

การเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือ

จักรพันธ์ วาศพุฒิสิต

TNI

สารนิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตรปริญญาวิทยาศาสตรมหาบัณฑิต

บัณฑิตวิทยาลัย สาขาเทคโนโลยีสารสนเทศ

สถาบันเทคโนโลยีไทย-ญี่ปุ่น

ปีการศึกษา 2561

OPTIMIZATION OF INDEXING ON DOCUMENT BOXES WITH
HANDWRITING RECOGNITION TECHNIQUES

Jakkaphan Whasphuttisit

TNI

A Term Paper Submitted in Partial Fulfillment of the Requirements
for the Degree of Master of Science Program in Information Technology

Graduate School
Thai-Nichi Institute of Technology

Academic Year 2018

หัวข้อสารนิพนธ์

การเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วย

เทคนิคการรู้จำอักขระภาพลายมือ

โดย

จักรพันธ์ วาศุพลิสิต

สาขาวิชา

เทคโนโลยีสารสนเทศ

อาจารย์ที่ปรึกษาสารนิพนธ์

ดร. สะพรั่งสิทธิ์ มฤตสาร

บัณฑิตวิทยาลัย สถาบันเทคโนโลยีไทย-ญี่ปุ่น อนุมัติให้รับสารนิพนธ์ฉบับนี้เป็น
ส่วนหนึ่งของการศึกษาตามหลักสูตรปริญญามหาบัณฑิต

.....คณบดีบัณฑิตวิทยาลัย

(รองศาสตราจารย์ ดร. พิชิต สุขเจริญพงษ์)

วันที่.....เดือน.....พ.ศ.....

คณะกรรมการสอบสารนิพนธ์

.....ประธานกรรมการ

(ผู้ช่วยศาสตราจารย์ ดร. ฐิติพร เลิศรัตน์เดชากุล)

.....กรรมการ

(ดร. ประมุข บุญเสียง)

.....อาจารย์ที่ปรึกษาสารนิพนธ์

(ดร. สะพรั่งสิทธิ์ มฤตสาร)

จักรพันธ์ วาศพุฒิสิริ : การเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือ. อาจารย์ที่ปรึกษา : ดร. สะพรั่งสิทธิ์ มฤตสาร, 105 หน้า.

สารนิพนธ์ฉบับนี้มีวัตถุประสงค์ดังนี้ 1) เพื่อจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้ 2) เพื่อสามารถแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้มีความถูกต้อง (Accuracy) อย่างน้อย 80% 3) เพื่อลดระยะเวลาที่ใช้ในขั้นตอนการบันทึกข้อมูลดัชนี (Index) เป็น 30% ของกระบวนการเดิม โดยการออกแบบระบบมีขั้นตอนในการทำงาน 8 ขั้นตอนดังนี้ 1) ศึกษาข้อมูลเกี่ยวกับหลักการบริหารจัดการเก็บเอกสารในรูปแบบกล่อง 2) ศึกษางานวิจัยที่เกี่ยวข้องกับหลักการ OCR (Optical Character Recognition) 3) ศึกษาวิธีการพัฒนาตัวต้นแบบโปรแกรม 4) วางแผนและออกแบบตัวต้นแบบโปรแกรม 5) จัดทำตัวต้นแบบสำหรับใช้พัฒนาโปรแกรมในส่วนของการประมวลผลภาพ การทำคลังข้อมูลภาพตัวอักษร การรู้จำและการแปลความหมาย 6) ทดสอบการใช้งานของตัวต้นแบบที่จะใช้พัฒนาโปรแกรม 7) วิเคราะห์ประเมินและสรุปผล 8) จัดทำรายงานและนำเสนอผลงาน

ผลการสังเคราะห์รูปแบบการจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้ประกอบด้วย 2 ส่วน 1) ส่วน Training Data 2) ส่วน Testing Data จากผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารมีความถูกต้อง (Accuracy) 87.30%

คำสำคัญ : Optical Character Recognition, การประมวลผลภาพ, การทำคลังข้อมูลภาพตัวอักษร, การรู้จำ, การแปลความหมาย

TNI

JAKKAPHAN WHASPHUTTISIT : OPTIMIZATION OF INDEXING ON DOCUMENT BOXES WITH HANDWRITING RECOGNITION TECHNIQUES. ADVISOR : DR. SAPRANGSIT MRUETUSATORN, 105 PP.

The objectives of this research are as follows, 1) To create a prototype for developing programs that can reduce the process of recording handwritten information on a document box. 2) To be able to convert handwritten information on the document box to be at least 80% of accuracy. 3) To reduce the time spent in the index data recording process to 30% of the original process. This framework is divided into 8 steps. The first step learns about the principles of document storage management in a box. The second step study research related to Optical Character Recognition principles. The third step study how to develop a prototype program. The fourth step is planning and prototype design. The fifth step creates prototypes of the image processing section, Text image data library, Recognition and interpretation for developing programs. The sixth step tests the prototype. The seventh step is analysis, evaluation and summarize the result of this research. The last step prepares the report and presentation.

This framework is divided into 2 sections, 1) Training data section, 2) Testing data section. The result of the assessment of the handwritten data on the document box is 87.30% of accuracy.

TNI

Graduate School

Field of Study Information Technology

Academic Year 2018

Student's Signature.....

Advisor's Signature.....

กิตติกรรมประกาศ

การทำสารนิพนธ์นี้สำเร็จลงได้นั้น ผู้วิจัยขอขอบพระคุณในความกรุณาของ อาจารย์ ดร. สะพรั่งสิทธิ์ มฤตสาธร ที่ได้เสียสละเวลาในการให้คำปรึกษา คำแนะนำ ตลอดระยะเวลาในการทำสารนิพนธ์ฉบับนี้

ทั้งนี้ผู้วิจัยขอขอบพระคุณคณะกรรมการสอบทุกท่าน ได้แก่ รองศาสตราจารย์ ดร. พิชิต สุขเจริญพงษ์ ผู้ช่วยศาสตราจารย์ ดร. ฐิติพร เลิศรัตน์เดชากุล และ ดร. ประมุข บุญเสียง ที่กรุณาตรวจสอบแก้ไขข้อบกพร่องของสารนิพนธ์ฉบับนี้ด้วยความเอาใจใส่ จนทำให้สารนิพนธ์ฉบับนี้สำเร็จลงได้

ขอขอบคุณทุกท่านในคณะเทคโนโลยีสารสนเทศ สถาบันเทคโนโลยีไทย-ญี่ปุ่น ที่ให้การสนับสนุนและกำลังใจในความช่วยเหลือที่ผ่านมา

สุดท้ายนี้ผู้วิจัยหวังเป็นอย่างยิ่งว่า สารนิพนธ์ฉบับนี้จะเป็นประโยชน์ต่อในการใช้งาน และ การศึกษารวมทั้งเป็นต้นแบบในการพัฒนา เพื่อประยุกต์ใช้งานต่อไปในอนาคต

จักรพันธ์ วาศพุฒิสสิทธิ์

TNI

THAI - JAPAN
INSTITUTE OF TECHNOLOGY

สารบัญ

	หน้า
บทคัดย่อภาษาไทย.....	ง
บทคัดย่อภาษาอังกฤษ.....	จ
กิตติกรรมประกาศ.....	ฉ
สารบัญ.....	ช
สารบัญตาราง.....	ฌ
สารบัญรูป.....	ญ
บทที่	
1 บทนำ.....	1
1.1 สภาวะความเป็นมา แนวทางเหตุผล และปัญหา.....	1
1.2 วัตถุประสงค์ของงานวิจัย.....	5
1.3 ขอบเขตของการวิจัย.....	5
1.4 ขั้นตอนการทำวิจัย.....	6
1.5 ประโยชน์ที่คาดว่าจะได้รับ.....	6
2 งานวิจัยและทฤษฎีที่เกี่ยวข้อง.....	7
2.1 งานวิจัยที่เกี่ยวข้อง.....	7
2.2 ทฤษฎีที่เกี่ยวข้องกับงานวิจัย.....	20
3 วิธีการดำเนินงานวิจัย.....	32
3.1 ขั้นตอนการวิเคราะห์และออกแบบระบบ.....	33
3.2 การนำเสนอแนวคิดเพื่อพัฒนาต้นแบบที่สามารถแปลงดัชนี (Index) บนกล่อง เอกสารให้เป็นตัวอักษร (Proposed Method).....	40
3.3 การวัดประสิทธิภาพและความถูกต้อง.....	45
3.4 เครื่องมือที่ใช้พัฒนา.....	45

สารบัญ (ต่อ)

บทที่	หน้า
4 ผลของการวิจัย	46
4.1 ผลการสังเคราะห์รูปแบบการจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถ ลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร	47
4.2 ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร	48
5 สรุปและอภิปรายผลการวิจัย ข้อเสนอแนะ	55
5.1 สรุปและอภิปรายผลการวิจัย	55
5.2 ปัญหาและอุปสรรคในการวิจัย	56
5.3 ข้อเสนอแนะ	58
บรรณานุกรม	60
ภาคผนวก	63
ภาคผนวก ก. ตัวอย่างข้อมูลลายมือที่จัดเก็บเพื่อทำ Training Data	64
ภาคผนวก ข. ตัวอย่างข้อมูลกล่องเอกสารที่นำมาใช้ทำเป็น Testing Data	69
ภาคผนวก ค. การเตรียมข้อมูลและวิธีการทดลอง	71
ภาคผนวก ง. เอกสารที่นำเสนอวันสอบป้องกันสารนิพนธ์	76
ประวัติผู้เขียนสารนิพนธ์	105

สารบัญตาราง

ตาราง	หน้า
2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา	12
3.1 ข้อมูลและรูปแบบการบันทึกข้อมูลของดัชนีกล่องเอกสาร	35
3.2 การแบ่งระดับของการเขียนพยานะภาษาไทย	36
3.3 การแบ่งระดับของการเขียนสระ วรรณยุกต์และตัวเลขภาษาไทย	37
3.4 การแบ่งระดับการเขียนภาษาอังกฤษตัวพิมพ์ใหญ่	39
3.5 การแบ่งระดับการเขียนภาษาอังกฤษตัวพิมพ์เล็ก ตัวเลขอารบิกและสัญลักษณ์	39
4.1 การประเมินผลภาพรวมของข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่อง	49
4.2 ผลข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า	50
4.3 ผลข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร	50
4.4 ผลความถูกต้องของข้อมูลลายมือที่ปรากฏในแต่ละช่องของกล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า	51
4.5 ผลความถูกต้องของข้อมูลลายมือที่ปรากฏในแต่ละช่องของกล่องบริษัทที่ให้บริการจัดเก็บเอกสาร	52
4.6 ผลการทำงานแบบเดิมเทียบกับการใช้โปรแกรมเข้ามาช่วยในการทำงาน	53

TNI

TAHITI - NICHI INSTITUTE OF TECHNOLOGY

สารบัญรูป

รูป	หน้า
1.1 ตัวอย่างรายละเอียดข้อมูลที่ต้องบันทึกด้านข้างกล่องเอกสาร	2
1.2 ตัวอย่างกล่องเอกสารที่เขียนด้วยอักษรภาษาไทยและตัวเลขอารบิก.....	2
1.3 ตัวอย่างกล่องเอกสารที่เขียนด้วยอักษรภาษาอังกฤษและตัวเลขอารบิก.....	3
1.4 ตัวอย่างกล่องเอกสารที่เขียนด้วยอักษรภาษาไทย ภาษาอังกฤษและตัวเลขอารบิก.....	3
2.1 พิกเซลขนาด 1x1	21
2.2 พิกเซลที่ปรากฏอยู่ในรูปภาพ.....	22
2.3 ตัวอย่างการเพิ่มขนาดของวัตถุด้วยวิธีการขยายภาพ (Dilation).....	25
2.4 ตัวอย่างการลดขนาดของวัตถุด้วยวิธีการกร่อนภาพ (Erosion).....	26
2.5 ตัวอย่างการขยายภาพที่ได้จากการกร่อนภาพด้วยวิธีการเปิดภาพ (Opening).....	26
2.6 ตัวอย่างการกร่อนภาพที่ได้จากการขยายภาพด้วยวิธีการปิดภาพ (Closing).....	27
2.7 หน้าเว็บไซต์ของ Tesseract OCR.....	30
2.8 ขั้นตอนการทำงานของ Tesseract OCR Engine	30
3.1 ขั้นตอนเดิมของกระบวนการบันทึกข้อมูลดัชนี.....	32
3.2 ขอบเขตขั้นตอนการดำเนินงานวิจัย	33
3.3 ตัวอย่างรูปแบบของกล่องที่จะจัดทำดัชนี.....	33
3.4 บรรทัดของการเขียนตัวอักษรภาษาไทย	35
3.5 ระดับการเขียนตัวอักษรภาษาอังกฤษแบบตัวพิมพ์ใหญ่.....	38
3.6 ระดับการเขียนตัวอักษรภาษาอังกฤษแบบตัวพิมพ์เล็ก.....	38
3.7 ระดับการเขียนตัวอักษรภาษาอังกฤษแบบตัวพิมพ์ใหญ่และตัวพิมพ์เล็ก	38
3.8 แผนภาพการนำเสนอแนวคิดเพื่อพัฒนาต้นแบบ	41
3.9 ขั้นตอนการเตรียมตัวอักษรต้นแบบ (Training Data).....	42
3.10 ขั้นตอนการทำงานการปรับปรุงคุณภาพของภาพ (Pre-Processing).....	43
3.11 การหา Segmentation ของตัวอักษรที่จะทดลอง	44
3.12 ขั้นตอนการตรวจสอบความถูกต้องโดยมนุษย์ (Human Intervention).....	44
4.1 แนวคิดการจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรม.....	47
4.2 ตัวอย่างกล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า	48
4.3 ตัวอย่างกล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร	49

สารบัญรูป (ต่อ)

รูป	หน้า
5.1 ภาพเปรียบเทียบกล่องเอกสารที่มีแสงเงาติดยึดในภาพ	57
5.2 ภาพกล่องเอกสารที่เขียนตัวอักษรในแต่ละช่องไม่เป็นไปตามประเภทที่ระบุไว้	57
5.3 ภาพกล่องเอกสารที่เขียนตัวอักษรต่อเนื่องกัน	58
ก.1 ตัวอย่างการเขียนตัวอักษรภาษาไทยที่เป็นพยัญชนะ	65
ก.2 ตัวอย่างการเขียนตัวอักษรภาษาไทยที่เป็นสระและวรรณยุกต์	65
ก.3 ตัวอย่างการเขียนตัวอักษรภาษาไทยที่เป็นตัวเลขไทย	66
ก.4 ตัวอย่างการเขียนตัวอักษรภาษาอังกฤษตัวพิมพ์ใหญ่	66
ก.5 ตัวอย่างการเขียนตัวอักษรภาษาอังกฤษตัวพิมพ์เล็ก	67
ก.6 ตัวอย่างการเขียนตัวอักษรที่เป็นเครื่องหมาย	67
ก.7 ตัวอักษรที่เป็นตัวเลขอารบิก	68
ข.1 ภาพแสดงกล่องที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า	70
ข.2 ภาพแสดงกล่องเอกสารของบริษัทที่ให้บริการจัดเก็บเอกสาร	70
ค.1 ขั้นตอนการทดลองส่วนที่ 1 Training Data และส่วนที่ 2 Testing Data	72
ค.2 ตัวอย่างการจัดกลุ่มไฟล์ภาพตัวอักษร	73
ค.3 การนำตัวอักษรเข้า MATLAB OCR Trainer	73
ค.4 การนำภาพของคำที่พบบ่อยเข้า MATLAB OCR Trainer	74
ค.5 การแปลงรูปภาพสีเป็นรูปภาพขาวดำ โดยใช้ Adaptive Thresholding	74
ค.6 การทำ Segmentation จะใช้วิธีเลือก Region of Interest (roi)	75
ค.7 การทำ Pattern Matching	75
ง.1 เอกสารนำเสนอสารนิพนธ์หน้าที่ 1	77
ง.2 เอกสารนำเสนอสารนิพนธ์หน้าที่ 2	77
ง.3 เอกสารนำเสนอสารนิพนธ์หน้าที่ 3	78
ง.4 เอกสารนำเสนอสารนิพนธ์หน้าที่ 4	78
ง.5 เอกสารนำเสนอสารนิพนธ์หน้าที่ 5	79
ง.6 เอกสารนำเสนอสารนิพนธ์หน้าที่ 6	79
ง.7 เอกสารนำเสนอสารนิพนธ์หน้าที่ 7	80
ง.8 เอกสารนำเสนอสารนิพนธ์หน้าที่ 8	80

