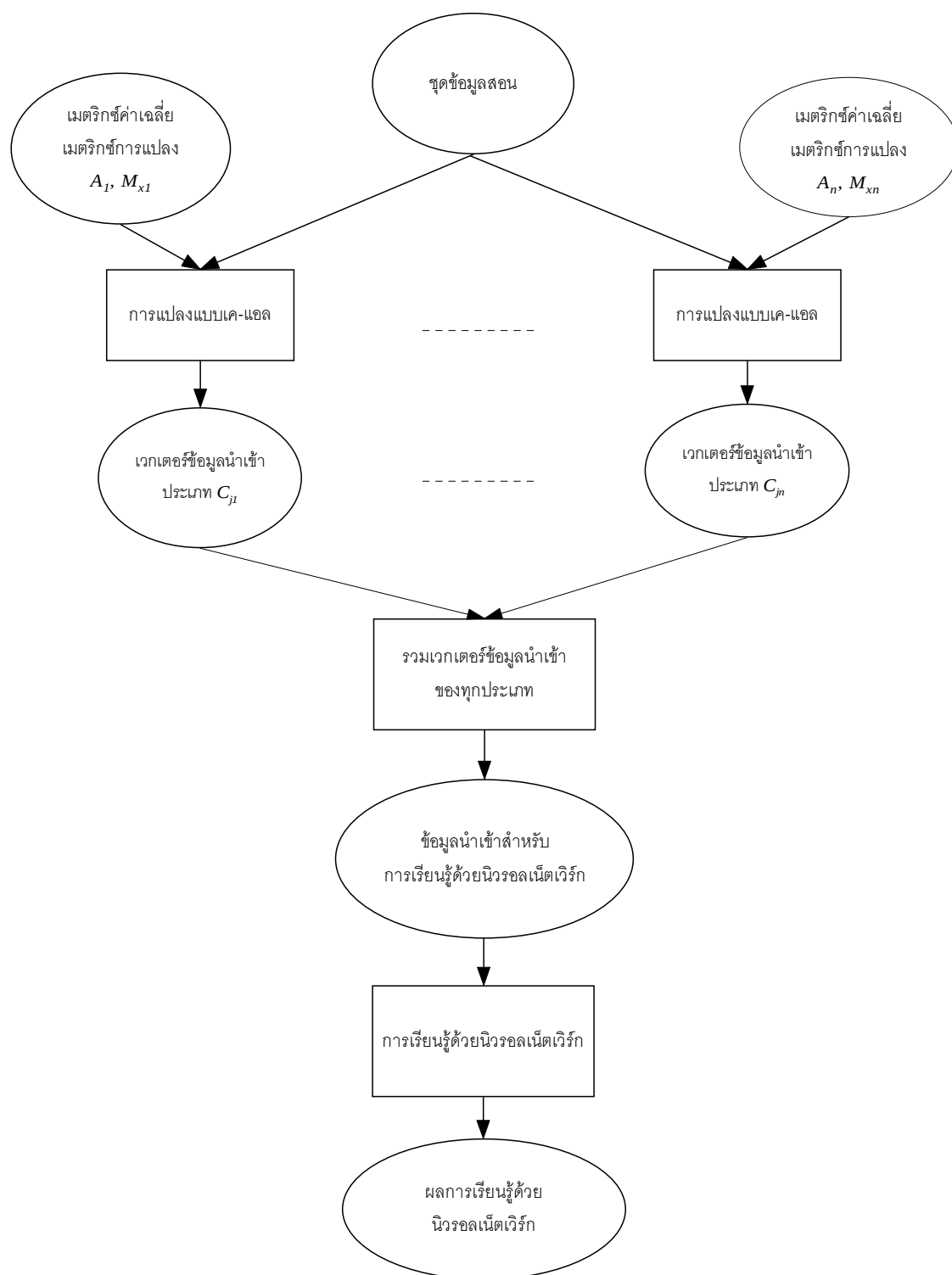
เมตริกซ์ไอเกนเวกเตอร์จะถูกเรียงลำดับตามค่าไอเกนของแต่ละไอเกนเวกเตอร์จากมากไปน้อย และเลือกไอเกนเวกเตอร์ตัวที่มีค่าไอเกนมากที่สุดได้เรียงมาจนได้จำนวนไอเกนเวกเตอร์ที่ต้องการซึ่งจะได้เมตริกซ์การแปลง A ขั้นตอนดังกล่าวสามารถแสดงขั้นตอนดังรูปที่ 3.2



รูปที่ 3.2 การหาเมตริกซ์การแปลงของการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท

3.1.1.2 การแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทและเรียนรู้ด้วยนิรอลเน็ตเวิร์ก

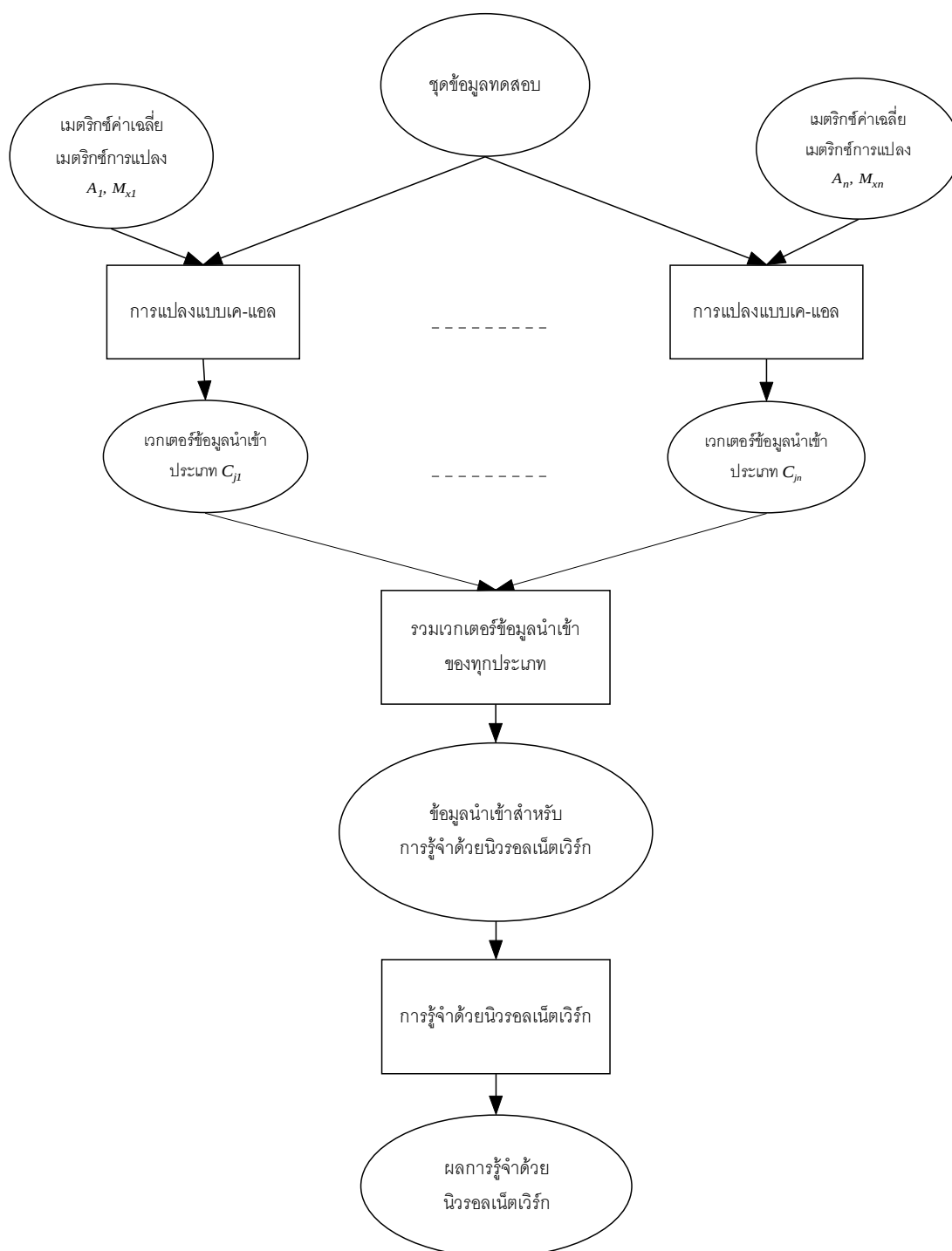
ทั้งเมตริกซ์การแปลงและเวกเตอร์ค่าเฉลี่ยจะถูกนำไปใช้สำหรับกระบวนการแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทโดยใช้การแปลงแบบเค-แอลในการแปลงข้อมูลชุดทดสอบตัวหนึ่ง กับเมตริกซ์การแปลงและเวกเตอร์เฉลี่ยของทุกประเภทซึ่งจะทำให้ได้กลุ่มข้อมูลนำเข้าของประเภท C_1 ถึง C_{68} (เนื่องจากมีจำนวนประเภททั้งหมด 68 ประเภท) ซึ่งข้อมูลทั้งหมดจะถูกเรียงลำดับตั้งแต่ประเภทที่ 1 ถึงประเภทที่ 68 ตามลำดับเป็นรวมกันเป็นข้อมูล 1 ชุดของตัวอักษรประเภทนั้นๆ เพื่อใช้เป็นข้อมูลนำเข้าในการเรียนรู้ด้วยนิรอลเน็ตเวิร์กต่อไป กระบวนการแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทแสดงได้ดังรูปที่ 3.3



รูปที่ 3.3 การแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทและเรียนรู้ด้วยนิวรอลเน็ตเวิร์ก

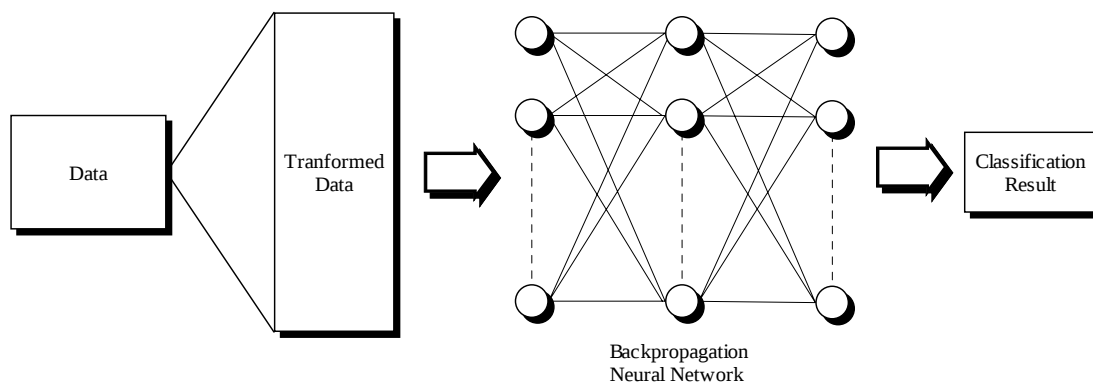
ขั้นตอนการรู้จำ

ในขั้นตอนการรู้จำ ทั้งเมตริกซ์การแปลงและเวกเตอร์เฉลี่ยที่ได้จากข้อมูลชุดสอนจะถูกนำไปใช้สำหรับการแปลงด้วยวิธีเค-แอลกับข้อมูลชุดทดสอบอีกทีหนึ่ง ซึ่งมีเกณฑ์ในการเลือกองค์ประกอบสำคัญแบบต่างๆ จะได้ข้อมูลนำเข้าสำหรับการรู้จำ ทำการรู้จำด้วยนิวรอลเน็ตเวิร์ก ผลลัพธ์ที่ได้คือประเภทของตัวอักษรที่ระบบทำการรู้จำได้ ขั้นตอนทั้งหมดแสดงดังรูปที่ 3.4



รูปที่ 3.4 การแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทและการรู้จำด้วยนิวรอลเน็ตเวิร์ก

ขั้นตอนโดยรวมของการเรียนรู้และรู้จำ สามารถแสดงด้วยรูปที่ 3.5 โดยข้อมูลตัวอักษรจะถูกแปลงด้วยการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท และทำการเรียนรู้หรือรู้จำด้วยนิวรอลเน็ตเวิร์ก สุดท้ายจะได้ผลลัพธ์ของการรู้จำ ดังแสดงในรูปที่ 3.5



รูปที่ 3.5 ภาพรวมการเรียนรู้และการรู้จำด้วยนิวรอลเน็ตเวิร์ก

3.2 เกณฑ์การเลือกองค์ประกอบสำคัญ

เกณฑ์การเลือกองค์ประกอบสำคัญมีความสำคัญอย่างยิ่ง ในกระบวนการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทนั้นเมื่อเราทำการแปลงข้อมูลเป็นข้อมูลนำเข้าของแต่ละประเภทจะได้จำนวนข้อมูล คือ ผลรวมของจำนวนกลุ่มข้อมูลนำเข้าของข้อมูลประเภท เมื่อจำนวนประเภทมี 68 ประเภท และจำนวนองค์ประกอบของข้อมูลดั้งเดิม 1024 ตัว จึงมีค่ามากที่สุดที่เป็นไปได้คือ 69,632 (68×1024) ซึ่งทำให้ทางปฏิบัติจำนวนข้อมูลดังกล่าวมากเกินไปที่จะใช้ในระบบการเรียนรู้ได้ ดังนั้นจึงต้องใช้เกณฑ์การเลือกองค์ประกอบสำคัญของแต่ละประเภทซึ่งเป็นวิธีการที่จะลดจำนวนองค์ประกอบของแต่ละประเภทลงโดยที่ต้องทำให้ประสิทธิภาพในการรู้จำตัวอักษรที่ดีทั้งในด้านการใช้ทรัพยากรคอมพิวเตอร์และความถูกต้องในการรู้จำด้วย เกณฑ์การเลือกองค์ประกอบสำคัญต่างๆ มีดังนี้

3.2.1 เกณฑ์กำหนดล่วงหน้า

งานวิจัยนี้ได้เลือกใช้ค่าจำนวนองค์ประกอบของแต่ละประเภทหรือ $P_n = 1$ ซึ่งจะได้ข้อมูลนำเข้า 68 ตัว และ $P_n = 5$ จะได้ข้อมูลนำเข้า 340 ตัว เฉพาะในกรณีเกณฑ์ $P_n = 1$ ถ้าให้จำนวนโหนดในชั้นซ่อนเป็นจำนวนครึ่งหนึ่งหรือ 34 โหนดจะมีจำนวนน้อยไปจึงกำหนดให้มีจำนวนโหนดชั้นซ่อนเท่ากับ 68 เฉพาะในกรณีนี้ จำนวนข้อมูลนำเข้าและจำนวนโหนดในชั้นซ่อนแสดงดังตารางที่ 3.1

ตารางที่ 3.1 จำนวนข้อมูลนำเข้าและจำนวนโหนดในชั้นซ่อน

จำนวนองค์ประกอบ ของแต่ละประเภท (P_n)	จำนวนโหนดในชั้น ซ่อน	จำนวนข้อมูลนำเข้า
1	68	68
5	170	340

3.2.2 เกณฑ์เปอร์เซ็นต์ของความแปรปรวน

เกณฑ์เปอร์เซ็นต์ของความแปรปรวนในงานวิจัยนี้ใช้ค่าขีดแบ่ง $L = 0.50$ และ 0.42 ตารางด้านล่างนี้จะแสดงตัวอย่างของการเลือกองค์ประกอบสำคัญโดยใช้ตารางค่าของไอเกน และผลรวมสะสมค่าของไอเกนช่วยในการเลือกองค์ประกอบสำคัญด้วยวิธีนี้ โดยกำหนดให้ค่าขีดแบ่ง $L = 0.50$ และแสดงเฉพาะตัวอักษร ก ซึ่งมีค่าไอเกนที่เรียงจากมากไปน้อยดังตารางที่ 3.2

ตารางที่ 3.2 ตัวอย่างค่าของไอเกนและผลรวมสะสมค่าของไอเกนของตัวอักษร ก

อันดับขององค์ประกอบสำคัญ	ค่าของไอเกน	ผลรวมสะสมค่าของไอเกน
1	0.2179	0.2179
2	0.1276	0.3456
3	0.0640	0.4096
4	0.0542	0.4638
5	0.0525	0.5163
6	0.0485	0.5648
7	0.0420	0.6068
8	0.0322	0.6390
9	0.0305	0.6695
10	0.0267	0.6962

จากตารางที่ 3.2 เมื่อพิจารณาจะพบว่าผลรวมสะสมค่าของไอเกนที่น้อยกว่า 0.50 จึงเลือกองค์ประกอบสำคัญตั้งแต่ตัวที่ 1 ถึงตัวที่ 4

การทดลองนี้ใช้ค่าขีดแบ่งเท่ากับ 0.50 และ 0.42 ตามลำดับ ซึ่งจะได้จำนวนข้อมูลนำเข้า คือ 335 ตัว และ 240 ตัว ดังแสดงในตารางต่อไปนี้