สารบัญรูป (ต่อ)

รูป	หน้า
ง.9 เอกสารนำเสนอสารนิพนธ์หน้าที่ 9.....	81
ง.10 เอกสารนำเสนอสารนิพนธ์หน้าที่ 10.....	81
ง.11 เอกสารนำเสนอสารนิพนธ์หน้าที่ 11.....	82
ง.12 เอกสารนำเสนอสารนิพนธ์หน้าที่ 12.....	82
ง.13 เอกสารนำเสนอสารนิพนธ์หน้าที่ 13.....	83
ง.14 เอกสารนำเสนอสารนิพนธ์หน้าที่ 14.....	83
ง.15 เอกสารนำเสนอสารนิพนธ์หน้าที่ 15.....	84
ง.16 เอกสารนำเสนอสารนิพนธ์หน้าที่ 16.....	84
ง.17 เอกสารนำเสนอสารนิพนธ์หน้าที่ 17.....	85
ง.18 เอกสารนำเสนอสารนิพนธ์หน้าที่ 18.....	85
ง.19 เอกสารนำเสนอสารนิพนธ์หน้าที่ 19.....	86
ง.20 เอกสารนำเสนอสารนิพนธ์หน้าที่ 20.....	86
ง.21 เอกสารนำเสนอสารนิพนธ์หน้าที่ 21.....	87
ง.22 เอกสารนำเสนอสารนิพนธ์หน้าที่ 22.....	87
ง.23 เอกสารนำเสนอสารนิพนธ์หน้าที่ 23.....	88
ง.24 เอกสารนำเสนอสารนิพนธ์หน้าที่ 24.....	88
ง.25 เอกสารนำเสนอสารนิพนธ์หน้าที่ 25.....	89
ง.26 เอกสารนำเสนอสารนิพนธ์หน้าที่ 26.....	89
ง.27 เอกสารนำเสนอสารนิพนธ์หน้าที่ 27.....	90
ง.28 เอกสารนำเสนอสารนิพนธ์หน้าที่ 28.....	90
ง.29 เอกสารนำเสนอสารนิพนธ์หน้าที่ 29.....	91
ง.30 เอกสารนำเสนอสารนิพนธ์หน้าที่ 30.....	91
ง.31 เอกสารนำเสนอสารนิพนธ์หน้าที่ 31.....	92
ง.32 เอกสารนำเสนอสารนิพนธ์หน้าที่ 32.....	92
ง.33 เอกสารนำเสนอสารนิพนธ์หน้าที่ 33.....	93
ง.34 เอกสารนำเสนอสารนิพนธ์หน้าที่ 34.....	93
ง.35 เอกสารนำเสนอสารนิพนธ์หน้าที่ 35.....	94

สารบัญรูป (ต่อ)

รูป	หน้า
ง.36 เอกสารนำเสนอสารนิพนธ์หน้าที่ 36.....	94
ง.37 เอกสารนำเสนอสารนิพนธ์หน้าที่ 37.....	95
ง.38 เอกสารนำเสนอสารนิพนธ์หน้าที่ 38.....	95
ง.39 เอกสารนำเสนอสารนิพนธ์หน้าที่ 39.....	96
ง.40 เอกสารนำเสนอสารนิพนธ์หน้าที่ 40.....	96
ง.41 เอกสารนำเสนอสารนิพนธ์หน้าที่ 41.....	97
ง.42 เอกสารนำเสนอสารนิพนธ์หน้าที่ 42.....	97
ง.43 เอกสารนำเสนอสารนิพนธ์หน้าที่ 43.....	98
ง.44 เอกสารนำเสนอสารนิพนธ์หน้าที่ 44.....	98
ง.45 เอกสารนำเสนอสารนิพนธ์หน้าที่ 45.....	99
ง.46 เอกสารนำเสนอสารนิพนธ์หน้าที่ 46.....	99
ง.47 เอกสารนำเสนอสารนิพนธ์หน้าที่ 47.....	100
ง.48 เอกสารนำเสนอสารนิพนธ์หน้าที่ 48.....	100
ง.49 เอกสารนำเสนอสารนิพนธ์หน้าที่ 49.....	101
ง.50 เอกสารนำเสนอสารนิพนธ์หน้าที่ 50.....	101
ง.51 เอกสารนำเสนอสารนิพนธ์หน้าที่ 51.....	102
ง.52 เอกสารนำเสนอสารนิพนธ์หน้าที่ 52.....	102
ง.53 เอกสารนำเสนอสารนิพนธ์หน้าที่ 53.....	103
ง.54 เอกสารนำเสนอสารนิพนธ์หน้าที่ 54.....	103
ง.55 เอกสารนำเสนอสารนิพนธ์หน้าที่ 55.....	104

บทที่ 1

บทนำ

1.1 สภาวะความเป็นมา แนวทางเหตุผล และปัญหา

ปัญหาเรื่องการบริหารจัดการเอกสาร เป็นเรื่องที่สามารถพบได้ทั่วไป ไม่ว่าจะเป็นที่บ้าน ที่ทำงาน ซึ่งพบว่าในหลายๆ บริษัทมีปริมาณการใช้งานเอกสารที่เพิ่มมากขึ้น ตามอัตราการเจริญเติบโตของบริษัท ทำให้มีเอกสารเป็นจำนวนมากที่จำเป็นจะต้องจัดเก็บตามกฎหมาย แต่จะต้องแบ่งแยกประเภทการจัดเก็บออกเป็นหลากหลายประเภท ทำให้การค้นหาเอกสารแต่ละครั้งนั้นมีความยุ่งยาก อันเกิดจากการจัดเก็บเอกสารปะปนกัน จัดเก็บไม่ถูกที่ ทำให้ค้นหาไม่สะดวก ใช้ระยะเวลาในการค้นหาเอกสารนาน บ่อยครั้งยังพบเอกสารชำรุด สูญหาย อันเนื่องจากการไม่ได้มีการดูแลอย่างเป็นระบบ

บริษัทหลายแห่งได้เล็งเห็นความสำคัญและความจำเป็นในการจัดเก็บเอกสารทั้งหมดให้มีความเป็นระเบียบเรียบร้อย เป็นหมวดหมู่ เพื่อความสะดวกในการสืบค้น รวมไปถึงการบริหารจัดการพื้นที่ในบริษัทเพื่อลดต้นทุนในการจัดเก็บเอกสารที่มีการใช้งานน้อย แต่ต้องเก็บรักษาเอกสารตามอายุของกฎหมาย เช่น เอกสารทางบัญชีและการเงิน มีระยะเวลาจัดเก็บ 10 ปี เอกสารประวัติพนักงาน 2 ปี เป็นต้น หลายๆ บริษัทจึงได้มีการมองหาแบบการจัดเก็บเอกสารนอกสถานที่ โดยใช้ผู้เชี่ยวชาญจากภายนอก (Outsource) ซึ่งจะอยู่ในรูปแบบของบริษัทที่ให้บริการบริหารจัดการเอกสารและข้อมูล เข้ามาให้บริการบริหารจัดการเอกสารและรูปแบบในการบริหารจัดการเอกสารอย่างเป็นระบบ เช่น วิธีการจัดเก็บเอกสาร การจัดทำดัชนีเพื่อสืบค้นเอกสาร การบริการขนส่งเอกสาร และการจัดทำรายงานสรุปของเอกสารที่ฝาก เป็นต้น เพื่อให้บริษัทสามารถควบคุมเรื่องค่าใช้จ่ายในการบริหารจัดการเอกสารได้ โดยส่วนมากบริษัทผู้ให้บริการจะให้การจัดเก็บเอกสารในรูปแบบกล่องเอกสาร กล่าวคือ บริษัทจะต้องทำการบรรจุเอกสารที่ต้องการจัดเก็บลงกล่องเอกสารแล้วจัดทำดัชนีสืบค้นก่อนที่จะแจ้งผู้ให้บริการเข้ามาทำรับเอกสารไปจัดเก็บ หลังจากที่ได้รับกล่องเอกสารจากบริษัทแล้ว ก่อนที่จะดำเนินการจัดเก็บเข้าชั้นวางกล่องเอกสารทางผู้ให้บริการจะมีการดำเนินการบันทึกข้อมูลรายละเอียดที่ปรากฏอยู่ด้านข้างกล่อง ได้แก่ เลขบาร์โค้ดกล่อง รหัสลูกค้า ชื่อกล่องลูกค้า วันที่เอกสารครบกำหนด และรายละเอียดของกล่องเอกสาร ดังรูปที่ 1.1

Barcode Position สำหรับติดบาร์โค้ด	Customer Code รหัสลูกค้า
Reference Code ชื่อกล่องลูกค้า	Destruction Date ปี(ค.ศ.)/เดือน/ทำลาช _____/____/____
Document Description รายละเอียดของกล่องเอกสาร	

รูปที่ 1.1 ตัวอย่างรายละเอียดข้อมูลที่ต้องบันทึกด้านข้างกล่องเอกสาร

ข้อมูลที่ปรากฏนั้น มีทั้งที่เขียนด้วยลายมือ ตัวพิมพ์จากเครื่องพิมพ์ และบาร์โค้ด ซึ่งในแต่ละกล่องจะมีรูปแบบของการจัดทำดัชนี (INDEX) ที่เป็นเอกลักษณ์เฉพาะบุคคล ซึ่งส่งผลกระทบทำให้การบันทึกข้อมูลต้องใช้เวลาในการจัดทำ ข้อมูลที่จะดำเนินการบันทึกในแต่ละวันจะแบ่งเป็น 3 กลุ่มดังต่อไปนี้

กลุ่มที่ 1 ประมาณ 10% เป็นข้อมูลที่มีตัวอักษรภาษาไทยและตัวเลขอารบิก ดังรูปที่ 1.2

Barcode Position สำหรับติดบาร์โค้ด 00018DG 10000001734	Customer Code รหัสลูกค้า 00018DG
Reference Code ชื่อกล่องลูกค้า สมอ 07	Destruction Date ปี(ค.ศ.)/เดือน/ทำลาช _____/____/____
Document Description รายละเอียดของกล่องเอกสาร กล่อง 1	

รูปที่ 1.2 ตัวอย่างกล่องเอกสารที่เขียนด้วยอักษรภาษาไทยและตัวเลขอารบิก

กลุ่มที่ 2 ประมาณ 50% เป็นข้อมูลที่มีภาษาอังกฤษและตัวเลขอารบิก ดังรูปที่ 1.3

B 00018DG  10000001735	Customer Code รหัสลูกค้า 00018DG
Reference Code ชื่อกล่องลูกค้า สมอ 07	Destruction Date ปี(ค.ศ.)/เดือน/ทำลาช ____/____/____
Document Description รายละเอียดของกล่องเอกสาร Box 2	

รูปที่ 1.3 ตัวอย่างกล่องเอกสารที่เขียนด้วยอักษรภาษาอังกฤษและตัวเลขอารบิก

กลุ่มที่ 3 ประมาณ 40% เป็นข้อมูลที่มีตัวอักษรภาษาไทย ตัวอักษรภาษาอังกฤษและตัวเลขอารบิก ดังรูปที่ 1.4

B 00018DG  10000001717	Customer Code รหัสลูกค้า 00018DG
Reference Code ชื่อกล่องลูกค้า เทียบคุณวุฒิ	Destruction Date ปี(ค.ศ.)/เดือน/ทำลาช ____/____/____
Document Description รายละเอียดของกล่องเอกสาร ปังบประมาณ 2558 ลำดับที่ 1-70	

รูปที่ 1.4 ตัวอย่างกล่องเอกสารที่เขียนด้วยอักษรภาษาไทย ภาษาอังกฤษและตัวเลขอารบิก

ขั้นตอนการจัดทำดัชนีบนกล่องเอกสารจะประกอบด้วย ขั้นแรก นำกล่องมาวางที่จุดที่จะทำการบันทึกข้อมูล ใช้เวลา 10 วินาที ขั้นที่ 2 พนักงานทำการบันทึกข้อมูลจากรูปภาพ ใช้เวลา 30 วินาที ขั้นที่ 3 พนักงานทำการ QC ข้อมูลที่ทำการบันทึกใช้เวลา 15 วินาที ขั้นที่ 4 นำกล่องไปวางที่จุด Outbound ใช้เวลา 5 วินาที รวมระยะเวลาที่ใช้ตั้งแต่ขั้นตอนที่ 1 ถึง 4 ใช้เวลาการดำเนินการใน 1 กล่อง ประมาณ 60 วินาที (1 นาที) ขั้นที่ 5 นำข้อมูลเข้าสู่ระบบ โดยการบันทึกข้อมูลจะใช้เวลาเฉลี่ย 50 ของทั้งกระบวนการ โดยในแต่ละวันจะมีกล่องเอกสารที่รับเข้ามาประมาณ 1,000 กล่อง ซึ่งทำให้ต้องใช้พนักงานจำนวน 3 คน ทำการบันทึกข้อมูลก่อนที่จะนำ เข้าระบบและนำกล่องเอกสาร จัดเก็บเข้าชั้นวางได้

ในปัจจุบันความก้าวหน้าทางด้านเทคโนโลยีด้านการประมวลผลภาพ (Digital Image Processing) ได้มีนักวิจัยได้นำหลักการ OCR (Optical Character Recognition) และ Handwriting Recognition มาดำเนินการวิจัยในรูปแบบและวิธีการที่แตกต่างกันออกไป ซึ่งสามารถแบ่งงานวิจัย ออกเป็น 2 ส่วน ได้แก่

ส่วนที่ 1 งานวิจัยที่ใช้หลักการ OCR (Optical Character Recognition) มาประมวลผล ภาพที่เป็นตัวอักษรที่มาจากเครื่องพิมพ์ให้อยู่ในรูปแบบตัวอักษรที่ใช้งานบนคอมพิวเตอร์ได้ ตัวอย่างเช่น มีงานวิจัยเรื่อง Optical Character Recognition ได้นำเสนอการทำรู้จำอักขระภาพ โดยใช้ Open Source Engine ของกูเกิล (Google) ที่ชื่อว่า TESSERACT สำหรับใช้งานใน สมาร์ทโฟนบนระบบปฏิบัติการแอนดรอยด์ โดย TESSERACT สามารถรู้จำอักขระได้ทั้งตัวอักษร ที่มาจากเครื่องพิมพ์และลายมือเขียน จากนั้นจะแปลงให้เป็นตัวอักษรที่คอมพิวเตอร์สามารถอ่านได้ แล้วจึงนำตัวอักษรไปแปลงให้อยู่ในรูปแบบของการออกเสียง (Text to Speech) ซึ่งปัจจุบัน Tesseract เป็นเครื่องมือพื้นฐานของระบบที่ใช้หลักการ OCR (Optical Character Recognition) ที่ประมวลผล ภาพตัวอักษรให้อยู่ในรูปแบบตัวอักษรที่ใช้งานบนคอมพิวเตอร์ได้

ส่วนที่ 2 งานวิจัยที่ใช้หลักการ Handwriting Recognition ในการประมวลผลภาพจาก การเขียนด้วยลายมือให้อยู่ในรูปแบบตัวอักษรที่ใช้งานบนคอมพิวเตอร์ได้ ตัวอย่างเช่น งานวิจัยด้าน การแยกภาพตัวอักษรลายมือเขียนภาษาไทยแบบอัตโนมัติโดยใช้การวิเคราะห์ถดถอยเชิงเส้นและการ เรียนรู้ด้วยต้นไม้ตัดสินใจ ได้นำเสนอหลักการสำหรับแยกภาพตัวอักษรลายมือเขียนที่อยู่ติดกันแบบ สัมผัสในเอกสารภาพตัวอักษรออกจากกัน โดยใช้การวิเคราะห์ถดถอยเชิงเส้น (Linear Regression Analysis) และความกว้างของมัธยฐานของตัวอักษร (Median Character Width) สำหรับหากลุ่ม ของตัวอักษรที่อยู่ติดกันในแนวตั้งและแนวนอนตามลำดับ และแยกตัวอักษรโดยใช้แบบจำลองต้นไม้ ตัดสินใจที่ได้มาจากการฝึกฝน พบว่าความถูกต้องของการแยกตัวอักษรลายมือเขียนภาษาไทยเป็น ร้อยละ 90.44 นอกจากนี้ยังมีงานวิจัยการจดจำลายมือเขียนตัวอักษรไทยด้วยแผนผังคุณลักษณะ จัดการตัวเอง โดยระบบดังกล่าวเป็นการผสมผสานระหว่าง เครือข่ายประสาทเทียมแบบการเรียนรู้

เวกเตอร์ควอนไทเซชันและเครือข่ายไปข้างหน้า ซึ่งเป็นระบบการจดจำลายมือเขียนตัวอักษรไทยที่สามารถจดจำลายมือเขียนภาษาไทยที่มีลักษณะรูปร่างผิดเพี้ยนไปจากเดิมได้ อีกทั้งยังสามารถเรียนรู้ลักษณะลายมือเขียนใหม่ๆ เพิ่มเติมได้ ซึ่งระบบนี้มีความถูกต้องแม่นยำในการรู้จำตัวอักษร 91.09%

จากตัวอย่างงานวิจัยดังกล่าวทำให้มีแนวความคิดที่จะหาวิธีการจัดเก็บข้อมูลดังกล่าวด้วยการนำหลักการของ OCR (Optical Character Recognition) และ Handwriting Recognition เข้ามาประยุกต์ใช้ในการแปลงและบันทึกข้อมูล โดยทำการเก็บข้อมูลเป็นรูปภาพและนำมาประมวลผลเพื่อทำการแปลงข้อความที่ปรากฏในรูปภาพนั้นให้เป็นตัวอักษรที่สามารถบันทึกลงในคอมพิวเตอร์ได้ทันที โดยจะสามารถแปลงข้อความที่มาจากเครื่องพิมพ์และลายมือเขียนได้ ซึ่งจะรองรับการใช้งานภาษาไทยและภาษาอังกฤษ

งานวิจัยนี้จะนำเสนอกรอบแนวคิดและต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถแปลงดัชนี (Index) ที่เป็นข้อความที่อยู่บนกล่องเอกสารให้เป็นตัวอักษรที่สามารถใช้งานได้บนคอมพิวเตอร์เพื่อลดขั้นตอนการดำเนินงานของผู้ปฏิบัติงานและต้นทุนด้านการดำเนินงานของฝ่ายปฏิบัติการ

1.2 วัตถุประสงค์ของงานวิจัย

1.2.1 เพื่อจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้

1.2.2 เพื่อสามารถแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้มีความถูกต้อง (Accuracy) อย่างน้อย 80%

1.2.3 เพื่อลดระยะเวลาที่ใช้ในขั้นตอนการบันทึกข้อมูลดัชนี (Index) เป็น 30% ของกระบวนการเดิม

1.3 ขอบเขตของการวิจัย

1.3.1 งานวิจัยเป็นเพียงต้นแบบ (Prototype) เพื่อใช้ทดสอบการแปลงข้อมูลดัชนี (Index) ที่ปรากฏอยู่บนกล่องเอกสารให้เป็นตัวอักษรที่สามารถใช้งานบนคอมพิวเตอร์ได้

1.3.2 รองรับการแปลงข้อมูลที่มีรูปแบบและข้อความที่ได้กำหนดไว้ทั้งภาษาไทยและภาษาอังกฤษ

1.3.3 ข้อมูลตัวอักษรที่มีการเขียนติดกันหรือถูกเขียนทับลงบนกรอบของช่องที่ให้เขียนจะไม่อยู่ในขอบเขตการวิจัยนี้

1.3.4 การนำเข้าข้อมูล (Input Data) จะใช้ไฟล์ข้อมูลที่ได้จากกล้องถ่ายภาพ

1.3.5 สภาพของกล่องเอกสารที่นำมาทำการทดสอบจะต้องเป็นกล่องที่อยู่สภาพที่สมบูรณ์

เท่านั้น

1.3.6 งานวิจัยนี้จะครอบคลุมเฉพาะเรื่องของการแปลงข้อมูลดัชนี (Index) ลายมือที่ปรากฏอยู่บนกล่องเท่านั้น

1.3.7 การทดสอบและประเมินผลจะทำการประเมินผลโดยการเปรียบเทียบผลของข้อมูลตัวอักษรที่แปลได้จากการทดสอบกับข้อมูลตัวอักษรของดัชนีที่ปรากฏในแต่ละกล่อง

1.4 ขั้นตอนการทำวิจัย

1.4.1 ศึกษาข้อมูลเกี่ยวกับหลักการบริหารจัดการเก็บเอกสารในรูปแบบกล่อง

1.4.2 ศึกษางานวิจัยที่เกี่ยวข้องกับหลักการ OCR (Optical Character Recognition)

1.4.3 ศึกษาวิธีการพัฒนาตัวต้นแบบโปรแกรม

1.4.4 วางแผนและออกแบบตัวต้นแบบโปรแกรม

1.4.5 จัดทำตัวต้นแบบสำหรับใช้พัฒนาโปรแกรมในส่วนของการประมวลผลภาพ การทำคลังข้อมูลภาพตัวอักษร การรู้จำและการแปลความหมาย

1.4.6 ทดสอบการใช้งานของตัวต้นแบบที่จะใช้พัฒนาโปรแกรม

1.4.7 วิเคราะห์ ประเมินและสรุปผล

1.4.8 จัดทำรายงานและนำเสนอผลงาน

1.4 ประโยชน์ที่คาดว่าจะได้รับ

1.4.1 เพื่อลดขั้นตอนการบันทึกข้อมูลดัชนี (Index) บนกล่องเอกสารด้วยบุคคล

1.4.2 เพื่อลดต้นทุนของฝ่ายปฏิบัติการ 50%

1.4.3 เพื่อลดข้อผิดพลาดจากการบันทึกข้อมูลด้วยพนักงาน

TNI

บทที่ 2

งานวิจัยและทฤษฎีที่เกี่ยวข้อง

ในบทนี้จะนำเสนอแนวคิด ทฤษฎี และงานวิจัยหรือวรรณกรรม เพื่อเป็นแนวทางในการศึกษาวิจัย การเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขรภาพลายมือ โดยได้ทำการได้ศึกษาเอกสาร ทฤษฎี งานวิจัย หรือวรรณกรรมที่เกี่ยวข้อง ซึ่งมีสาระสำคัญ ดังต่อไปนี้

2.1 งานวิจัยที่เกี่ยวข้อง

งานวิจัยที่เกี่ยวข้องกับหลักการ Handwriting Recognition ในการประมวลผลภาพจากการเขียนด้วยลายมือให้อยู่ในรูปแบบตัวอักษรที่ใช้งานบนคอมพิวเตอร์ ได้มีงานวิจัยที่ผ่านมาได้ทำการวิจัยการประมวลผลภาพจากการเขียนด้วยลายมือ โดยสามารถแบ่งกลุ่มตามขั้นตอนการประมวลผลภาพออกเป็น 3 กลุ่ม ได้แก่ กลุ่มที่ 1 ใช้อัลกอริทึมและเทคนิคที่กระบวนการ Pre-Processing กลุ่มที่ 2 ใช้อัลกอริทึมและเทคนิคที่กระบวนการ Processing และกลุ่มที่ 3 ใช้อัลกอริทึมและเทคนิคที่กระบวนการ Post-Processing โดยมีรายละเอียด ดังนี้

กลุ่มที่ 1 ใช้อัลกอริทึมและเทคนิคที่กระบวนการ Pre-Processing

วิเชษฐ์รจน์ เอี่ยมสำอางค์ และรัชฎา คงคะจันทร์ [1] ได้ทำการศึกษาการแยกภาพตัวอักษรลายมือเขียนภาษาไทยแบบอัตโนมัติ โดยใช้การวิเคราะห์ถดถอยเชิงเส้น (Linear Regression Analysis) และการเรียนรู้ด้วยต้นไม้ตัดสินใจ (Decision Tree Learning) ซึ่งเป็นกระบวนการเตรียมพร้อมของภาพข้อมูลก่อนที่จะดำเนินการรู้จำลายมือเขียน (Handwritten Recognition) ลักษณะของการเขียนภาษาไทยมีความแตกต่างจากภาษาอังกฤษ จะแบ่งออกบรรทัดของตัวอักษรที่เขียนได้เป็น 4 ระดับ โดยสามารถติดกันได้ในระดับเดียวกันและข้ามระดับทั้งในแนวนอนและแนวตั้ง จึงได้มีวิธีการวิเคราะห์ตัวอักษรด้วยคุณลักษณะต่างๆ ของตัวอักษรไทย เพื่อแยกตัวอักษรที่ติดกันในแนวนอนและแนวตั้ง โดยใช้เส้นการวิเคราะห์การถดถอยสำหรับตัดแบ่งระดับพยัญชนะกับสระ และทำการตัดแบ่งตัวอักษร (Segmentation) ก่อนการรู้จำตัวอักษรตามหลักของการรู้จำตัวอักษรไทยโดยใช้เทคนิคการใช้ต้นไม้ตัดสินใจเข้ามาช่วย โดยพบว่ามีความถูกต้องแม่นยำในการแปลงข้อมูลอยู่ที่ 90.44% ซึ่งอุปสรรคที่ส่งผลกระทบทำให้การแปลมีความแม่นยำลดลงเกิดจาก การเขียนตัวอักษรที่ไม่มีหัวทำให้เกิดความคลาดเคลื่อนในการรู้จำลดลง กรณีมีสระข้างปนอยู่ เช่น สระเอ สระโอ จะทำให้เกิดความผิดพลาดในการตัดตัวอักษรแต่ละตัวได้

พาฝัน ดวงไพศาล และคณะ [2] ได้ทำการศึกษาการจดจำการเขียนลายมือตัวเลขไทยด้วยเทคนิคโครงข่ายการส่งข้อมูลแบบไม่ย้อนกลับ (Feed-Forward Neural Network) เปรียบเทียบกับ

เทคนิค K-Nearest Neighbor (KNN) โดยสามารถจดจำลายมือเขียนตัวเลขไทยที่มีลายมือแตกต่างกันไป จากเดิมได้ โดยภาพที่ Input เข้าไปจะใช้เครื่องสแกนเนอร์ สแกนที่ความละเอียด 200 dpi แล้วใช้โปรแกรม Photoshop ทำการตัดภาพตัวอักษรให้มีขนาด 200 x 200 พิกเซล แล้วใช้โปรแกรม MatLab ในการคำนวณเพื่อประมวลผล พบว่ามีความถูกต้องแม่นยำอยู่ที่ 87.86% แต่ขนาดข้อมูลก็นำมาคำนวณมีขนาดใหญ่เกินไป ทำให้ต้องใช้หน่วยความจำสูงมากในการประมวลผล โดยได้มีข้อเสนอแนะว่าสามารถจะทำการลดขนาดของข้อมูลได้ 2 วิธีวิธีที่ 1 ลดขนาดภาพโดยตรงโดยใช้เทคนิค Nearest-Neighbor Interpolation และ Bilinear Interpolation และวิธีที่ 2 ลดขนาดของข้อมูล (Dimension Reduction)

ธีรเดช ราชไพบูลย์ และนุชนาฏ สัตยาภิ [3] ได้ทำการศึกษาการรู้จำลายมือเลขประจำตัวผู้สอบในกระดาษคำตอบของข้อสอบแบบปรนัย โดยมีการกำหนดให้ผู้เข้าสอบจะต้องเขียนตัวเลขอารบิกด้วยปากกาถูกลิ้นสีน้ำเงินไว้ที่มุมบนซ้ายที่กระดาษคำตอบ โดยใช้เทคนิคโปรเจกชันโปรไฟล์ 4 แกน ได้แก่ แนวแถว (Horizontal) แนวคอลัมน์ (Vertical) แนวทแยงซ้าย (Left Diagonal) และแนวทแยงขวา (Right Diagonal) มาทำการวิเคราะห์โครงสร้าง (Structural Analysis) ของตัวเลขที่ปรากฏอยู่ในภาพกระดาษคำตอบ เพื่อจะดำเนินการสกัดเอาส่วนที่สำคัญ (Feature Extraction) ของตัวเลขออกมา แล้วนำมาจับคู่กับรูปแบบ (Template Matching) ที่ได้ทำการกำหนดไว้ในคลังข้อมูลจากการใช้เทคนิคดังกล่าวพบว่ามีความถูกต้องแม่นยำอยู่ที่ 77% สาเหตุที่ทำให้ค่าความแม่นยำลดลงเนื่องจากภาพที่นำมาประมวลผลมีขนาดภาพใหญ่ทำให้ประมวลผลช้า ลายมือที่เขียนตัวเลขบนกระดาษคำตอบนั้นใช้ปากกาถูกลิ้นเขียน ความหนักเบาในการเขียนและชนิดหมึกของปากกาที่ใช้ ออกไม่เท่ากัน ทำให้เวลาที่สแกนเข้าไปอาจทำให้เส้นขาดหายได้ จึงควรใช้ปากกาเคมีหรือดินสอเขียนได้ ในขั้นตอนทำ Feature Extraction ใช้วิธีการพิจารณาแบบเส้นตรง (Linear) อาจจะไม่ดีเท่าการใช้การพิจารณาแบบเชิงต้นไม้ (Tree) เพื่อให้ทำการจัดกลุ่มของเลขบางชุดให้อยู่ในกลุ่มเดียวกันก่อนแล้วค่อยนำไปพิจารณาต่อจนได้ค่าที่เหมาะสมที่สุด และสุดท้ายการวิเคราะห์คำตอบจากโปรเจกชัน 4 แกน โดยใช้ค่าเฉลี่ยจะมีความน่าเชื่อถือน้อยที่สุด ดังนั้นจึงควรใช้แค่ค่าฐานนิยมและมัธยฐานแทนซึ่งจะมีความน่าเชื่อถือมากกว่า