ตารางที่ 3.3 จำนวนข้อมูลนำเข้าและโหนดในชั้นซ่อนของเกนท์เปอร์เซ็นต์ของความแปรปรวน

ค่าขีดแบ่ง	จำนวนโหนดในชั้น ซ่อน	จำนวนข้อมูลนำเข้า
0.50	167	335
0.42	120	240

3.2.3 เกณฑ์การทดสอบกองเศษหินขอบภูเขา

ทำการพล็อตกราฟค่าไอเอนของข้อมูลชุดสอนกับทุกประเภท ด้วยการสังเกตจุดเปลี่ยน ความชันของกราฟเราจะได้จำนวนองค์ประกอบของข้อมูลแต่ละประเภทดังตารางที่ 3.4 ต่อไปนี้

ตารางที่ 3.4 จำนวนองค์ประกอบสำคัญของเกณฑ์การทดสอบกองเศษหินขอบภูเขา

ประเภท	ตัวอักษร	จำนวนองค์ประกอบสำคัญ
1	ก	3
2	ข	3
3	ฅ	3
4	ค	5
5	ค	4
6	ฅ	4
7	ง	4
8	จ	5
9	ฉ	5
10	ช	5
11	ฅ	3
12	ณ	5
13	ญ	2

ตารางที่ 3.4 จำนวนองค์ประกอบสำคัญของเกณฑ์การทดสอบกองเศษหินขอบภูเขา (ต่อ)

ประเภท	ตัวอักษร	จำนวนองค์ประกอบสำคัญ
14	ฎ	4
15	ฏ	4
16	ฐ	4
17	ฑ	4
18	ฒ	2
19	ณ	3
20	ด	3
21	ต	4
22	ถ	3
23	ท	3
24	ธ	4
25	น	2
26	บ	3
27	ป	3
28	ผ	3
29	ฝ	3
30	พ	3
31	ฟ	4
32	ภ	2
33	ม	3
34	ย	4
35	ร	4
36	ล	2
37	ว	5
38	ศ	3
39	ษ	3
40	ส	2
41	ห	2
42	ฬ	5
43	อ	3

ตารางที่ 3.4 จำนวนองค์ประกอบสำคัญของเกณฑ์การทดสอบกองเศษหินขอบภูเขา (ต่อ)

ประเภท	ตัวอักษร	จำนวนองค์ประกอบสำคัญ
44	ฮ	3
45	ะ	3
46	า	2
47	า	3
48	า	4
49	า	3
50	า	3
51	ง	2
52	ง	4
53	ง	4
54	ง	3
55	ง	5
56	เ	2
57	โ	3
58	ไ	3
59	ใ	3
60	อ	2
61	ฤ	2
62	ฎ	2
63	ๆ	3
64	ๆ	3
65	'	2
66	๖	2
67	๗	2
68	+	2

จากข้อมูลชุดสอนจะได้จำนวนข้อมูลนำเข้ารวมทั้งสิ้น 305 ตัว

ตารางที่ 3.5 จำนวนโหนดชั้นซ้อนและข้อมูลนำเข้าของเกณฑ์การทดสอบกองเศษหินขอบภูเขา

จำนวนโหนดในชั้นซ้อน	จำนวนข้อมูลนำเข้า
152	305

3.2.4 เกณฑ์การทดสอบการแบ่งกลุ่ม

เราได้สังเกตความคล้ายคลึงของตัวอักษร จำนวนองค์ประกอบสำคัญของแต่ละประเภท สามารถคำนวณได้จากสมการ (2.7) เมื่อรวมจำนวนองค์ประกอบสำคัญของแต่ละประเภทจะได้จำนวนข้อมูลนำเข้าทั้งหมด 262 ตัว ซึ่งสามารถแสดงการจับกลุ่มและจำนวนองค์ประกอบสำคัญ ดังตารางที่ 3.6

ตารางที่ 3.6 การจับกลุ่มด้วยการสังเกตตัวอักษรที่คล้ายกัน

กลุ่มที่	ตัวอักษร	จำนวนตัวอักษรที่คล้ายกันในกลุ่ม	จำนวนองค์ประกอบสำคัญ
1	ก ฅ	2	5
2	ข ฃ	2	5
3	ค ฅ	2	5
4	ฌ	1	3
5	ง	1	3
6	จ	1	3
7	ฉ	1	3
8	ช ฌ	2	5
9	ฌ ฌ	2	5
10	ญ	1	3
11	ฎ ฏ	2	5
12	ฐ	1	3
13	ฑ	1	3
14	ฒ	1	3
15	ด ต	2	5
16	ท	1	3
17	ธ	1	3

ตารางที่ 3.6 การจับกลุ่มด้วยการสังเกตตัวอักษรที่คล้ายกัน (ต่อ)

กลุ่มที่	ตัวอักษร	จำนวนตัวอักษรที่คล้ายกันในกลุ่ม	จำนวนองค์ประกอบสำคัญ
18	น	1	3
19	บ	1	3
20	ป	1	3
21	ผ	1	3
22	ฝ	1	3
23	พ	1	3
24	ฟ	1	3
25	ภ	1	3
26	ม	1	3
27	ย	1	3
28	ร	1	3
29	ล	1	3
30	ว	1	3
31	ศ	1	3
32	ษ	1	3
33	ส	1	3
34	ห	1	3
35	ฬ	1	3
36	อ	1	3
37	ฮ	1	3

ตารางที่ 3.6 การจับกลุ่มด้วยการสังเกตตัวอักษรที่คล้ายกัน (ต่อ)

กลุ่มที่	ตัวอักษร	จำนวนตัวอักษรที่คล้ายกันในกลุ่ม	จำนวนองค์ประกอบสำคัญ
38	ะ	1	3
39	า	1	3
40	ั	1	3
41	๒ ๓ ๔	3	7
42	ง	1	3
43	ฃ	1	3
44	ฅ	1	3
45	ง	1	3
46	ฉ	1	3
47	เ	1	3
48	โ ไ ใ	3	7
49	°	1	3
50	ฤ	1	3
51	ฦ	1	3
52	๗ ๘	2	5
53	'	1	3
54	๙	1	3
55	๖	1	3
56	+	1	3

3.3 ข้อกำหนดต่างๆ ของนิรอลเน็ตเวิร์ก

เมื่อได้ข้อมูลนำเข้าที่จะทำการเรียนรู้และรู้จำด้วยระบบนิรอลเน็ตเวิร์กแล้ว เรากำหนดให้นิรอลเน็ตเวิร์กเป็นแบบแบ็กพรอปาเกชันนิรอลเน็ตเวิร์ก ซึ่งมีจำนวนโหนดในชั้นซ่อนเป็นครึ่งหนึ่งของจำนวนข้อมูลนำเข้า ยกเว้นเกณฑ์กำหนดล่วงหน้าที่มีค่าที่กำหนดล่วงหน้ามีค่าเท่ากับหนึ่งเท่านั้นที่มีจำนวนโหนดในชั้นซ่อนเท่ากับจำนวนข้อมูลนำเข้า

1. ค่าเริ่มต้นของค่าน้ำหนักและค่าไบแอส

การกำหนดค่าเริ่มต้นของค่าน้ำหนักและค่าไบแอส ใช้การสุ่มเลขจำนวนจริงระหว่าง $(-0.5, 0.5)$

2. ฟังก์ชันการกระตุ้น

ฟังก์ชันการกระตุ้นใช้ฟังก์ชันซิกมอยด์ ซึ่งให้ค่าของฟังก์ชันอยู่ระหว่าง $(0, 1)$ ทุกโหนด

3. ค่าอัตราการเรียนรู้ (Learning Rate)

ค่าอัตราการเรียนรู้ เท่ากับ 0.05

4. จำนวนโหนดในชั้นทางเข้าและชั้นซ่อน

จำนวนโหนดขึ้นอยู่กับเกณฑ์การเลือกองค์ประกอบสำคัญแต่ละแบบ ซึ่งสามารถแสดงได้ดังตารางที่ 3.7

5. จำนวนโหนดในชั้นทางออก

เนื่องจากมีจำนวนประเภททั้งหมด 68 ประเภท จึงกำหนดให้จำนวนโหนดในชั้นทางออกมีจำนวน 68 โหนด

6. เงื่อนไขในการลู่เข้า (convergence)

กำหนดให้มีค่าความผิดพลาดไม่ต่ำกว่า 0.0001

7. จำนวนรอบในการเรียนรู้สูงสุด

กำหนดให้มีจำนวนรอบในการเรียนรู้สูงสุดคือ 200,000 รอบ

ตารางที่ 3.7 จำนวนโหนดชั้นข้อมูลนำเข้าและจำนวนโหนดใช้ซ้อนจากเกณฑ์การเลือก
องค์ประกอบสำคัญแบบต่างๆ

เกณฑ์การเลือกจำนวน ข้อมูลที่มีลักษณะสำคัญ	จำนวนโหนด ชั้นข้อมูล นำเข้า	จำนวนโหนด ในชั้นซ้อน
เกณฑ์กำหนดล่วงหน้า $P_n = 1$	68	68
เกณฑ์กำหนดล่วงหน้า $P_n = 5$	340	170
เกณฑ์เปอร์เซ็นต์ของความ แปรปรวน 0.50	335	167
เกณฑ์เปอร์เซ็นต์ของความ แปรปรวน 0.42	240	120
เกณฑ์ทดสอบกองเศษหิน ขอบภูเขา	305	152
เกณฑ์การทดสอบการ แบ่งกลุ่ม	262	131

บทที่ 4

การทดลอง

บทนี้กล่าวถึงการทดลองและผลการทดลองที่ได้จากการรู้จำตัวอักษรภาษาไทยโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทและนิเวศเน็ตเวิร์ก โดยเนื้อหาในบทนี้จะแบ่งออกเป็น รายละเอียดของข้อมูลที่ใช้ในการทดลอง วิธีการทดลอง ผลการทดลองและวิเคราะห์ผลการทดลอง

4.1 ชุดข้อมูลที่ใช้ในการทดลอง

ข้อมูลที่ใช้ทดสอบการรู้จำตัวพิมพ์อักษรไทยเป็นข้อมูลภาพที่เก็บไว้ในแฟ้มข้อมูล 1 แฟ้ม ต่อหนึ่งตัวอักษร การจัดเก็บข้อมูลภาพของตัวอักษรจัดเก็บโดยใช้วิธีการจัดเก็บแบบ Windows Bitmap (BMP) โดยตัวอักษรที่ใช้ในการวิจัยนี้ใช้รูปแบบของทรูไทยที่มีอยู่บนระบบปฏิบัติการ วินโดวส์ โดยการทดลองนี้ใช้ตัวอักษร 6 รูปแบบคือ AngsanaUPC, BrowaliaUPC, CordiaUPC, DilleniaUPC, EucrosiaUPC และ FreesiaUPC โดยรูปแบบแต่ละรูปแบบจะมีตัวอักษรทั้งหมด 68 ตัวแบ่งออกเป็น 2 กลุ่มใหญ่ๆ ได้ดังนี้

พยัญชนะ 44 ตัวอักษร ได้แก่

ก ข ข ค ค ฅ ง จ ฉ ช ซ ฌ ญ ฎ ฏ ฐ ฑ ฒ ณ ด ต ถ ท ธ น บ ป ผ ฝ พ ฟ ภ ม ย ร ล ว ศ
๒๓ ห พ ๒๔

สระและวรรณยุกต์รวม 24 ตัวอักษร ได้แก่

$\frac{d}{dt} \left(\frac{1}{r^2} \right) = -\frac{2}{r^3} \frac{dr}{dt}$

นอกจากนี้แต่ละรูปแบบของตัวอักษรจะประกอบด้วยตัวอักษรขนาด 16, 18, 20, 22, 24, 26, 28 และ 36 จุด ดังนั้นในหนึ่งชุดตัวอักษรจะมีทั้งหมด 3264 ตัวอักษร ตัวอย่างตัวอักษรที่ใช้สามารถดูได้ที่ภาคผนวก ก

ข้อมูลมีทั้งหมด 3 ชุดดังนี้

1. **ชุดสอน:** ชุดข้อมูลต้นแบบสำหรับใช้สอน ได้มาจากการพิมพ์ตัวอักษรทั้งหมด 3264 ตัวอักษรจากโปรแกรม Microsoft Word ด้วยเครื่องพิมพ์แบบเลเซอร์ที่มีความละเอียด 600 จุดต่อ

นิ้ว แล้วนำเอกสารที่ได้ไปทำการอ่านด้วยเครื่องสแกนเนอร์โดยใช้ความละเอียด 200 จุดต่อนิ้ว แล้วทำการตัดตัวอักษรเก็บลงแฟ้มข้อมูล 1 แฟ้มต่อหนึ่งตัวอักษร

2. ชุดทดสอบที่ 1: ข้อมูลสำหรับใช้ทดสอบชุดที่ 1 ได้มาจากการนำเอกสารที่ได้มาจากชุดสอนมาทำการถ่ายเอกสารด้วยเครื่องถ่ายเอกสารโดยปรับความเข้มของการถ่ายให้จางลง แล้วนำเอกสารที่ได้จากการถ่ายเอกสารไปทำการอ่านด้วยเครื่องสแกนเนอร์ที่ความละเอียด 200 จุดต่อนิ้ว แล้วทำการตัดตัวอักษรเก็บลงแฟ้มข้อมูล 1 แฟ้มต่อหนึ่งตัวอักษร

3. ชุดทดสอบที่ 2: ข้อมูลสำหรับใช้ทดสอบชุดที่ 2 ได้มาจากการนำเอกสารที่ได้มาจากชุดสอนมาทำการถ่ายเอกสารด้วยเครื่องถ่ายเอกสารโดยปรับความเข้มของการถ่ายให้เข้มขึ้น แล้วนำเอกสารที่ได้จากการถ่ายเอกสารไปทำการอ่านด้วยเครื่องสแกนเนอร์ที่ความละเอียด 200 จุดต่อนิ้ว แล้วทำการตัดตัวอักษรเก็บลงแฟ้มข้อมูล 1 แฟ้มต่อหนึ่งตัวอักษร

4.2 วิธีการทดลอง

การทดลองแบ่งออกเป็น 2 ขั้นตอนคือ

4.2.1 ขั้นตอนการเรียนรู้

1. ทำการประมวลผลภาพเบื้องต้นและแปลงข้อมูลรูปภาพตัวอักษรเก็บข้อมูลในรูปแบบไฟล์ตัวอักษร (text file) โดยจัดเก็บในรูปแบบ 0 แทนจุดขาว และ 1 แทนจุดดำ ในการทดลองนี้มีข้อมูลแต่ละประเภทจำนวน 48 ตัวอย่างด้วยกัน โดยจัดเก็บแยกตามประเภทเป็นไฟล์โดยชื่อ 01.txt ถึง 68.txt
2. ทำการวิเคราะห์องค์ประกอบสำคัญหาเมตริกซ์ค่าไอเกนและเวกเตอร์ค่าเฉลี่ย ซึ่งจะได้ไฟล์ข้อมูลที่เก็บค่าเมตริกซ์ไอเกนเวกเตอร์ไฟล์ชื่อ eval_01.txt ถึง eval_68.txt และไฟล์ที่เก็บเวกเตอร์ค่าเฉลี่ยมีชื่อ mean_01.txt ถึง mean_68.txt
3. ทำการวิเคราะห์องค์ประกอบสำคัญของข้อมูลชุดทดสอบของแต่ละประเภท แล้วจัดเก็บโดยแยกตามเกณฑ์ในแฟ้มที่ต่างกัน 6 ชนิด ดังนี้
 - เกณฑ์กำหนดล่วงหน้า โดย $P_n = 1$
 - เกณฑ์กำหนดล่วงหน้า โดย $P_n = 5$
 - เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.50
 - เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.42

- เกณฑ์การทดสอบกองเศษหินขอบภูเขา
 - เกณฑ์การทดสอบการแบ่งกลุ่ม
4. จากชุดข้อมูลสอนทำการแปลงแบบเค-แอลกับข้อมูลแต่ละประเภทด้วยเกณฑ์ต่างๆ กัน จะได้กลุ่มข้อมูลที่ได้ทำการแปลงแล้วของแต่ละประเภท เก็บเป็นไฟล์ชื่อ OCR_Train.txt ซึ่งในไฟล์ดังกล่าวเก็บข้อมูลนำเข้าชุดสอนและระบุประเภทที่ถูกต้องของข้อมูลชุดสอน
 5. ทำการสร้างนิเวศเน็ตเวิร์ก ซึ่งในงานวิจัยนี้ใช้ชุดคำสั่งสำหรับการพัฒนานิเวศเน็ตเวิร์ก ที่มีชื่อว่า Fast Artificial Neural Network Library (FANN) เวอร์ชัน 2.0 ดาวน์โหลดได้จากเว็บไซต์ <http://fann.sourceforge.net> ทำการคอมไพล์บนระบบปฏิบัติการวินโดวส์ เอ็กซ์พี ด้วย Microsoft Visual C++ เวอร์ชัน 6 และจะได้โปรแกรมชื่อ OCR_Trainer.exe
 6. ทำการเรียนรู้โดยโปรแกรม OCR_Trainer.exe โดยจะรับไฟล์ข้อมูลนำเข้า OCR_Train.txt โดยกำหนดเงื่อนไขในการเรียนรู้ต่างๆ ตามที่ได้ออกแบบในบทที่ 3 ให้ทำงานเรียนรู้อัตโนมัติไปเรื่อยๆ จนถึงเงื่อนไขในการรู้เข้า หรือจนถึงจำนวนรอบที่กำหนดไว้ โดยโปรแกรมได้กำหนดให้โปรแกรมทำการบันทึกค่าน้ำหนักและค่าไบแอสทุก 1000 รอบ จัดเก็บเป็นไฟล์ชื่อ OCR_<จำนวนรอบการเรียนรู้>.net เช่น OCR_10000.net โดยข้อมูลดังกล่าวจะใช้ในระบบการรู้จำต่อไป