อมรเทพ วิจิตรเจริญ และพฤษดี ศิริแสงตระกูล [4] ได้ทำการศึกษาการสกัดคุณลักษณะเด่นแบบผสมเพื่อรู้จำตัวเขียนอักษรธรรมอีสานด้วยโครงข่ายประสาทเทียมแบบแพร่ย้อนกลับด้วยวิธีผสมระหว่างการแบ่งโซนภาพ และฮิสโตแกรมโปรเจกชัน โดยโครงข่ายประสาทเทียมแบบแพร่ย้อนกลับ เป็นอัลกอริทึมการเรียนรู้แบบมีผู้ฝึกสอน เพื่อปรับปรุงค่าน้ำหนักของโครงข่ายประสาทเทียมให้เหมาะสมกับโครงข่ายนั้นๆ ด้วยการคำนวณหาค่าผิดพลาด (Error) จากผลลัพธ์ (Output) ซึ่งกระบวนการของการศึกษานี้ได้เน้นการใช้เทคนิค Feature Extraction เพื่อสกัดเอาคุณลักษณะเด่นของตัวอักษรออกมา โดยใช้อัลกอริทึม 2 แบบมาใช้ในขั้นตอนนี้ ได้แก่ อัลกอริทึมที่ 1 ความหนาแน่น

ของจุดพิกเซลด้วยวิธีแบ่งโซน อัลกอริทึมที่ 2 ฮิสโตแกรมโปรเจกชันในแนวตั้งและแนวนอน หลังจากที่ได้สกัดตัวอักษรได้แล้ว ได้นำเทคนิค เคโฟลด์ครอสวาเลชัน (K-Fold Cross Validation) มาทำการรู้จำตัวอักษรโดยแบ่งข้อมูลที่ใช้ในการทดสอบออกเป็นกลุ่มจำนวน K กลุ่ม (K-Folds) เพื่อให้ข้อมูลทุกตัวมีโอกาสเป็นชุดทดสอบและชุดสอนเพื่อป้องกันปัญหาการเลือกข้อมูลที่ดีและง่ายมาเป็นข้อมูลชุดทดสอบ จากการทดสอบพบว่าได้มีการเปรียบเทียบค่าความถูกต้องแม่นยำที่ได้เป็น 4 วิธี ได้แก่ วิธีการแบ่งโซน มีค่าความแม่นยำ 62.22% วิธีฮิสโตแกรมโปรเจกชัน มีค่าความแม่นยำ 52.24% วิธีผสมระหว่างการแบ่งโซนและฮิสโตแกรมโปรเจกชันแนวนอน มีค่าความแม่นยำ 63.87% และวิธีผสมระหว่างการแบ่งโซนและฮิสโตแกรมโปรเจกชันแนวตั้ง มีค่าความแม่นยำ 71.22% จะเห็นได้ว่าวิธีการผสมจะทำให้ได้ประสิทธิภาพในการประมวลผลดีกว่า สาเหตุที่ค่าความแม่นยำมีค่าไม่สูงมากเนื่องจากเอกสารต้นฉบับเป็นใบลานเก่าทำให้เกิด Noise รบกวนมากซึ่งทำให้ประสิทธิภาพการรู้จำลดลง และควรจะเน้นขั้นตอนการประมวลผลขั้นต้นให้มีความละเอียดยิ่งขึ้น ก่อนเข้าสู่กระบวนการสกัด

V. S. Tapkir and S. D. Shelke [5] ได้ทำการศึกษาการทำ OCR ลายมือเขียนตัวอักษร Marathi ซึ่งเทคนิค OCR แบบดั้งเดิมจะพบปัญหาเรื่องคุณภาพในการแปล โดยความผิดพลาดจะเกิดขึ้นในขั้นตอนการแยกส่วนของตัวอักษร (Character Segmentation) ซึ่งเป็นกระบวนการที่สำคัญในการทำ OCR เนื่องจากความถูกต้องแม่นยำในการแปลจะขึ้นอยู่กับ การใช้อัลกอริทึมการแยกส่วนของตัวอักษรที่ใช้กับส่วนนี้ จากการศึกษาพบว่า การแยกส่วนของภาพ (Segmentation) การเขียนอักษรเทวนาครี (Devanagari) ด้วยลายมือจะทำได้ยากกว่าเมื่อเทียบกับอักษรเทวนาครี (Devanagari) หรือตัวอักษรภาษาอังกฤษที่พิมพ์จากเครื่อง เนื่องจากการเขียนจะมีความหลากหลายกว่าและมีการเพิ่มขึ้นของรูปแบบตัวอักษร ซึ่งลายมือแต่ละคนเขียนตัวอักษรเดียวกัน แต่ระบบจะทำการรู้จำเป็นแบบใหม่ (New Pattern) และอาจจะมีการเขียนซ้อนกันเกิดขึ้น ในงานวิจัยนี้ได้นำเสนอการทำ การแยกส่วนของภาพ (Segmentation) การเขียนอักษรเทวนาครี (Devanagari) ด้วยอัลกอริทึม Projection Profiles โดยมีการนำเทคนิคการแยกส่วนของภาพมาใช้งานออกเป็น 3 ส่วน ส่วนที่ 1 เทคนิค Line Segmentation คือนำภาพ Original มาทำการหาและแยกบรรทัดของแต่ละข้อความ ส่วนที่ 2 เทคนิค Word Segmentation หลังจากที่ได้แยกเป็นแต่ละบรรทัดแล้วจะนำบรรทัดข้อความที่ได้มาทำการตัดเป็นคำๆ ส่วนที่ 3 เทคนิค Character Segmentation นำเอาคำแต่ละคำมาทำการแยกเป็นแต่ละตัวอักษร ซึ่งได้ค่าความแม่นยำในการแปลผลข้อมูลโดยแบ่งเป็น เทคนิค Line Segmentation = 100% และ Character Segmentation = 98%

M. Nayak and A. K. Nayak [6] ได้ทำการศึกษาการทำ OCR ตัวอักษร Odia โดยใช้ Tesseract OCR Engine และมี Training Data เพื่อช่วยในการแปลข้อมูล ซึ่งการ Training Tesseract Engine เพื่อทำให้ระบบรู้รูปแบบของตัวอักษร Odia ที่จะทำการแปลตัวอักษรจากพจนานุกรม

ข้อมูลที่มีการสร้างไว้ ซึ่งผลการทดสอบพบว่า ตัวอักษร Odia ขนาด 10pt 12pt 16pt มีความแม่นยำในการแปล 100% และตัวอักษร Utkal มีความแม่นยำ 98% (มีการทดลองผสมตัวอักษร Utkal เข้าไปในเอกสาร Odia) แต่งานวิจัยนี้ยังไม่สามารถแปลตัวอักษรภาษาอื่นที่ปรากฏอยู่บนเอกสารตัวอักษร Odia ได้

R. K. Bawa and G. K. Sethi [7] ได้นำเสนอการใช้เทคนิค Binarization สำหรับการสกัดตัวอักษรเทวนาครีออกจากภาพ ด้วยการใช้เทคนิค Edge Detection และ Morphological Dilation ในขั้นตอนการทำ Pre-Processing โดยเริ่มจากการนำเข้าภาพถ่ายด้วยกล้องถ่ายรูป จากนั้นจะแปลงภาพให้เป็น Gray Scale Image ในขั้นตอนนี้จะใช้ Low-Pass Wiener Filter เพื่อปรับแต่งภาพ จากนั้นใช้ Canny Edge ทำ Edge Detection เพื่อหาขอบภาพ แล้วใช้ Morphological Dilation ในการเติม Pixel บริเวณขอบของภาพตัวอักษรที่จะทำการ Capture ออกมา จากนั้นจะทำการวิเคราะห์หาขอบภาพ (Edge Analysis) ของตัวอักษร ถ้ารูปร่างของข้อมูลที่ได้ไม่มีรูปร่างคล้ายกับตัวอักษรเทวนาครี จะทำการกำจัดออกไป แล้วทำการแปลงภาพให้เป็น Binary Image หลังจากนั้นขั้นตอนการทำ Post-Processing จะใช้ Shirk and Swell Filter เข้ามาช่วย Shirk Filter จะใช้สำหรับ Reduce Noise จาก Background ภาพ ส่วน Swell Filter จะใช้ในการเติม Pixel ใน Foreground จากการทดลองดังกล่าวมีการใช้เทคนิค Canny Edge Detection และ Morphological Dilation เพื่อให้รูปร่างตัวอักษรมีความชัดเจนขึ้น เพื่อให้สกัดออกมาเป็นข้อความได้ง่ายขึ้น และได้นำเทคนิค Shirk และ Swell Filter เข้ามาช่วยแบ่งแยกส่วนที่ไม่ใช่ตัวอักษรกับตัวอักษรออกจากกัน ในงานวิจัยนี้จะไม่มีการบอกค่าความแม่นยำ เนื่องจากจะเป็นการเปรียบเทียบวิธีการด้วยเทคนิคต่างๆ ในขั้นตอน Pre-Processing เท่านั้น

อุษณีย์ สังฆธรรม [8] ได้ทำการศึกษาการรู้จำตัวอักษรพิมพ์ไทยโดยใช้วิธีคอนดิชันนัลแรนดอมฟิวส์และระยะเซนทรอยด์แบบลำดับชั้น ซึ่งเน้นการสกัดคุณลักษณะเด่นและการจำแนกตัวอักษร การสกัดคุณลักษณะเด่นจะใช้เทคนิคการกระจายทิศทางสำหรับรูปภาพตัวอักษร เพื่อวิเคราะห์ความแตกต่างของตัวอักษรแต่ละตัว ในส่วนของการจำแนกตัวอักษรจะใช้วิธีคอนดิชันนัลแรนดอมฟิวส์ ในการเลือกคุณสมบัติที่ใช้ในการจำแนกตัวอักษรภายในแต่ละกลุ่ม และใช้วิธีระยะเซนทรอยด์แบบลำดับชั้น ทำการเอาตัวอักษรที่แยกกลุ่มแล้วไปทำนายเปรียบเทียบกับข้อมูลทดสอบจากการทดสอบพบว่าได้มีความถูกต้อง 96.96%

กลุ่มที่ 2 ใช้อัลกอริทึมและเทคนิคที่กระบวนการ Processing และกลุ่มที่ 3 ใช้อัลกอริทึมและเทคนิคที่กระบวนการ Post-Processing

ประสิทธิ์ บุญเอนก และอาทิตย์ ศรีแก้ว [9] ได้นำเสนอการจดจำลายมือเขียนตัวอักษรไทยด้วยแผนผังคุณลักษณะจัดการตัวเอง โดยระบบดังกล่าวเป็นการผสมผสานระหว่างเทคนิคเครือข่ายประสาทเทียมแบบการเรียนรู้เวกเตอร์คอนโตนไทเซชันและเทคนิคเครือข่ายไปข้างหน้า ซึ่งเป็นระบบ

การจดจำลายมือเขียนตัวอักษรไทยที่สามารถจดจำ ลายมือเขียนภาษาไทยที่มีลักษณะรูปร่างผิดเพี้ยนไปจากเดิมได้ และยังสามารถเรียนรู้ลักษณะลายมือเขียนใหม่ ๆ ที่เพิ่มเติมได้ จากการใช้เทคนิคดังกล่าวพบว่ามีความถูกต้องแม่นยำอยู่ที่ 91.09%

S. P. Godara and P. S. Patwal [10] ได้นำเสนอการสกัดตัวอักษรละตินจากภาพเอกสารที่มีอักษรเทวนาครีและอักษรละตินปนกันอยู่ งานวิจัยนี้เน้นขั้นตอนการทำ Pre-Processing ในเรื่องของปรับปรุงคุณภาพของภาพที่จะนำมาแปลข้อมูล โดยใช้เทคนิค Binarization, Dilation, Erosion และ Skew Correction และใช้อัลกอริทึม Horization และ Vertical Projection Profile ทำการสกัดเอาตัวอักษรที่อยู่ในรูปภาพออกมา ผลการทดลองพบจะแบ่งออกเป็น 4 กรณี ได้แก่ กรณีที่ 1 เอกสารที่เป็นตัวอักษรละตินอย่างเดียวจะมีค่าความแม่นยำในการแปลข้อมูลอยู่ที่ 99.35% กรณีที่ 2 เอกสารที่มีตัวอักษรละตินส่วนใหญ่และมีตัวอักษรเทวนาครีปะปนอยู่ในจะค่าความแม่นยำ 0.65% กรณีที่ 3 เอกสารที่เป็นตัวอักษรเทวนาครีอย่างเดียวจะมีค่าความแม่นยำอยู่ที่ 99.80% กรณีที่ 4 เอกสารที่มีตัวอักษรเทวนาครีเป็นส่วนใหญ่และมีตัวอักษรละตินปะปนอยู่ในจะค่าความแม่นยำอยู่ที่ 0.65% จากผลการทดลองพบว่ามีข้อจำกัดในกรณีที่เอกสารมี 2 ภาษาปะปนกันจะทำให้ความแม่นยำในการแปลภาษาที่ไม่ใช่ภาษาหลัก มีค่าความแม่นยำน้อยมาก

ศุภรัตน์ อภิวงค์โสภณ [11] ได้ทำการศึกษาการรู้จำข้อมูลลายมือจากการกรอกแบบฟอร์ม โดยใช้ทำการแบ่งระดับข้อความและหาขอบภาพของแต่ละตัวอักษร จากนั้นนำมาแบ่งครึ่งตัวอักษร แล้วทำการเปรียบเทียบกับค่ารหัสลูกโซ่ 8 ทิศ (Octet Chain Code) ว่าเป็นตัวอักษรอะไร แล้วนำข้อมูลที่ได้มาจัดเก็บในระบบฐานข้อมูล ถ้ามีข้อมูลที่เป็นค่าประเภทเดียวกันที่เคยรู้จำและมีในฐานข้อมูลอยู่แล้ว ระบบจะทำการกรอกแบบฟอร์มให้โดยอัตโนมัติ จากการทดลองดังกล่าวพบว่ามีค่าความแม่นยำอยู่ที่ 91% โดยข้อผิดพลาดที่เกิดขึ้นจะเกิดจากลายมือของผู้ที่กรอกแบบฟอร์มเขียนตัวอักษรสองตัวลากเส้นติดกัน

จากงานวิจัยข้างต้นสามารถสรุปเป็นตาราง ดังตารางที่ 2.1

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
1	การแยกภาพตัวอักษรลายมือเขียนภาษาไทยแบบอัตโนมัติ โดยใช้การวิเคราะห์ถดถอยเชิงเส้นและการเรียนรู้ด้วยต้นไม้ตัดสินใจ	Automatic Thai Handwritten Characters Segmentation Based on Linear Regression Analysis and Decision Tree Learning	วิเชษฐ์รจน์ เอี่ยมสำอางค์ และรัชฎา คงคะจันทร์ [1]	Thai Journal of Science and Technology ปีที่ 2 ฉบับที่ 2 พฤษภาคม-สิงหาคม 2556	OCR: ลายมือ ภาษา: ไทย วัตถุประสงค์ : คัดแยกตัวอักษรที่ติดกัน	การรับภาพเอกสารลายมือเขียนมาคัดแยกให้เป็นตัวอักษรเดี่ยว และตัวอักษรติดกัน วิเคราะห์ตัวอักษรด้วยคุณลักษณะต่างๆ ของตัวอักษรไทย เพื่อแยกตัวอักษรที่ติดกันในแนวนอนและแนวตั้ง ใช้เส้นการวิเคราะห์การถดถอยสำหรับคัดแบ่งระดับพยัญชนะกับสระ ขั้นตอนนี้ทำการตัดแบ่งตัวอักษรก่อนการรู้จำตัวอักษรตามหลักของการรู้จำตัวอักษรไทย การใช้ต้นไม้ตัดสินใจ	90.44%	✓		
2	การจดจำการเขียนลายมือตัวเลขไทยด้วยโครงข่ายการส่งข้อมูลแบบไม่ย้อนกลับ	Classification of Thai Number Handwriting by Using Feed-Forward Neural Network	พาฝัน ดวงไพศาล และคณะ [2]	Joint conference on ACTIS & NCOPA 2015, Jan 30-31, Nakorn Phanom, Thailand. ISSN: 1906-9006	OCR: ลายมือ ภาษา: ไทย วัตถุประสงค์ภาพ: 1. รู้จำตัวเลขภาษาไทยที่เขียนด้วยลายมือที่แตกต่างไปจากเดิมได้ 2. เรียนรู้ลักษณะของลายมือใหม่ๆ	1. โครงข่ายการส่งข้อมูลแบบไม่ย้อนกลับ (Feed-Forward Neural Network) 2. K-Nearest Neighbor (KNN) 3. การเลือกสุ่มข้อมูลแบบ K-Fold Cross Validation	87.86%	✓	✓	

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา (ต่อ)

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
3	การรู้จำตัวเลขอารบิกจากลายมือสำหรับระบุรหัสผู้สอบบนกระดาษคำตอบแบบปรนัย	Handwritten Digits OCR for Identifying Examinee Number on Objective Test Answer Sheet	ธีรเดช ราชไพญ์ และ นุชนาฏ สัตยาภิ [3]	Naresuan University Journal: Science and Technology (NUJUST), [S.I.], v. 22, n. 2, p. 94-105, dec. 2014. ISSN 2539-553X	OCR: ลายมือ ภาษา: เลขอารบิก วัตถุประสงค์: 1. รู้จำตัวเลขของช่องรหัสผู้สอบในกระดาษคำตอบ	1. เทคนิคโครงข่ายประสาทเทียม 4 แกน ได้แก่ แนวแถว (Horizontal) แนวคอลัมน์ (Vertical) แนวทแยงซ้าย (Left Diagonal) แนวทแยงขวา (Right Diagonal) 2. การรู้จำตัวอักษร (Template Matching, Feature Extraction, Structural Analysis, Neural Network) 3. Algorithm แบบ Flood Fill เป็นตัวที่ใช้ในการหาพิกเซลสีที่ต้องการกับสีที่แตกต่าง	77%	✓		

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา (ต่อ)

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
4	การสกัดคุณลักษณะเด่นแบบผสมเพื่อรู้จำตัวเขียนอักษรธรรมอีสานด้วยโครงข่ายประสาทเทียม	A Hybrid Features Extraction for Isarn Dharma Handwritten Character Recognition Using Neural Network	อมรเทพ วิจิตรเจริญ และ พุทธดี ศรีแสงตระกูล [4]	2012 4th Conference on Knowledge and Smart Technology (KST)	OCR: ลายมือ ภาษา: อักษรอีสาน วัตถุประสงค์: รู้จำตัวเขียนอักษรธรรมอีสาน	โครงข่ายประสาทเทียมแบบแพร่ย้อนกลับ เป็นอัลกอริทึมการเรียนรู้แบบมีผู้ฝึกสอน เพื่อปรับปรุงค่าน้ำหนักของโครงข่ายให้เหมาะสมกับโครงข่ายนั้นๆ โดยการคำนวณหาความผิดพลาด (Error) จากผลลัพธ์ (Output) เทคนิคการรู้จำ 1. Pre-processing (Binarization, Edge-Detection, Neural Network) 2. Feature Extraction (อัลกอริทึมที่ 1 ความหนาแน่นของจุดพิกเซลด้วยวิธีแบ่งโซน อัลกอริทึมที่ 2 อีสโตแกรมโปรเจกชันในแนวตั้งและแนวนอน)	1.วิธีการแบ่งโซน 62.22% 2. วิธีฮิสโตแกรมโปรเจกชัน 52.24% 3. วิธีผสมระหว่างการแบ่งโซนและฮิสโตแกรมโปรเจกชันแนวนอน	✓	✓	✓

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา (ต่อ)

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
						3. การรู้จำตัวอักษร มีการใช้เคโฟลด์ครอสวาเลชัน (K-Fold Cross Validation) ในการแบ่งข้อมูลที่ใช้ในการทดสอบออกเป็นกลุ่มจำนวน K กลุ่ม (K-Folds) เพื่อให้ข้อมูลทุกตัวมีโอกาสเป็นชุดทดสอบและชุดสอนเพื่อป้องกันปัญหาการจากการเลือกข้อมูลที่ดีและง่ายมาเป็นข้อมูลชุดทดสอบ				
5		OCR For Handwritten Marathi Script	V. S. Tapkir and S. D. Shelke [5]	International Journal of Scientific & Engineering Research Volume 3, Issue 8, August-2012 1 ISSN 2229-5518	OCR: ลายมือ ภาษา: Marathi วัตถุประสงค์: สกัดตัวอักษร (Features Extraction) ออกมาจากรูปโดยใช้เทคนิค Projection Histogram และ Zoning Density features	1. ใช้ Algorithm Projection profiles 2. Features Extraction 3. Zoning Feature 4. Projection Histogram Features 5. Euclidean Distance Classifiers	Line Segmentation = 100% Character Segmentation = 98%	✓		

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา (ต่อ)

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
6		Odia Characters Recognition by Training Tesseract OCR Engine	M. Nayak and A. K. Nayak [6]	International Journal of Computer Applications (0975 – 8887) International Conference in Distributed Computing & Internet Technology (ICDCIT-2014)	OCR: ลายมือ ภาษา: Odia วัดประสิทธิภาพ: ใช้ Tesseract OCR Engine เพื่อรู้จำตัวอักษร	ใช้ Training Procedure สำหรับกำหนดการทำ Train Data โดยมีขั้นตอนดังนี้ 1. Prepare Training Data Image 2. Prepare Box File 3. Training the box file 4. Compute the Character Set File 5. Prepare Font_Properties File 6. Clustering 7. Prepare Dictionary 8. Last Step (ขั้นตอนนี้ Tesseract จะรู้จำตัวอักษรที่กำหนดไว้แล้ว)	1. 100% สำหรับตัวอักษร Odia ขนาด 10pt 12pt 16pt 2. 98% สำหรับตัวอักษร Utkal (มีการทดลองผสม Font Type Utkal เข้าไปในเอกสาร Odia)	✓		
7		A Binarization Technique For Extraction Of Devanagari Text From Camera Based Images	R. K. Bawa and G. K. Sethi [7]	Signal & Image Processing: An International Journal (SIPIJ) Vol.5, No.2, April 2014	OCR: ตัวพิมพ์ ภาษา: Devanagari วัดประสิทธิภาพ: ใช้เทคนิค Binarization ในการสกัดตัวอักษรเทวนาครีออกจากภาพ	1. Edge Detection 2. Morphological Dilation 3. Wiener Filter (Low-Pass) 4. Canny Edge 5. Shirk & Swell Filter	N/A	✓		

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา (ต่อ)

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
8	การรู้จำตัวอักษรพิมพ์ไทยโดยใช้วิธีคอนดิชันนัลแรนดอมฟิวส์และระยะเซนทรอยด์แบบลำดับชั้น		อุษณีย์ สังฆธรรม [8]	2556 การรู้จำตัวอักษรพิมพ์ไทยโดยใช้วิธีคอนดิชันนัลแรนดอมฟิวส์และระยะเซนทรอยด์แบบลำดับชั้น วิทยานิพนธ์ปริญญา มหาบัณฑิต สถาบัน บัณฑิตพัฒนบริหาร ศาสตร์	OCR: ตัวพิมพ์ ภาษา: ไทย วัตถุประสงค์: สกัดคุณลักษณะเด่น (Features Extraction) การจำแนกตัวอักษร (Classification) เทคนิคการกระจายทิศทาง	1. ใช้วิธีคอนดิชันนัลแรนดอมฟิวส์ ในการเลือกคุณสมบัติที่ใช้ในการจำแนกตัวอักษรภายในแต่ละกลุ่ม 2. ใช้วิธีระยะเซนทรอยด์แบบลำดับชั้น ทำการเอาตัวอักษรที่แยกกลุ่มแล้วไปทำนายเปรียบเทียบกับข้อมูลทดสอบ	96.96%	✓	✓	✓
9	การจดจำลายมือเขียนตัวอักษรไทยด้วยแผนผังคุณลักษณะจัดการตัวเอง	THAI HAND-WRITING CHARACTER RECOGNITION USING SELFORGANIZING FEATURE MAP	ประสิทธิ์ บุญเอนก และ อาทิตย์ ศรีแก้ว [9]	2551 การจดจำลายมือเขียนตัวอักษรไทยด้วยแผนผังคุณลักษณะจัดการตัวเอง วิทยานิพนธ์ ปริญญา มหาบัณฑิต มหาวิทยาลัยเทคโนโลยีสุรนารี	OCR: ลายมือ ภาษา: ไทย วัตถุประสงค์: จดจำลายมือเขียนตัวอักษรไทยที่สามารถจดจำลายมือเขียนภาษาไทยที่มีลักษณะรูปร่างผิดเพี้ยนไป	เครือข่ายประสาทเทียมแบบการเรียนรู้เวกเตอร์คอนโเทชัน เครือข่ายไปข้างหน้า	91.09%		✓	

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา (ต่อ)

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
10		Latin Script Detection and Removal from Devanagari Document Image for OCR	S. P. Godara and P. S. Patwal [10]	International Journal of Computer & Organization Trends – Volume 6 Number 1 – Mar 2014	OCR: ลายมือ ภาษา: Latin Devanagari วัตถุประสงค์: เน้นการทำ Pre-Processing (Binarization, Dilation and Erosion, Skew Correction) ก่อนที่จะทำการสกัดเอาตัวอักษรที่อยู่ในรูปภาพออกมา	ใช้ Algorithm Horization & Vertical Projection Profile	เอกสาร Latin Script Latin & Latin = 99.35% Latin & Devanagari = 0.65% เอกสาร Devanagari Script Devanagari & Devanagari = 99.80% Devanagari & Latin = 0.65%		✓	✓

ตารางที่ 2.1 เปรียบเทียบงานวิจัยที่ได้ทำการศึกษา (ต่อ)

ลำดับ	เรื่อง (ไทย)	เรื่อง (อังกฤษ)	ผู้เขียน	Conference or Journal	วัตถุประสงค์ของงานวิจัย	เทคนิคหรือหลักการที่ใช้	ความแม่นยำ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3
11	การรู้จำตัวอักษรภาษาไทยที่เขียนด้วยลายมือในแบบฟอร์ม	Thai Character Recognition of Handwritten in Forms	ศุภรัตน์ อภิวงค์โสภณ [11]		OCR: ลายมือ ภาษา: ไทย วัตถุประสงค์: จำข้อมูลลายมือจากการกรอกแบบฟอร์ม โดยใช้ทำการแบ่งระดับข้อความ หาขอบภาพของแต่ละตัวอักษรเพื่อนำมาแบ่งครึ่งตัวอักษรแล้วมาเปรียบเทียบกับค่ารหัสลูกโซ่ 8 ทิศ	1. เทคนิครหัสลูกโซ่ 8 ทิศ (Octet Chain Code) 2. เทคนิค Divide and Conquer	91%		✓	✓

สิ่งที่ได้รับจากการศึกษางานวิจัยที่เกี่ยวข้อง สามารถนำสิ่งที่ได้ศึกษาทั้งข้อดีและข้อจำกัด มาออกแบบการทดลองของงานที่นำเสนอ โดยนำข้อมูลเรื่องต่างๆ ดังนี้

1. การเตรียมข้อมูล Training Data เพื่อใช้สร้างตัวอักษรต้นแบบมาเตรียมไว้สำหรับการทำ Pattern Matching
2. การทำ Pre-Processing ของฝั่ง Testing Data ที่มีขั้นตอนการทำ Binarization, Features Extraction และ Segmentation
3. การทำ Processing โดยวิธี Pattern Matching ระหว่างข้อมูล Training Data กับ Testing Data
4. รู้กระบวนการทำงานและวิธีการใช้เครื่องมือ Tesseract OCR เข้ามาช่วยในกระบวนการทำงานให้มีประสิทธิภาพและความแม่นยำมากยิ่งขึ้น
5. มีการออกแบบการทดลองที่มี Training Data รายการคำศัพท์ที่ใช้ประจำเข้ามาอยู่ในข้อมูลด้วย เนื่องด้วยงานที่นำเสนอนี้เป็นงานที่ความเฉพาะของข้อมูลที่น่ามาทดลองและเป็นข้อมูลลายมือตัวอักษรของผู้ปฏิบัติงานที่ทำการเขียนอยู่เป็นประจำและมีคำศัพท์ซ้ำๆ ในแต่ละกล่อง แต่งานวิจัยส่วนใหญ่จะเน้นการสร้างต้นแบบจากแบบอักษรที่มีเพื่อให้แปลงออกมาเป็นตัวอักษรแต่ละตัว จึงทำให้งานที่นำเสนอตรงจุดนี้เป็นข้อสมมติฐานที่น่ามาออกแบบการทดลองเพื่อ พิสูจน์ให้เห็นว่างานที่เสนอนั้นสามารถทำได้จริง

2.2 ทฤษฎีที่เกี่ยวข้องกับงานวิจัย

2.2.1 ชนิดของภาพ

ภาพถ่าย [12] หมายถึง รูปภาพที่ได้จากการถ่ายภาพจากกล้องถ่ายภาพ ไม่ว่าจะเป็นการถ่ายภาพด้วยกล้องฟิล์มหรือถ่ายภาพด้วยกล้องดิจิทัล และสามารถนำภาพถ่ายมาพิมพ์หรือนำมาใช้งานในรูปแบบไฟล์ โดยภาพถ่ายด้วยกล้องฟิล์มสามารถนำมาแปลงเป็นไฟล์ภาพดิจิทัลได้โดยใช้เครื่องสแกนเนอร์ เพื่อที่จะสามารถนำไฟล์ภาพที่ได้ไปทำการประมวลผลภาพในรูปแบบต่างๆ

ไฟล์ภาพถ่ายดิจิทัลจะถูกนำมาใช้ในการคำนวณด้วยการสร้างแบบจำลองทางเรขาคณิตหรือสูตรทางคณิตศาสตร์ ทำให้สามารถนำมาสังเคราะห์ให้มีความเหมาะสมยิ่งขึ้น โดยทั่วไปภาพดิจิทัลสามารถแบ่งประเภทได้ ดังนี้