4.2.2 ขั้นตอนการรู้จำ

1. ทำการประมวลผลภาพเบื้องต้นและแปลงข้อมูลรูปภาพตัวอักษรชุดทดสอบโดยจัดเก็บในรูปแบบ 0 แทนจุดขาว และ 1 แทนจุดดำ โดยจัดเก็บแยกตามประเภทเป็นไฟล์โดยชื่อ 01.txt ถึง 68.txt
2. ทำการแปลงแบบเค-แอลกับข้อมูลชุดทดสอบแต่ละประเภทด้วยเกณฑ์ต่างๆ กัน โดยใช้เมตริกซ์การแปลงที่ได้จากขั้นตอนการเรียนรู้ จะได้กลุ่มข้อมูลที่ได้ทำการแปลงแล้วของแต่ละประเภท เก็บเป็นไฟล์ชื่อ OCR_Test.txt ซึ่งในไฟล์ดังกล่าวเก็บข้อมูลนำเข้าชุดสอนและระบุประเภทที่ถูกต้องของข้อมูลชุดสอน
3. ทำการสร้างโปรแกรมสำหรับการรู้จำโดยชุดคำสั่งสำหรับการพัฒนานิเวศเน็ตเวิร์ก Fast Artificial Neural Network Library ด้วยคอมไพเลอร์ Microsoft Visual C++ จะได้โปรแกรมชื่อ TestOCR.exe
4. ทำการรู้จำกับข้อมูลชุดทดสอบโดยโปรแกรม TestOCR.exe ทำงานจะอ่านไฟล์ค่าน้ำหนักและค่าไบแอสจากไฟล์ชื่อ OCR_<จำนวนรอบการเรียนรู้>.net โดยเลือกใช้ไฟล์ค่า

น้ำหนักและค่าไบแอสที่มีค่าความถูกต้องจากการทดสอบด้วยข้อมูลชุดสอนมากที่สุด และรับข้อมูลนำเข้าที่ได้ทำการแปลงแบบเค-แอลด้วยเกณฑ์แบบต่างๆ แล้วบันทึกผลการรู้จำเป็นไฟล์ชื่อ result.txt โปรแกรมจะคำนวณเปอร์เซ็นต์ความถูกต้องจากการรู้จำข้อมูลชุดทดสอบ โดยเก็บข้อมูลผลลัพธ์การรู้จำเป็นเปอร์เซ็นต์ความถูกต้องแยกตามตัวอักษรแต่ละประเภท และเปอร์เซ็นต์ความถูกต้องโดยเฉลี่ยทั้งหมด

4.3 ผลการทดลอง

ผลความถูกต้องในการรู้จำแต่ละวิธีแสดงผลในตารางที่ 4.1 และตารางที่ 4.2 ส่วนจำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิรวลเน็ตเวิร์กแยกตามตัวอักษรแสดงโดยตารางที่ 4.3 และตารางที่ 4.4

ตารางที่ 4.1. ผลการทดลองที่ได้จากการใช้วิธีการเลือกข้อมูลที่มีลักษณะสำคัญด้วยวิธีต่างๆ กับชุดทดสอบที่ 1

เกณฑ์การเลือกจำนวนข้อมูลที่มีลักษณะสำคัญ	จำนวนองค์ประกอบสำคัญทั้งสิ้น (ตัว)	จำนวนตัวอย่างที่รู้จำผิด (เปอร์เซ็นต์ความถูกต้อง)
เกณฑ์กำหนดล่วงหน้า $P_n = 1$	68	86 (97.37%)
เกณฑ์กำหนดล่วงหน้า $P_n = 5$	340	81 (97.52%)
เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.50	335	78 (97.61%)
เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.42	240	81 (97.52%)
เกณฑ์ทดสอบกองเศษหิน ขอบภูเขา	305	73 (97.76%)
เกณฑ์การทดสอบการแบ่งกลุ่ม	262	70 (97.86%)

ตารางที่ 4.2 ผลการทดลองที่ได้จากการใช้วิธีการเลือกข้อมูลที่มีลักษณะสำคัญด้วยวิธีต่างๆ กับชุดทดสอบที่ 2

เกณฑ์การเลือกจำนวนข้อมูลที่มีลักษณะสำคัญ	จำนวนองค์ประกอบสำคัญทั้งสิ้น (ตัว)	จำนวนตัวอย่างที่รู้จำผิด (เปอร์เซ็นต์ความถูกต้อง)
เกณฑ์กำหนดล่วงหน้า $P_n = 1$	68	116 (96.45%)
เกณฑ์จำนวนของไอเกิน $P_n = 5$	340	107 (96.72%)
เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.50	335	108 (96.69%)
เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.42	240	105 (96.78%)
เกณฑ์ทดสอบกองเศษหิน ขอบภูเขา	305	113 (96.54%)
เกณฑ์การทดสอบการแบ่งกลุ่ม	262	104 (96.81%)

ตารางที่ 4.3 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิรวลเน็ตเวิร์กแยกตามตัวอักษรบนชุดทดสอบที่ 1

ตัวอักษร	จำนวนตัวอย่างทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความแปรปรวน 0.50	เปอร์เซ็นต์ความแปรปรวน 0.42	การทดสอบกองเศษหินขอบภูเขา	การทดสอบการแบ่งกลุ่ม
ก	48	0	0	0	3	0	1
ข	48	5	4	4	4	4	3
ฃ	48	4	5	3	2	2	3
ค	48	3	0	2	2	3	1
ฅ	48	3	4	2	3	4	3

ตารางที่ 4.3 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิรอลเน็ตเวิร์กแยกตามตัวอักษรบน
ชุดทดสอบที่ 1 (ต่อ)

ตัวอักษร	จำนวน ตัวอย่าง ทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความ แปรปรวน 0.50	เปอร์เซ็นต์ความ แปรปรวน 0.42	การทดสอบกองเศษ หินขอบภูเขา	การทดสอบ การแบ่งกลุ่ม
ฮ	48	1	0	0	0	0	1
ง	48	0	2	0	1	0	0
จ	48	1	3	1	0	1	1
ฉ	48	0	0	0	0	2	0
ช	48	1	0	0	0	1	1
ซ	48	3	3	3	3	3	3
ฌ	48	2	2	0	2	2	3
ญ	48	1	0	0	0	0	1
ฎ	48	4	3	4	3	3	2
ฏ	48	4	2	3	4	3	3
ฐ	48	0	0	0	0	0	0
ฑ	48	0	0	0	0	0	1
ฒ	48	0	0	2	0	0	0
ณ	48	0	0	0	0	0	0
ด	48	1	2	2	1	3	1
ต	48	2	2	3	2	1	3
ถ	48	0	1	0	0	0	0
ท	48	2	2	2	2	2	2
ธ	48	2	1	2	2	1	2
น	48	0	2	1	0	0	0
บ	48	2	2	2	2	2	2
ป	48	0	0	0	1	1	0
ผ	48	0	0	0	1	0	1
ฝ	48	0	1	1	2	0	1
พ	48	0	0	0	0	0	0
ฟ	48	3	2	1	1	2	1
ภ	48	0	0	0	0	0	0

ตารางที่ 4.3 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิรอลเน็ตเวิร์กแยกตามตัวอักษรบน
ชุดทดสอบที่ 1 (ต่อ)

ตัวอักษร	จำนวน ตัวอย่าง ทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความ แปรปรวน 0.50	เปอร์เซ็นต์ความ แปรปรวน 0.42	การทดสอบกองเศษ หินขอบภูเขา	การทดสอบ การแบ่งกลุ่ม
ม	48	0	0	0	0	0	0
ย	48	2	2	2	2	2	1
ร	48	0	0	0	1	0	0
ล	48	1	0	2	2	1	1
ว	48	1	2	2	1	2	1
ศ	48	0	0	0	1	1	1
ษ	48	0	1	0	0	0	0
ส	48	0	0	0	0	0	0
ห	48	1	0	0	0	0	0
พ	48	0	1	1	1	0	1
อ	48	1	1	0	0	1	0
ฮ	48	2	0	1	2	1	1
ะ	48	2	2	1	1	2	1
า	48	1	1	1	1	0	1
อ	48	2	2	1	2	2	3
า	48	1	0	1	1	0	2
า	48	2	2	2	2	2	2
า	48	2	2	2	2	2	2
ง	48	1	1	2	1	2	0
อ	48	2	2	1	2	2	2
อ	48	1	0	1	0	1	0
อ	48	2	2	2	3	1	1
อ	48	0	0	0	0	0	0
เ	48	2	2	2	1	1	0
โ	48	0	0	1	3	0	0
ไ	48	1	1	3	3	2	2

ตารางที่ 4.3 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิวรอลเน็ตเวิร์กแยกตามตัวอักษรบน
ชุดทดสอบที่ 1 (ต่อ)

ตัวอักษร	จำนวน ตัวอย่าง ทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความ แปรปรวน 0.50	เปอร์เซ็นต์ความ แปรปรวน 0.42	การทดสอบกองเศษ หินขอบภูเขา	การทดสอบ การแบ่งกลุ่ม
ใ	48	1	2	2	1	0	0
°	48	2	3	3	2	2	3
ฤ	48	2	1	2	0	0	0
ฎ	48	1	1	1	0	0	0
ฏ	48	2	2	1	1	2	1
ช	48	0	1	0	0	0	0
'	48	2	0	0	1	0	0
๒	48	0	1	0	0	1	1
๓	48	0	1	0	0	1	0
+	48	5	2	3	3	2	3
รวม	3264	86	81	78	81	73	70

ตารางที่ 4.4 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิวรอลเน็ตเวิร์กแยกตามตัวอักษรบน
ชุดทดสอบที่ 2

ตัวอักษร	จำนวน ตัวอย่าง ทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความ แปรปรวน 0.50	เปอร์เซ็นต์ความ แปรปรวน 0.42	การทดสอบกองเศษ หินขอบภูเขา	การทดสอบ การแบ่งกลุ่ม
ก	48	0	1	1	0	2	2
ข	48	4	4	3	4	4	3
ฃ	48	5	4	5	6	5	3
ค	48	5	2	3	4	3	1
ค	48	2	3	4	4	4	3
ฌ	48	0	4	0	3	2	2
ง	48	2	1	0	0	2	0
จ	48	3	2	2	2	3	2
ฉ	48	2	2	2	1	2	3

ตารางที่ 4.4 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิวรอลเน็ตเวิร์กแยกตามตัวอักษรบน
ชุดทดสอบที่ 2 (ต่อ)

ตัวอักษร	จำนวน ตัวอย่าง ทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความ แปรปรวน 0.50	เปอร์เซ็นต์ความ แปรปรวน 0.42	การทดสอบกองเศษ หินขอบภูเขา	การทดสอบ การแบ่งกลุ่ม
ช	48	4	3	2	2	2	4
ซ	48	2	2	3	1	3	1
ฌ	48	0	0	0	0	0	1
ญ	48	0	1	0	0	0	0
ฎ	48	3	2	3	1	3	2
ฏ	48	2	1	2	2	2	4
ฐ	48	0	2	1	1	0	0
ฑ	48	0	1	0	1	0	0
ฒ	48	0	1	3	0	2	0
ณ	48	0	1	0	1	2	2
ด	48	4	3	3	2	2	1
ต	48	4	3	3	3	3	3
ถ	48	3	2	2	2	2	3
ท	48	1	1	1	0	2	0
ธ	48	2	1	2	1	2	2
น	48	2	2	2	3	2	2
บ	48	0	0	0	1	0	0
ป	48	0	0	0	0	1	0
ผ	48	2	2	2	2	2	2
ฝ	48	0	0	1	1	1	0
พ	48	1	1	2	2	2	3
ฟ	48	0	1	0	2	2	0
ภ	48	0	1	2	1	0	0
ม	48	2	1	2	1	2	4
ย	48	2	3	2	2	2	3
ร	48	2	2	2	2	2	2

ตารางที่ 4.4 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิวรอลเน็ตเวิร์กแยกตามตัวอักษรบน
ชุดทดสอบที่ 2 (ต่อ)

ตัวอักษร	จำนวน ตัวอย่าง ทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความ แปรปรวน 0.50	เปอร์เซ็นต์ความ แปรปรวน 0.42	การทดสอบกองเศษ หินขอบภูเขา	การทดสอบ การแบ่งกลุ่ม
ด	48	0	1	2	1	1	0
ว	48	2	2	2	2	2	5
ศ	48	1	1	2	2	1	0
ษ	48	2	1	2	1	2	2
ส	48	1	2	2	1	2	1
ห	48	1	0	1	0	2	0
พ	48	0	0	0	1	0	1
อ	48	1	1	2	2	2	3
ฮ	48	1	1	1	2	1	1
ง	48	2	2	2	1	2	2
า	48	1	2	1	2	2	1
อ	48	0	0	0	1	1	0
า	48	0	0	0	1	0	0
า	48	1	1	1	1	2	1
า	48	0	0	0	0	0	0
า	48	3	2	1	1	1	1
า	48	3	2	2	1	2	2
า	48	4	2	3	3	2	2
า	48	4	2	2	1	2	2
า	48	3	2	3	2	2	3
า	48	4	3	3	2	1	5
า	48	0	1	0	0	0	4
า	48	0	0	0	0	0	0
า	48	3	2	3	2	2	2
า	48	3	1	0	1	1	0
า	48	2	1	0	1	0	0

ตารางที่ 4.4 จำนวนตัวอย่างที่รู้จำผิดพลาดด้วยนิวรอลเน็ตเวิร์กแยกตามตัวอักษรบน
ชุดทดสอบที่ 2 (ต่อ)

ตัวอักษร	จำนวน ตัวอย่าง ทั้งหมด	จำนวนตัวอย่างที่รู้จำผิดพลาด					
		กำหนดล่วงหน้า $P_n = 1$	กำหนดล่วงหน้า $P_n = 5$	เปอร์เซ็นต์ความ แปรปรวน 0.50	เปอร์เซ็นต์ความ แปรปรวน 0.42	การทดสอบกองเศษ หินขอบภูเขา	การทดสอบ การแบ่งกลุ่ม
ภ	48	3	2	2	3	1	2
ฏ	48	3	3	4	4	3	3
ช	48	2	2	2	1	2	0
'	48	2	2	1	2	0	2
ข	48	0	2	1	2	2	0
ฅ	48	1	1	1	1	1	0
+	48	3	3	2	2	3	1
รวม	3264	116	107	108	105	113	104

4.4 วิเคราะห์ผลการทดลอง

ดังแสดงในตารางที่ 4.4 ตัวอักษรที่คล้ายกัน (เช่น ข-ช, ฏ-ฐ และ ด-ต เป็นต้น) มักมีอัตราผิดพลาดที่สูงกว่าตัวอักษรอื่น ตัวอย่างเช่นตัวอักษร ข มีความถูกต้องในการรู้จำเฉลี่ยที่ต่ำเพียง 91.84% เนื่องจากตัวอักษรมีความคล้ายคลึงกันมากทำให้เป็นปัญหาในการจำแนกประเภท แต่เมื่อมีการเพิ่มจำนวนองค์ประกอบสำคัญโดยเฉพาะในประเภทที่มีลักษณะคล้ายกันเหล่านี้ ด้วยวิธีการทดสอบการแบ่งกลุ่มจะพบว่ามีความถูกต้องในการรู้จำที่สูงขึ้น ในการทดลองของตัวอักษร ข เกณฑ์การทดสอบการแบ่งกลุ่มมีเปอร์เซ็นต์ความถูกต้อง 93.75% ทั้งนี้เนื่องจากมีจำนวนองค์ประกอบสำคัญมากขึ้น ทำให้ความถูกต้องในการรู้จำประเภทตัวอักษรที่คล้ายกันมีความถูกต้องมากยิ่งขึ้น

ตารางที่ 4.5 ค่าความถูกต้องของการรู้จำจากการวิเคราะห์ห้องคำประกอบสำคัญแบบดั้งเดิม
เปรียบเทียบกับวิธีการวิเคราะห์ห้องคำประกอบสำคัญแบบหลายประเภท

เกณฑ์ที่ใช้	ชุดทดสอบที่ 1	ชุดทดสอบที่ 2	ค่าเฉลี่ย
การวิเคราะห์ห้องคำประกอบสำคัญแบบดั้งเดิม	97.15%	96.54%	96.86%
เกณฑ์กำหนดล่วงหน้า $P_n = 1$	97.35%	96.45%	96.80%
เกณฑ์กำหนดล่วงหน้า $P_n = 5$	97.51%	96.73%	97.12%
เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.50	97.59%	96.67%	97.13%
เกณฑ์เปอร์เซ็นต์ของความแปรปรวน 0.42	97.51%	96.79%	97.15%
เกณฑ์ทดสอบกองเศษหินขอบภูเขา	97.77%	96.53%	97.15%
เกณฑ์การทดสอบการแบ่งกลุ่ม	97.85%	96.82%	97.34%