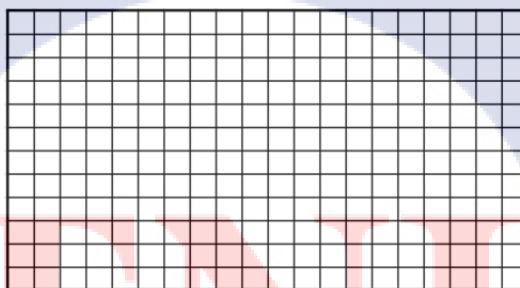
ภาพสี (Color Image) โดยทั่วไปภาพดิจิทัล จะเก็บค่าของความเข้มเป็นปริมาณเวกเตอร์ โดยแต่ละจุดของรูปภาพจะมีการแทนที่ด้วยค่ามากกว่าหนึ่งค่า เช่น แบบจำลองสี RGB จะมีค่าความเข้มของสีแดง (R:Red) สีเขียว (G:Green) และสีน้ำเงิน (B:Blue) ผสมกันในอัตราส่วนที่แตกต่างกันเพื่อประกอบกันเป็นรูปภาพ 1 รูป

ภาพสีเทา (Grayscale Image) เป็นภาพที่มีค่าความเข้มเป็นปริมาณสเกลาร์ บางครั้งเรียกว่า “ภาพสีเดียว” ภาพโทนสีเทา เกิดจากการกำหนดค่าให้กับความเข้มแสงของแต่ละจุดภาพโดย สีขาวจะถูกแทนด้วย “255” และสี ดำจะถูกแทนด้วยค่า “0” ในขณะที่ ระดับของโทนสีจะเข้มข้นเรื่อยๆ จากสีขาวจนถึงสีดำ

ภาพขาวดำ (Binary Image) เป็นภาพดิจิทัลที่ แต่ละจุดภาพสามารถมีค่าที่เป็นไปได้เพียงแค่ 2 ค่า คือ ค่าสีขาวแทนค่า ด้วย “1” และค่าสีดำแทนค่าด้วย “0” บางครั้งเรียกว่า “ภาพสองระดับ” (Bi-level) ภาพขาวดำ (Black-and-white) หรือภาพโมโนโครม (Monochrome)

2.2.2 พิกเซล (Pixel)

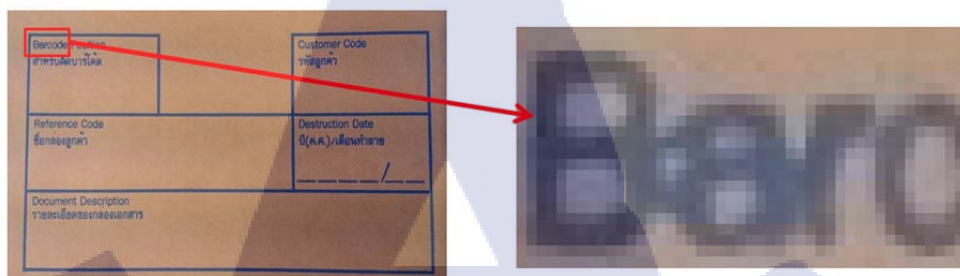
พิกเซล (Pixel) [12] หรือจุดภาพ เป็นคำที่มีความหมายมาจาก พิกเจอร์ (Picture) หรือรูปภาพและคำว่า อีเลเมนต์ (Element) หมายถึง ส่วนหรือองค์ประกอบ ปกติแล้วภาพ 1 ภาพ จะประกอบด้วยพิกเซลเป็นจำนวนมากซึ่งมีความหนาแน่นของพิกเซลหรือที่เรียกว่าความละเอียดแตกต่างกันไป พิกเซลจะใช้สำหรับการบอกคุณสมบัติและความละเอียดของภาพ อุปกรณ์แสดงผลภาพ โดยภาพที่มีจำนวนพิกเซลมากจะมีความละเอียดสูง จะนิยมบอกจำนวนค่าพิกเซลในลักษณะแนวนอนคูณแนวตั้ง ($X \times Y$) เช่น ขนาดหน้าจอแสดงผลภาพ 1024×768 พิกเซล จะมีจำนวนพิกเซลเท่ากับ 786,432 พิกเซล เป็นต้น ดังรูปที่ 2.1



รูปที่ 2.1 พิกเซลขนาด 1x1

โครงสร้างทั่วไปของรูปภาพดิจิทัล ภาพกราฟิกจากคอมพิวเตอร์หรือภาพบิตแมป (Bitmap) จะแสดงเป็นตารางสี่เหลี่ยมจัดรัสของพิกเซลหรือจุดสีหลายๆ จุดวางต่อเนื่องกัน และมีความเข้มของแต่ละพิกเซลจะแตกต่างกันตามระบบของภาพสี ซึ่งจำนวนสีที่แตกต่างกันจะถูกแทนด้วยพิกเซล ซึ่งจะขึ้นอยู่กับจำนวนบิตต่อพิกเซล โดยภาพ 1 บิตต่อพิกเซล หมายถึง พิกเซลแต่ละจุดมีค่าเท่ากับ 1 บิต พิกเซลสามารถเป็นได้ทั้งค่า 0 เป็นค่าที่แสดงถึงพิกเซลนั้นมีความเข้มแสงต่ำสุด

หรือที่เรียกว่าด้านมืด หรือค่า 1 เป็นค่าแสดงถึงพิกเซลนั้นมีความเข้มแสงสูงสุดหรือเป็นด้านสว่าง ดังรูปที่ 2.2



รูปที่ 2.2 พิกเซลที่ปรากฏอยู่ในรูปภาพ

เมื่อภาพมีการเพิ่มจำนวนสีขึ้นเรื่อยๆ แต่ละบิตจะเพิ่มค่าเป็น 2 เท่า เช่น ภาพ 1 บิตต่อพิกเซล จะมี 2 สี (ขาวดำ) หรือภาพที่มี 2 บิต จะมี 4 สีให้เลือก เป็นต้น ตัวอย่างสีและการคำนวณของจำนวนสีของภาพ 8 บิตต่อพิกเซล มีค่าเท่ากับ 2^8 หรือ 256 สี ซึ่งค่าจะเริ่มที่ 0 ไปจนถึง 255 ค่า 0 จะเป็นพิกเซลที่มีความเข้มแสงต่ำสุด (ด้านมืด) และค่า 255 เป็นพิกเซลที่มีความเข้มแสงสูงสุด (ด้านสว่าง)

2.2.3 การประมวลผลภาพ (Digital Image Processing)

การประมวลผลภาพ [12] จะเป็นการจัดการกับพิกเซลของภาพอินพุต เพื่อให้ได้ภาพเอาต์พุตก่อนที่จะนำภาพที่ได้ไปปรับแต่งให้สามารถใช้งานได้ในขั้นตอนต่อไป วิธีการที่จะทำการประมวลผลภาพจะมีหลากหลายวิธี วิธีการจัดการกับพิกเซลถือเป็นกระบวนการประมวลผลขั้นต้น (Pre-Processing) อย่างหนึ่ง ที่เพิ่มหรือลดความสว่างของภาพ โดยจะทำการแปลงพิกเซลแต่ละพิกเซลของภาพอินพุตเข้ากับพิกเซลที่สอดคล้องกันของภาพเอาต์พุต ซึ่งถือว่าเป็นกระบวนการเพื่อปรับปรุงคุณภาพของภาพ เพื่อให้รายละเอียดในส่วนที่มองเห็นไม่ชัดเจนในภาพต้นฉบับ สามารถมองเห็นได้ชัดเจนขึ้น

การประมวลผลภาพที่ทำให้คอมพิวเตอร์สามารถรู้จักวัตถุในภาพได้ แบ่งได้เป็น 2 ระดับ ได้แก่ การประมวลผลภาพในระดับต่ำ (Low-Level Image Processing) และการประมวลผลภาพในระดับสูง (High-Level Image Processing)

การประมวลผลภาพในระดับต่ำ (Low-Level Image Processing) เป็นการประมวลผลเชิงตัวเลขเกือบทั้งหมด โดยการประมวลผลจะใช้ค่าพิกเซลของจุดภาพ เพื่อทำการหาตัวแปรต่างๆ มา

อธิบายข้อมูลของรูปภาพ แล้วนำตัวแปรต่างๆ เหล่านั้นไปใช้ในขั้นตอนการประมวลผลภาพในระดับสูงต่อไป ตัวอย่างของการประมวลผลภาพในระดับต่ำ เช่น การประมวลผลภาพก่อน (Pre-Processing) ด้วยการกำจัดสัญญาณรบกวน การหาขอบภาพเพื่อให้ภาพมีความคมชัดขึ้น เป็นต้น การประมวลผลภาพในระดับสูง (High-Level Image Processing) เป็นการนำผลหรือข้อมูลที่ได้การประมวลผลภาพในระดับต่ำมาประมวลผลเพื่อให้คอมพิวเตอร์สามารถรู้จักและเข้าใจภาพได้ โดยจะแสดงในรูปของสัญลักษณ์หรือสิ่งที่ปรากฏอยู่ในภาพ

2.2.4 Binarization

ไบนารีเซชัน (Binarization) [12] หรือการทำภาพขาวดำ 2 ระดับ เป็นการแปลงข้อมูลภาพสีมาทำให้เป็นสีเทาแล้วนำภาพสีเทามาทำให้เป็นภาพขาวดำที่มีค่าระดับความสว่างเพียงแค่ 0 กับ 1 เท่านั้นซึ่งเรียกว่า ภาพไบนารี โดยอาศัยจุดเทรชโฮลด์ทำการกำหนดระดับขาวหรือดำ ดังสมการที่ 2.1 เพื่อช่วยให้การวิเคราะห์ภาพและการคำนวณเป็นไปได้ง่าย

สมการที่ 2.1 สมการไบนารีเซชัน

$$g(x,y) = \begin{cases} 1 \Rightarrow f(x,y) \geq T \\ 0 \Rightarrow f(x,y) < T \end{cases}$$

โดยที่ $f(x,y)$ คือ ภาพระดับเทาของตำแหน่ง x,y

T คือ ค่าเทรชโฮลด์

$g(x,y)$ คือ ภาพขาวดำสองระดับที่ได้ของตำแหน่ง x,y

ลักษณะของภาพขาวดำสองระดับที่ได้ขึ้นกับค่าเทรชโฮลด์ ดังนั้นค่าเทรชโฮลด์มีค่าสูงภาพที่ได้จะเป็นสีดำมาก แต่ถ้าค่าเทรชโฮลด์มีค่าต่ำ ภาพที่ได้จะเป็นสีขาวมาก ซึ่งควรจะมีการกำหนดค่าเทรชโฮลด์ที่เหมาะสมจะช่วยให้ภาพขาวดำสองระดับที่มีความถูกต้องใกล้เคียงมากที่สุด โดยทั่วไปจะกำหนดค่าเทรชโฮลด์ด้วยการเลือกค่าพิกเซล 2 วิธี

วิธีที่ 1 เลือกค่าพิกเซลสูงกว่าเทรชโฮลด์ ถ้าค่าพิกเซลมากกว่าเทรชโฮลด์ (วัตถุมีความสว่างมากกว่าพื้นหลัง) พิกเซลภาพจะมีค่า 255 (สีขาว) และพิกเซลพื้นหลังมีค่า 0 (สีดำ) เรียกว่า พิกเซลสูงกว่าเทรชโฮลด์

วิธีที่ 2 เลือกค่าพิกเซลต่ำกว่าเทรชโฮลด์ ถ้าค่าพิกเซลต่ำกว่าเทรชโฮลด์ พิกเซลภาพจะมีค่า 0 (สีดำ) และพิกเซลพื้นหลังมีค่า 255 (สีขาว)

ปกติค่าพิกเซลของวัตถุจะมีค่าเท่ากับ 1 (สีขาว) และพิกเซลของพื้นหลังมีค่าเท่ากับ 0 (สีดำ) ภาพไบนารีจะถูกสร้างขึ้นโดยพิกเซลสีขาวหรือสีดำ ซึ่งขึ้นอยู่กับระดับของพิกเซล

การนำทฤษฎีเทรชโฮลด์มาใช้จะเป็นการหาค่าความเข้มข้นให้ค่าที่มีความแตกต่างระหว่างตัววัตถุและพื้นหลัง เช่น กรณีภาพกล่องเอกสารที่มีการเขียนตัวอักษรเพื่อทำดัชนี จะมีความเข้มข้นของตัวอักษรเป็น 0 (สีดำ) และความเข้มข้นของพื้นหลังเป็น 255 (สีขาว) เป็นต้น ซึ่งในทฤษฎีเทรชโฮลด์ยังสามารถแบ่งการหาค่าได้โดยใช้ โกลบอลเทรชโฮลด์ (Global Thresholding) และอะแดปทีฟเทรชโฮลด์ (Adaptive Thresholding)

โกลบอลเทรชโฮลด์ (Global Thresholding) เป็นการทำให้ค่าเทรชโฮลด์เป็นค่าคงที่หนึ่งค่า โดยใช้ฮิสโตแกรมจากรูปภาพที่นำมาทดสอบ ซึ่งค่าเทรชโฮลด์ที่ใช้ควรเลือกค่าฮิสโตแกรมที่อยู่จุดต่ำสุดที่อยู่ระหว่างจุดสูงสุด (Peaks) ดังสมการที่ 2.2

สมการที่ 2.2 โกลบอลเทรชโฮลด์

$$g(x,y) = \begin{cases} 1 & \text{if } (x,y) > T \\ 0 & \text{Otherwise} \end{cases}$$

โดยที่ $g(x,y)$ คือ ข้อมูลภาพ ณ ตำแหน่ง x,y
 T คือ ค่าเทรชโฮลด์

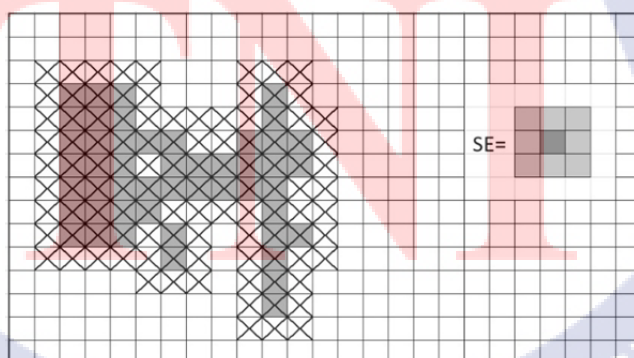
อะแดปทีฟเทรชโฮลด์ (Adaptive Thresholding) เป็นการหาค่าเทรชโฮลด์ที่มีได้หลายค่า โดยจะทำการหาค่าเทรชโฮลด์จากค่าความเข้มของสีในบริเวณกรอบสี่เหลี่ยมของภาพที่กำหนด โดยมีเงื่อนไขว่ากรอบสี่เหลี่ยมจะต้องมีขนาดที่เหมาะสม แต่ถ้าภาพไม่มีลักษณะเป็นขอบสูงชันพอ ค่าเทรชโฮลด์จะมีผลกระทบต่อตำแหน่งขอบเขตและขนาดของวัตถุจริง ทำให้การวัดขนาดและการคำนวณหาพื้นที่ที่คลาดเคลื่อนได้ จะนิยมใช้สำหรับการทำให้ภาพสีหรือภาพ Grayscale ที่มีแสงสว่างไม่สม่ำเสมอ (Nonuniform Illumination) อยู่ในรูปแบบภาพขาวดำ (Binary Image) เทคนิควิธีการหาค่าเทรชโฮลด์ทำได้โดยการสร้างหน้าต่าง (Window) ขึ้นมาขนาด $N \times N$ เช่น 3×3 โดยที่ N นั้นควรเป็นเลขคี่จากนั้นนำหน้าต่าง (Window) นี้ไปวางไว้ที่บริเวณหนึ่งของภาพแล้วนำค่า Gray Level ของทุกพิกเซลที่อยู่ในขอบเขตของหน้าต่าง (Window) มาบวกกันแล้วหารด้วยจำนวนช่องทั้งหมดของหน้าต่าง (Window) จะได้ค่าเทรชโฮลด์ที่อยู่ภายในหน้าต่าง (Window) นั้น จากนั้นทำการหาค่าเทรชโฮลด์เช่นนี้ไปเรื่อยๆ กับบริเวณที่ไม่ซ้ำกันจนกระทั่งได้มีการกำหนดค่าเทรชโฮลด์ครบในทุกๆ พิกเซล ถ้าค่า Gray Level ของพิกเซลนั้นมีค่ามากกว่าค่าเทรชโฮลด์ของพิกเซลนั้นแล้วจะกำหนดให้เป็นสีขาว แต่ถ้าค่า Gray Level ของพิกเซลนั้นมีค่าน้อยกว่าค่าเทรชโฮลด์ของพิกเซลนั้นแล้วจะ