ผลการทดลองของงานวิจัยนี้ดังตารางที่ 4.5 เมื่อเทียบกับวิธีการวิเคราะห์ห้องคำประกอบสำคัญโดยนิรอรเน็ตเวิร์กแบบดั้งเดิม พบว่าผลการทดลองด้วยวิธีต่างๆ โดยส่วนใหญ่แล้วให้ผลการรู้จำเฉลี่ยที่ดีขึ้นเมื่อเทียบกับวิธีดั้งเดิม ส่วนวิธีที่ให้ความถูกต้องในการรู้จำตัวอักษรภาษาไทยสูงที่สุดคือเกณฑ์การทดสอบการแบ่งกลุ่มที่ถึงแม้จำนวนข้อมูลนำเข้าน้อยกว่าวิธีเปอร์เซ็นต์ความแปรปรวน 0.50 (น้อยกว่า 30 ตัว) แต่ก็มีค่าความถูกต้องในการรู้จำที่สูงกว่า ซึ่งได้เปอร์เซ็นต์ความถูกต้องสำหรับชุดทดสอบที่ 1 คือ 97.85% เปอร์เซ็นต์ความถูกต้องสำหรับชุดทดสอบที่ 2 คือ 96.82% เปอร์เซ็นต์ความถูกต้องเฉลี่ยคือ 97.34%

จากการทดลองเราสามารถสรุปข้อดี-ข้อเสียของเกณฑ์การวิเคราะห์ห้องคำประกอบสำคัญแบบต่างๆ ได้ดังตารางที่ 4. 6

ตารางที่ 4. 6 ข้อดี-ข้อเสียของเกณฑ์การวิเคราะห์ห้องคำประกอบสำคัญแบบต่างๆ

เกณฑ์ในการเลือกองค์ประกอบสำคัญ	ข้อดี	ข้อเสีย
กำหนดล่วงหน้า	เป็นกระบวนการที่ง่ายที่สุด	อาจมีประสิทธิภาพน้อยกว่าวิธีอื่น
เปอร์เซ็นต์ของความแปรปรวน	มีเกณฑ์ในการที่เลือกที่ง่าย รู้จำตัวอักษรที่มีความซับซ้อนได้ดี	ตัวอักษรที่คล้ายกันอาจให้การรู้จำมีประสิทธิภาพที่ไม่ดี

ตารางที่ 4. 6 ข้อดี-ข้อเสียของเกณฑ์การวิเคราะห์องค์ประกอบสำคัญแบบต่างๆ (ต่อ)

เกณฑ์ในการเลือกองค์ประกอบสำคัญ	ข้อดี	ข้อเสีย
การทดสอบกองเศษหินขอบภูเขา	มีเกณฑ์ในการที่เลือกที่ง่าย รู้จำตัวอักษรที่มีความซับซ้อนได้ดี	เสียเวลาในการพล็อตกราฟและ สังเกตตัวอักษรทั้ง 68 ประเภท และตัวอักษรที่คล้ายกันอาจให้ ประสิทธิภาพในการรู้จำที่ไม่ดี
การทดสอบการแบ่งกลุ่ม	มีประสิทธิภาพมากที่สุด โดยที่ ตัวอักษรที่คล้ายกันให้ความ ถูกต้องในการรู้จำดีที่สุด	เสียเวลาในการเลือกตัวอักษร และถ้าเลือกไม่ดีอาจทำให้ ประสิทธิภาพในการรู้จำไม่ดี

บทที่ 5

สรุปผลการวิจัยและข้อเสนอแนะ

5.1 สรุปผลการวิจัย

งานวิจัยนี้ได้นำเสนอวิธีการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทซึ่งเป็นวิธีใหม่ในการรู้จำตัวอักษรภาษาไทย มีข้อดีคือสามารถใช้เกณฑ์การเลือกจำนวนองค์ประกอบสำคัญของข้อมูลแต่ละประเภทแบบต่างๆ ทำให้มีความยืดหยุ่นในปรับจำนวนข้อมูลที่มีลักษณะสำคัญที่เหมาะสมโดยส่วนใหญ่จำนวนข้อมูลองค์ประกอบสำคัญในแต่ละประเภทถ้ามีจำนวนมากกว่าจะให้ผลการรู้จำที่ถูกต้องมากกว่าแต่ทั้งนี้ยังขึ้นกับเกณฑ์การเลือกองค์ประกอบสำคัญด้วย

ขั้นตอนการทดลองแบ่งออกเป็น 2 ขั้นตอน คือ ขั้นตอนการสอนและขั้นตอนการทดสอบ ในขั้นตอนการสอนใช้ข้อมูลชุดสอนซึ่งประกอบด้วยตัวอักษรต้นแบบจำนวน 3264 ตัว ทำการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทด้วยเกณฑ์ต่างๆ โดยเราได้เสนอเกณฑ์ที่แตกต่างกัน 4 วิธีและทำการเรียนรู้ด้วยนิวรอลเน็ตเวิร์กแบบแบ็กพรอพาเกชัน ส่วนขั้นตอนการเรียนรู้ใช้ข้อมูลชุดทดสอบซึ่งประกอบด้วยตัวอักษร 2 ชุด ซึ่งได้จากการนำตัวอักษรต้นแบบไปถ่ายเอกสารให้จางลงและเข้มขึ้นแต่ละชุดมีจำนวน 3264 ตัว รวมทั้งหมด 6528 ตัว ทำการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทโดยใช้เวกเตอร์ค่าเฉลี่ยและเมตริกซ์ไอเกนเวกเตอร์ที่ได้จากข้อมูลชุดสอนด้วยเกณฑ์ต่างๆ แล้วทำการรู้จำด้วยนิวรอลเน็ตเวิร์กแบบแบ็กพรอพาเกชันที่ได้จากขั้นตอนการเรียนรู้ เพื่อวัดความถูกต้องในการรู้จำ

โครงสร้างของระบบการรู้จำตัวอักษรภาษาไทยโดยใช้วิธีการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทสามารถแบ่งออกเป็น 2 ส่วน คือ ส่วนการวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภท และส่วนการรู้จำตัวอักษรด้วยนิวรอลเน็ตเวิร์กแบบแบ็กพรอพาเกชัน ในส่วนของ การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทเราได้เสนอวิธีการในการเลือกองค์ประกอบสำคัญที่แตกต่างด้วยกัน 4 วิธี ซึ่งแต่ละวิธีมีข้อดีข้อเสียดังตารางที่ 4. 6 โดยเกณฑ์กำหนดล่วงหน้าเป็นวิธีการที่ง่ายที่สุดแต่มีประสิทธิภาพในการรู้จำน้อยที่สุดเนื่องจากไม่มีการพิจารณาความซับซ้อนหรือความคล้ายคลึงของตัวอักษร เกณฑ์เปอร์เซ็นต์ความแปรปรวนและเกณฑ์กองเศษหิน ขอบกฎเรามีการวิเคราะห์ความซับซ้อนของตัวอักษรทำให้มีประสิทธิภาพดีขึ้น เกณฑ์การทดสอบการแบ่งกลุ่มมีการวิเคราะห์ความคล้ายคลึงตัวอักษรทำให้มีประสิทธิภาพที่ดีเช่นกัน ในส่วนการรู้จำตัวอักษรด้วยนิวรอลเน็ตเวิร์กแบบแบ็กพรอพาเกชัน นิวรอลเน็ตเวิร์กมีจำนวนโหนดของชั้น

นำเข้าที่น้อยที่สุดคือ 68 โหนดและมากที่สุด 335 โหนด จำนวนโหนดชั้นชอนเป็นครึ่งหนึ่งของจำนวนโหนดชั้นนำเข้ายกเว้นเกณฑ์กำหนดล่วงหน้า $P_n = 1$ ที่มีจำนวนโหนดชั้นชอนเท่ากับจำนวนโหนดชั้นนำเข้า

ผลความถูกต้องของการรู้จำโดยใช้การวิเคราะห์องค์ประกอบสำคัญแบบหลายประเภทให้ผลการรู้จำเฉลี่ยที่ดีขึ้นเมื่อเทียบกับการวิเคราะห์องค์ประกอบสำคัญแบบดั้งเดิม ผลความถูกต้องของการรู้จำขึ้นอยู่กับการเลือกองค์ประกอบสำคัญ จำนวนโหนดในชั้นนำเข้า จำนวนโหนดในชั้นชอน และข้อกำหนดอื่นๆ ของนิรอลเน็ตเวิร์ก โดยเกณฑ์การทดสอบการแบ่งกลุ่มมีผลการรู้จำโดยเฉลี่ยดีที่สุดคือ 97.34%

5.2 ข้อเสนอแนะ

งานวิจัยนี้มีจุดประสงค์เพื่อเพิ่มประสิทธิภาพการรู้จำซึ่งมีข้อเสนอแนะดังต่อไปนี้

1. การใช้เกณฑ์การเลือกองค์ประกอบสำคัญแบบอื่นซึ่งต้องมีการค้นคว้าหรือวิจัยเพิ่มเติมเพื่อเพิ่มประสิทธิภาพในการรู้จำให้ดียิ่งขึ้น
2. การปรับเพิ่มหรือลดจำนวนองค์ประกอบในแต่ละเกณฑ์การเลือกองค์ประกอบสำคัญ ให้มากขึ้นหรือน้อยลงตามความเหมาะสม เช่น เมื่อต้องการความรวดเร็วมากขึ้นและใช้หน่วยความจำน้อยลง เราอาจลดจำนวนองค์ประกอบสำคัญลงได้ เป็นต้น
3. การปรับปรุงจำนวนโหนดในชั้นชอนของนิรอลเน็ตและจำนวนรอบการเรียนรู้ เพิ่มประสิทธิภาพในการรู้จำให้ดียิ่งขึ้น
4. การใช้จำนวนตัวอักษรชุดสอนจำนวนมากขึ้น และมีความหลากหลายที่มากขึ้นกว่าที่ใช้ในการวิจัยนี้ (การวิจัยนี้ใช้จำนวน 3264 ตัวอักษร) ในการเรียนรู้เพื่อให้การเรียนรู้ทำได้ดีขึ้น

รายการอ้างอิง

- [1] Mitchell, T. M. Machine Learning. New York: McGraw-Hill, 1997.
- [2] Anderson, T.W. An introduction to multivariate statistical analysis, 3rd edition. Hoboken, N.J. Wiley-Interscience, 2003.
- [3] Timm, N. H. Applied multivariate analysis. New York: Springer, 2002.
- [4] Freund, R. J. and Wilson, W. J. Statistical methods. Boston: Academic Press, 1993.
- [5] Therrien, C. W. Eigenvalue properties of projection operators and their application to the subspace method of feature extraction. Lexington, Mass.,: M.I.T. Lincoln Laboratory, 1974.
- [6] Cattell, R. B. Factor analysis : an introduction and manual for the psychologist and social scientist. New York: Harper, 1952.
- [7] Bishop, C. M. Neural networks for pattern recognition. Oxford New York: Clarendon Press ; Oxford University Press, 1995.
- [8] ธเนศ ศรีวิรุฬห์ชัย, การรู้จำตัวอักษรพิมพ์ภาษาไทย โดยใช้เทคนิคด้านการวิเคราะห์องค์ประกอบสำคัญและนิวรอลเน็ตเวิร์ก, วิทยานิพนธ์ปริญญาโทบริหารธุรกิจ สาขาวิชาวิทยาศาสตร์คอมพิวเตอร์ จุฬาลงกรณ์มหาวิทยาลัย ,2541.
- [9] สุรพันธ์ เชื้อไพบูลย์, การจดจำลายมือเขียนไทยโดยการพิจารณาหัวของตัวอักษร , วิทยานิพนธ์ปริญญาโทบริหารธุรกิจ สาขาเทคโนโลยีพระจอมเกล้าเจ้าคุณทหารลาดกระบัง ,2531.
- [10] พัทธมน หิรัญยวณิชกร และ มณฑา บุญสุวรรณ, การรู้จำตัวอักษรไทยหลายรูปแบบโดยวิธีไดนามิกโปรแกรมมิ่ง , สถาบันบัณฑิตพัฒนบริหารศาสตร์, บริษัท การบินไทย จำกัด.
- [11] สรณยา เมรินทร์, การศึกษาการรู้จำตัวอักษรพิมพ์ภาษาไทยโดยวิธีซินแทกติก, วิทยานิพนธ์ปริญญาโทบริหารธุรกิจ สาขาวิชาวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย ,2537.

- [12] เดชา รัตนธาร, การรู้จำตัวอักษรพิมพ์ภาษาไทยโดยใช้เทคนิคแบบฟัซซีโลจิก และวิธีซินแทกติก, วิทยานิพนธ์ปริญญามหาบัณฑิต สาขาวิศวกรรมไฟฟ้า จุฬาลงกรณ์มหาวิทยาลัย, 2538.
- [13] อภิญา สุพรรณวรราช, การประยุกต์ใช้การโปรแกรมตรรกะเชิงอุปนัยในการรู้จำตัวพิมพ์อักษรภาษาไทย, วิทยานิพนธ์ปริญญาวิทยาสตรมหาบัณฑิต สาขาวิชาวิทยาศาสตร์คอมพิวเตอร์ จุฬาลงกรณ์มหาวิทยาลัย ,2540.
- [14] สุขสา พิชิตเดช, การรู้จำตัวอักษรพิมพ์ภาษาไทยโดยการใช้กลุ่มก้อนของนิวรอลเน็ตเวิร์ก, วิทยานิพนธ์ปริญญามหาบัณฑิต สาขาวิชาวิทยาศาสตร์คอมพิวเตอร์ จุฬาลงกรณ์มหาวิทยาลัย ,2540.

ขนาด 36 จุด

ก ข ฃ ค ฅ ฆ ง จ ฉ ช ซ ฌ ญ ฎ ฏ
 ฐ ฑ ฒ ณ ด ต ถ ท ธ น บ ป ผ ฝ พ
 ฟ ภ ม ย ร ล ว ศ ษ ส ห ฬ อ ฮ ะ
 ั ิ ี ึ ุ ฤ ฦ ๅ ๆ ็ ่ ้ ๐
 ๑ ๒ ๓ ๔ ๕ ๖ ๗ ๘ ๙ +
 ฤ ฦ ๅ ๆ ็ ่ ้ ๐

ตัวอักษรต้นแบบรูปแบบ DilleniaUPCขนาด 16 จุด

ก ข ฃ ค ฅ ฆ ง จ ฉ ช ซ ฌ ญ ฎ ฏ ฐ ฑ ฒ ณ ด ต ถ
 ท ธ น บ ป ผ ฝ พ ฟ ภ ม ย ร ล ว ศ ษ ส ห ฬ อ ฮ
 ะ ั ิ ี ึ ุ ฤ ฦ ๅ ๆ ็ ่ ้ ๐ ฤ ฦ ๅ ๆ ็ ่ ๙ +

ขนาด 18 จุด

ก ข ฃ ค ฅ ฆ ง จ ฉ ช ซ ฌ ญ ฎ ฏ ฐ ฑ ฒ ณ ด ต ถ
 ท ธ น บ ป ผ ฝ พ ฟ ภ ม ย ร ล ว ศ ษ ส ห ฬ อ ฮ
 ะ ั ิ ี ึ ุ ฤ ฦ ๅ ๆ ็ ่ ้ ๐ ฤ ฦ ๅ ๆ ็ ่ ๙ +

ขนาด 28 จุด

ก ข ข ค ค ฅ ง จ ฉ ช ช ฌ ญ ฎ ฏ ฐ ฑ ฒ ณ ด ต ถ

ท ฒ บ ป ผ ฝ พ ฟ ภ ม ย ร ล ว ศ ษ ส ห ฬ อ ฮ

ะ ั ิ ึ ุ ฤ ฦ ๅ ๆ ็ ่ ้ ๐ ๑ ๒ ๓ ๔ ๕ ๖ ๗ ๘ ๙ +

ขนาด 36 จุด

ก ข ข ค ค ฅ ง จ ฉ ช ช ฌ ญ ฎ ฏ ฐ ฑ ฒ ณ ด ต ถ

ท ฒ บ ป ผ ฝ พ ฟ ภ ม ย ร ล ว ศ ษ ส ห ฬ อ ฮ

ะ ั ิ ึ ุ ฤ ฦ ๅ ๆ ็ ่ ้ ๐ ๑ ๒ ๓ ๔ ๕ ๖ ๗ ๘ ๙ +

๐ ๑ ๒ ๓ ๔ ๕ ๖ ๗ ๘ ๙ +

๐ ๑ ๒ ๓ ๔ ๕ ๖ ๗ ๘ ๙ +

ประวัติผู้เขียนวิทยานิพนธ์

นายอุดม สถาพรชัยสิทธิ์

เกิดวันที่ 17 กรกฎาคม พ.ศ. 2522

จังหวัดกรุงเทพมหานคร

สำเร็จการศึกษาปริญญาวิศวกรรมศาสตรบัณฑิต สาขาวิชาวิศวกรรมคอมพิวเตอร์
ภาควิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย ปีการศึกษา
2544

ศึกษาต่อในหลักสูตรวิศวกรรมศาสตรมหาบัณฑิต สาขาวิชาวิศวกรรมคอมพิวเตอร์
ภาควิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย ปีการศึกษา
2547