กำหนดให้เป็นสีดำ ทำเช่นนี้จนครบทุกพิกเซลจะได้ผลลัพธ์เป็นภาพขาวดำ งานวิจัยนี้ได้เลือกใช้ Local Threshold สำหรับการปรับปรุงภาพที่มีแสงเงาติดมาในตอนถ่ายรูป โดยหลักการทำงานของ Local Threshold จะแบ่งบริเวณออกเป็นส่วนย่อยๆ เพื่อหาเทรชโฮลด์ที่เหมาะสมที่สุดสำหรับบริเวณนั้นๆ ซึ่งอาจจะแตกต่างจากบริเวณอื่นๆ ได้ การเลือกค่าเทรชโฮลด์ที่เหมาะสมจะขึ้นอยู่กับความเข้มแสงของบริเวณที่เป็นข้อมูลที่น่าสนใจและบริเวณที่เป็นฉากหลังมีความแตกต่างกันพอสมควร ซึ่งจะต้องสามารถแบ่งฉากหลังและวัตถุออกจากกันได้เป็นอย่างดี

2.2.5 Morphology (dilation erosion)

มอร์โฟโลยี (Morphology) [12] เป็นกระบวนการทางคณิตศาสตร์ตามรูปแบบหลักทฤษฎีแลตติซ (Lattice Theory) เพื่อประมวลผลภาพในคอมพิวเตอร์ บนพื้นฐานทฤษฎีเซต (Set Theory) ทฤษฎีตาข่าย โครงสร้างของเครือข่ายและฟังก์ชันแบบสุม นิยมใช้งานเพื่อประยุกต์กับภาพดิจิทัล หลักการพื้นฐานของมอร์โฟโลยีจะเป็นการพิจารณาเฉพาะจุดของรูปภาพที่เป็นข้อมูลที่ต้องการเท่านั้น ตัวอย่างเช่น มีรูปภาพตัวอักษรสีดำและพื้นหลังมีสีขาว มอร์โฟโลยีจะพิจารณาเฉพาะตัวอักษรสีดำที่เป็นส่วนของข้อมูลเท่านั้น ส่วนพื้นหลังสีขาวปรากฏอยู่บนภาพนั้น จะไม่นำ ส่วนนี้ไปพิจารณา หลักการของมอร์โฟโลยีมีรูปแบบที่สำคัญ คือ การขยายภาพ (Dilation) การกร่อนภาพ (Erosion) การเปิดภาพ (Closing) และการปิดภาพ (Opening)

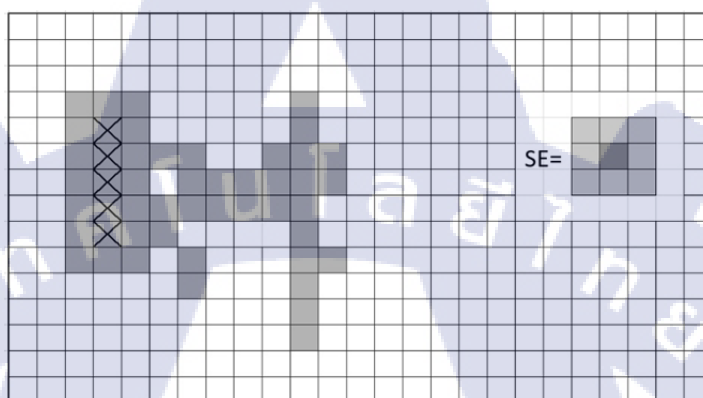
การขยายภาพ (Dilation) เป็นขั้นตอนการเพิ่มขนาดของวัตถุโดยใช้องค์ประกอบของโครงสร้าง (Structuring Element : SE) สำหรับการตรวจสอบและขยายรูปทรงที่มีอยู่ในอินพุต เช่น ภาพตัวอักษรที่สแกนออกมาพบว่าการขาดหายไปบางส่วน โดยสามารถเชื่อมต่อบริเวณที่ขาดหายไปด้วยช่องว่างที่มีขนาดเล็กกว่าขนาดขององค์ประกอบของโครงสร้าง ดังรูปที่ 2.3



รูปที่ 2.3 ตัวอย่างการเพิ่มขนาดของวัตถุด้วยวิธีการขยายภาพ (Dilation)

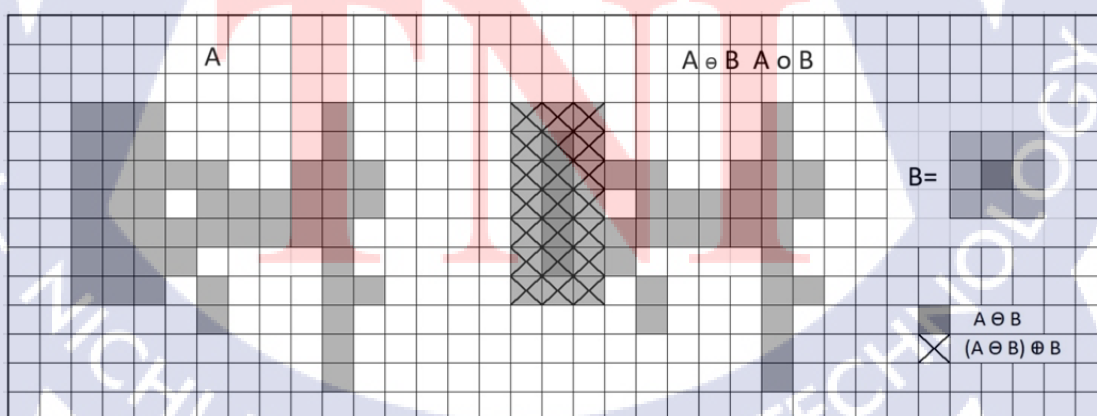
บริเวณที่เป็นเครื่องหมายกากบาทในตารางจะเป็นผลลัพธ์ที่ได้จากการนำ SE (Structuring Element) เข้ามาทำในกระบวนการ

การกร่อนภาพ (Erosion) เป็นขั้นตอนการลดขนาดของวัตถุและลบความผิดปกติเล็กๆ โดยการลบวัตถุด้วยรัศมีที่มีขนาดเล็กกว่าองค์ประกอบของโครงสร้าง จะได้ผลลัพธ์ตามที่มีเครื่องหมายกากบาทในตาราง ดังรูปที่ 2.4



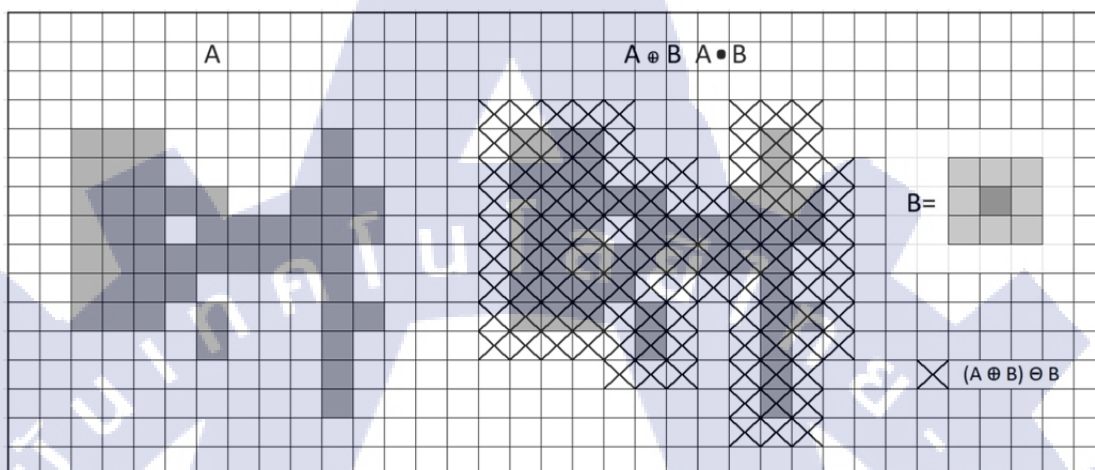
รูปที่ 2.4 ตัวอย่างการลดขนาดของวัตถุด้วยวิธีการกร่อนภาพ (Erosion)

การเปิดภาพ (Opening) เป็นการขยายภาพที่ได้จากการกร่อนภาพ โดยจะทำงานร่วมกับการปิดภาพ หลักการทำงานของการทำงานนั้นจะเป็นการกำจัดวัตถุขนาดเล็กออกจากภาพ หรือพิกเซลที่มืดของภาพ และนำไปวางไว้ในพื้นหลังของภาพ ทำให้พิกเซลของภาพถูกเปิดกว้างมากขึ้น จะได้ผลลัพธ์ตามที่มีเครื่องหมายกากบาทในตาราง ดังรูปที่ 2.5



รูปที่ 2.5 ตัวอย่างการขยายภาพที่ได้จากการกร่อนภาพด้วยวิธีการเปิดภาพ (Opening)

การปิดภาพ (Closing) เป็นการกร่อนภาพที่ได้จากการขยายภาพ โดยการปิดภาพ จะกำจัดช่องโหว่เล็กๆ ออกจากภาพ ซึ่งจะทำให้การเปลี่ยนพื้นหลังให้เป็นพื้นหน้า ซึ่งจะทำให้ภาพมีการเชื่อมต่อกันมากขึ้น และการปิดภาพจะทำให้พิกเซลของภาพถูกเชื่อมต่อกันมากขึ้น จะได้ผลลัพธ์ ตามที่มีเครื่องหมายกากบาทในตาราง ดังรูปที่ 2.6



รูปที่ 2.6 ตัวอย่างการกร่อนภาพที่ได้จากการขยายภาพด้วยวิธีการปิดภาพ (Closing)

โดยเทคนิคมอร์โฟโลยี (Morphology) จะนำไปใช้ขั้นตอนก่อนที่จะทำการรูปจำตัวอักษรเพื่อหารูปภาพตัวอักษรสีดำและพื้นหลังที่มีสีขาว

2.2.6 Canny Edge Detection Algorithm

การหาขอบภาพแคนนี่ (Canny Edge Detection) [12] เป็นการหาขอบภาพโดยใช้วิธีการหลายๆ ขั้นตอน เพื่อจับขอบเขตความกว้างของขอบภาพในภาพ โดยมีจุดมุ่งหมายเพื่อตรวจจับขอบภาพที่ถูกต้อง โดยทำเครื่องหมายที่ขอบภาพจริงในภาพที่เป็นไปได้ และได้ตำแหน่งที่ถูกต้อง ด้วยเครื่องหมายที่ขอบภาพจะใกล้เคียงที่สุดกับขอบภาพจริงในภาพ และมีการตอบสนองต่ำที่สุด ด้วยขอบภาพจะทำเครื่องหมายในครั้งเดียวและสัญญาณรบกวนในภาพจะไม่สร้างขอบภาพเทียม

หลักการของวิธีแคนนี่ [12] จะเริ่มต้นจากการปรับภาพให้เรียบ (Smoothing) ด้วยตัวกรองเกาส์เซียน (Gaussian Filter) เพื่อกำจัดสัญญาณรบกวน หลังจากนั้น จะทำการคำนวณค่าขนาด (Magnitude) และทิศทาง (Orientation) ของ Gradient โดยใช้การหาอนุพันธ์อันดับหนึ่ง ในถัดมาจึงใช้ Nonmaxima Suppression กับ Gradient Magnitude เพื่อทำให้ได้ขอบที่บางลง และ

ในขั้นตอนสุดท้ายใช้ Double Thresholding Algorithm เพื่อระบุพิกเซลที่เป็นขอบและช่วยเชื่อมต่อขอบ

2.2.7 Pattern recognition

การรู้จำแบบ [13] เป็นแขนงหนึ่งของวิทยาการคอมพิวเตอร์ มีจุดประสงค์เพื่อใช้ในการจำแนกวัตถุ (Objects) การแบ่งประเภท (Classes) ตามรูปแบบของวัตถุ โดยหาจากคุณลักษณะทั่วไป รูปแบบ ความสัมพันธ์ของสิ่งที่เราให้ความสนใจ หลักการของการรู้จำแบบ จะใช้วิธีการแยกองค์ประกอบของวัตถุ เพื่อให้สามารถคำนวณหารูปแบบและคุณลักษณะของวัตถุนั้นๆ ได้ โดยการคำนวณจะมีการใช้เทคนิคจากสาขาอื่นๆ มากมาย เช่น การประมวลผลสัญญาณ ปัญญาประดิษฐ์และสถิติ เป็นต้น

รูปแบบ (Pattern) หมายถึง คุณลักษณะ รูปทรงของวัตถุที่เราให้ความสนใจ โดยวัตถุนั้นอาจรูปแบบที่มีความเหมือนหรือมีความแตกต่างกันได้

องค์ประกอบหลักสำคัญของระบบวิธีการรู้จำแบบ มีดังนี้

1. ลักษณะเด่น (Features) เป็นข้อมูลที่ป้อนให้ตัวแยกประเภท เพื่อให้ตัวแยกประเภทสามารถทำการแยกวัตถุหรือข้อมูล ออกเป็นประเภทต่างๆ ได้ตามที่ผู้ออกแบบได้กำหนดหรือคาดหวังเอาไว้

2. ตัวแยกประเภท (Classifiers) เป็นผู้ตัดสินใจแยกกลุ่ม ของวัตถุ ตามข้อมูลลักษณะเด่น โดยทั่วไปนิยมแบ่งออกเป็น 2 ประเภทด้วยกัน ได้แก่

- การแยกกลุ่มตามประเภทที่รู้ล่วงหน้าแล้ว (Prior Knowledge) และใช้ประโยชน์จากข้อมูลนั้นในการออกแบบตัวแยกประเภท ซึ่งจะเรียกว่า ตัวแยกแบบมีผู้สอน (Supervised Classifier)

- การแยกประชากรวัตถุ ออกจากกันโดยไม่มีข้อมูลของกลุ่มการแบ่งล่วงหน้า แต่จะแบ่งโดยใช้ลักษณะที่มีร่วมกันในกลุ่มย่อยแต่ละกลุ่มของประชากร ซึ่งจะเรียกว่า ตัวแยกแบบไม่มีผู้สอน (Unsupervised Classifier)

- โดยเทคนิคนี้จะนำไปใช้ขั้นตอนที่จะทำการรู้จำรูปแบบของตัวอักษรของข้อมูล Testing Data กับ Training Data

2.2.8 Optical Character Recognition

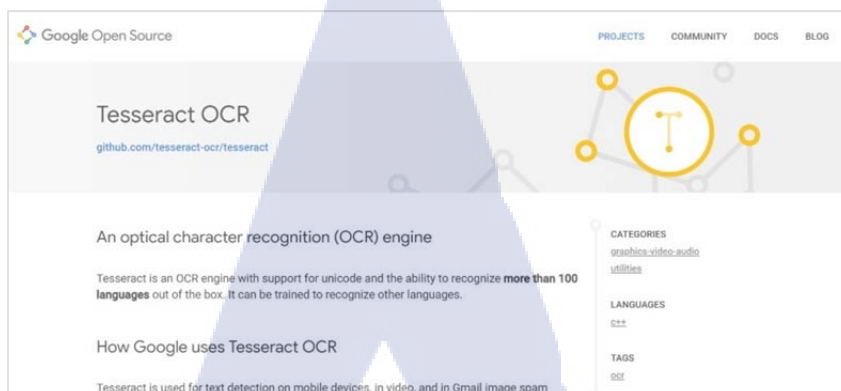
การรู้จำอักขระทางภาพ [14] หรือเรียกอย่างย่อว่า โอซีอาร์ (OCR) คือกระบวนการเพื่อแปลงรูปภาพของข้อความจากการพิมพ์หรือการเขียน ไปเป็นข้อความหรือตัวอักษรที่ผู้ใช้งานสามารถแก้ไขได้ในคอมพิวเตอร์ โดยที่ผู้ใช้งานไม่ทำการพิมพ์เอกสารเหล่านั้นให้เสียเวลาและยัง

สามารถจัดเก็บลงเครื่องคอมพิวเตอร์ได้ทันที การทำโอซีอาร์จะเริ่มกระบวนการแปลงข้อมูลในรูปแบบกระดาษที่จับต้องได้หรือข้อความที่อยู่บนพื้นผิวที่ไม่ใช่กระดาษให้อยู่ในรูปแบบของรูปภาพที่เป็นไฟล์อิเล็กทรอนิกส์ โดยใช้กล้องจากโทรศัพท์มือถือ เครื่องสแกนเนอร์ เครื่องถ่ายเอกสารแบบมัลติฟังก์ชัน กล้องดิจิทัล จากนั้นจะนำไฟล์ภาพที่ได้มาเข้าสู่กระบวนการของซอฟต์แวร์โอซีอาร์ โดยขั้นแรกจะทำการตรวจจับภาพว่ามีข้อมูลอยู่ในบรรทัดใดบ้าง ขั้นที่ 2 จะทำการตรวจสอบว่าบรรทัดที่มีข้อมูลมีข้อมูลใดที่สามารถแยกเป็นคำได้ ขั้นที่ 3 จะทำการแบ่งคำที่ได้ให้เป็นแต่ละตัวอักษรเพื่อให้สามารถนำตัวอักษรที่ได้ไปทำการเปรียบเทียบกับข้อมูลต้นแบบที่มีในระบบ เพื่อแปลงตัวอักษรในไฟล์รูปภาพเหล่านั้นมาเป็นตัวอักษรที่สามารถใช้งานได้บนคอมพิวเตอร์

เทคโนโลยีโอซีอาร์ ได้มีการพัฒนาให้มีความสามารถที่แปลงข้อมูลที่มาจากเครื่องพิมพ์ได้แม่นยำมากยิ่งขึ้น แต่ยังมีข้อจำกัดสำหรับข้อมูลที่เกิดจากการเขียนด้วยลายมือที่ยังไม่สามารถแปลงข้อมูลได้แม่นยำ เนื่องจากลายมือและวิธีการเขียนของแต่ละคนไม่เหมือนกัน บางคนเขียนแบบแยกแต่ละตัวที่สามารถอ่านได้ง่าย บางคนอาจจะเขียนตัวอักษรแบบเชื่อมต่อกันหรือเขียนต่อเนื่องกัน จึงทำให้ต้องมีการสอนตัวอย่างอักษรแต่ละตัวเพื่อให้ซอฟต์แวร์สามารถรู้จำกลุ่มตัวอักษรที่ทำการสอนไว้ เพื่อให้ซอฟต์แวร์สามารถอ่านข้อมูลลายมือที่คุ้นเคยแล้วแปลงออกมาเป็นตัวอักษรที่ถูกต้องได้ แต่การสอนแบบนี้มีข้อจำกัดในเรื่องของกลุ่มลายมือ เพราะถ้าเป็นลายมือที่แปลกจากตัวต้นแบบ จะทำให้การแปลงข้อมูลอาจถูกในบางกลุ่มลายมือที่มีลายมือใกล้เคียงกันหรือผิดไปเลยถ้าลายมือเขียนแตกต่างออกไป จึงทำให้ต้องมีการสอนตัวอย่างเพื่อให้สามารถประมวลผลแม่นยำมากยิ่งขึ้น โดยเทคนิคนี้จะนำไปใช้การรู้จำตอนที่ทำ Testing Data

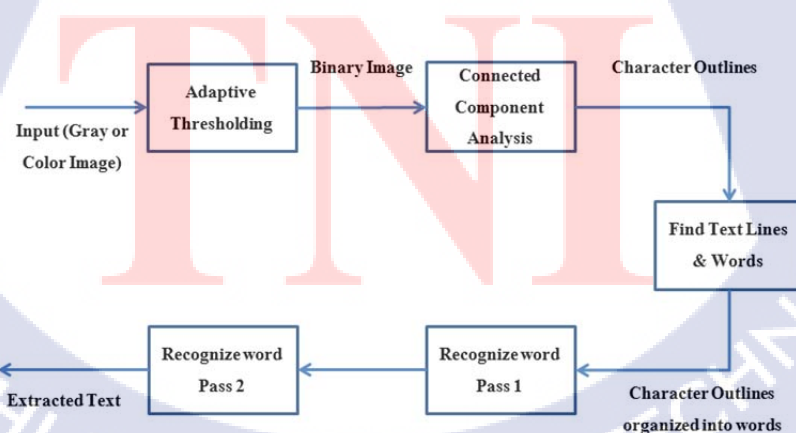
2.2.9 Tesseract OCR

Tesseract OCR [15] เป็น Google Open Source OCR Engine ที่ถูกพัฒนาขึ้นครั้งแรกระหว่างปี 1984-1994 ที่บริษัท เอช พี (HP) จากนั้นได้มีการส่งไปทดสอบประสิทธิภาพด้านความถูกต้องแม่นยำ (Accuracy) ที่ UNLV (University of Nevada, Las Vegas) จากนั้นจึงได้ทำเป็นโครงการร่วมมือระหว่าง HP Labs Bristol กับ HP's Scanner Division ที่เมือง Colorado ในปี 2005 Tesseract OCR ได้ถูกเผยแพร่เป็น Open Source ปัจจุบัน Tesseract เป็น Open-Source OCR engine ที่อยู่ภายใต้การดูแลของบริษัท กูเกิล (Google) ที่เปิดโอกาสให้ผู้ที่มีความสนใจในด้านการทำ OCR เข้ามาดาวน์โหลดและนำไปใช้เพื่อพัฒนาต่อยอดได้ ดังรูปที่ 2.7



รูปที่ 2.7 หน้าเว็บไซต์ของ Tesseract OCR

รูปแบบการทำงานของ Tesseract จะเริ่มจากการนำภาพที่เป็น Gray Scale หรือภาพสี เข้ามาปรับภาพให้เป็นภาพไบนารี (Binary Image) ด้วยการใช้เทคนิคอะแดปทีฟเทรชโฮลดิ้ง (Adaptive Thresholding) จากนั้นจะทำการวิเคราะห์การฟิกเซลที่เชื่อมต่อกันของตัวอักษร (Connected Component Analysis) เพื่อให้ได้แบบร่างของตัวอักษรออกมา จากนั้นจะทำการหาบรรทัดของตัวอักษรและหาคำ ซึ่งเมื่อได้เป็นคำศัพท์ออกมาแล้วจะทำการรู้จำคำศัพท์นั้นๆ 2 รอบ ดังนี้ รอบที่ 1 Recognize Word Pass 1 Tesseract ทำการรู้จำคำศัพท์ โดยส่งคำศัพท์ที่ได้ไปยัง Adaptive Classifier ที่เป็น Training Data เพื่อจะช่วยให้การประมวลผลการรู้จำคำศัพท์ได้ถูกต้องแม่นยำ แต่จะมีคำศัพท์จำนวนหนึ่งที่ไม่สามารถทำการจดจำได้ในรอบที่ 1 Tesseract จะส่งคำศัพท์เหล่านั้นไปประมวลผลซ้ำรอบที่ 2 (Recognize Word Pass 2) เพื่อให้สามารถสกัดตัวอักษร (Extracted Text) ออกมาได้ตามภาพตั้งต้น ดังรูปที่ 2.8



รูปที่ 2.8 ขั้นตอนการทำงานของ Tesseract OCR Engine

โดย Tesseract OCR Engine จะอยู่ในฟังก์ชัน OCR Trainer ของ MATLAB จะนำไปใช้กับการทำ Training Data กับ Testing Data

2.2.10 กระบวนการ OCR ในโปรแกรม MATLAB

กระบวนการ OCR ใน MATLAB [16] จะถูกแบ่งออกเป็น 2 ส่วน ได้แก่

ส่วนที่ 1 Training Data จะมีกระบวนการ ดังนี้

Pre-Processing ปรับปรุงคุณภาพของตัวข้อมูลที่น่าเข้ามา เพื่อเตรียมไว้สำหรับขั้นตอนการทำ Feature Extraction ได้ ในขั้นตอนนี้จะมีการทำ Binarization, Morphological และ Segmentation

Feature Extraction ลดปริมาณตัวข้อมูล โดยทำการแยกเฉพาะข้อมูลที่เกี่ยวข้องเท่านั้น ในขั้นตอนนี้จะมีการทำ Hough Transform คือการหาเส้นตรงในรูปภาพ จากจุดต่างๆ ในรูปภาพที่มีความน่าจะเป็นเส้นตรงได้ Chain Code Transform คือใช้แสดงทิศทางของเส้นตรงที่มีการเชื่อมต่อกันของแต่ละจุด โดยมี Contour เป็นเส้นขอบของรูปร่างข้อมูลที่น่าเข้าไป เพื่อให้สกัดเอาคุณลักษณะ รูปร่างของข้อมูลที่น่าเข้าไป

Characters Estimation หลังจากทำ Feature Extraction เสร็จจะได้ชุดข้อมูลตัวต้นแบบที่จะนำไปใช้ในขั้นตอนทำ Pattern Matching

ส่วนที่ 2 Testing จะมีกระบวนการ ดังนี้

Pre-Processing ปรับปรุงคุณภาพของข้อมูลทดสอบที่น่าเข้ามา เพื่อเตรียมทำ Feature extraction โดยใช้หลักการเดียวกันกับการทำ Training Data

Feature Extraction ใช้หลักการเดียวกันกับการทำ Training Data

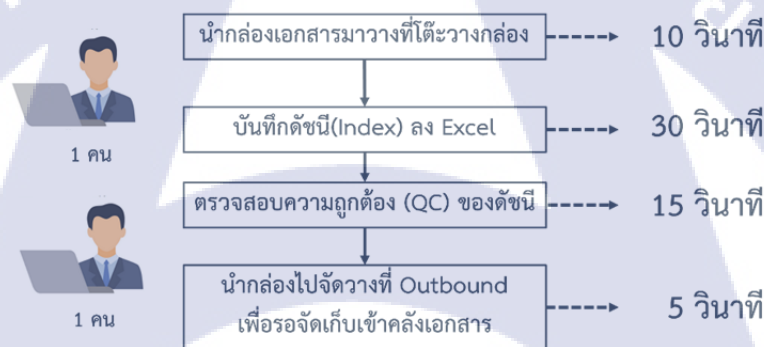
Classification เป็นการเปรียบเทียบข้อมูลของตัวข้อมูลทดสอบกับข้อมูลตัวต้นแบบเพื่อหาข้อมูลที่มีความใกล้เคียงที่สุด จากนั้นจะทำการแปลผลออกมา

บทที่ 3

วิธีการดำเนินงานวิจัย

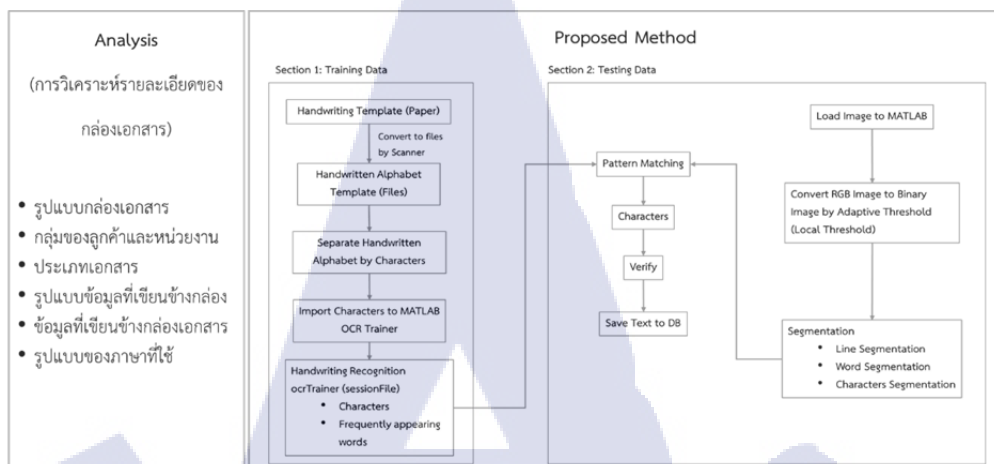
บทนี้จะกล่าวถึงวิธีการดำเนินงานวิจัยของการเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือ ซึ่งแต่เดิมการจัดทำดัชนีบนกล่องเอกสารจะประกอบด้วยขั้นแรก นำกล่องมาวางที่จุดที่จะทำการบันทึกข้อมูล ใช้เวลา 10 วินาที ขั้นที่ 2 พนักงานทำการบันทึกข้อมูลจากรูปภาพ ใช้เวลา 30 วินาที ขั้นที่ 3 พนักงานทำการ QC ข้อมูลที่ทำการบันทึกใช้เวลา 15 วินาที ขั้นที่ 4 นำกล่องไปวางที่จุด Outbound ใช้เวลา 5 วินาที ขั้นที่ 5 นำเข้าสู่ระบบ

รวมระยะเวลาที่ใช้เฉพาะขั้นตอนที่ 1 ถึง 4 ใช้เวลาการดำเนินการใน 1 กล่อง ประมาณ 60 วินาที (1 นาที) โดยการบันทึกข้อมูลจะใช้เวลาร้อยละ 50 ของทั้งกระบวนการ ดังรูปที่ 3.1



รูปที่ 3.1 ขั้นตอนเดิมของกระบวนการบันทึกข้อมูลดัชนี

งานวิจัยนี้จะนำเสนอกรอบแนวคิดและการพัฒนาต้นแบบ (Prototype) ที่สามารถแปลงดัชนี (Index) ที่เป็นข้อความที่อยู่บนกล่องเอกสารให้เป็นตัวอักษรที่สามารถใช้งานได้บนคอมพิวเตอร์ เพื่อลดขั้นตอนการดำเนินงานของผู้ปฏิบัติงานและต้นทุนด้านการดำเนินงานของฝ่ายปฏิบัติการ โดยมีแบ่งกรอบแนวคิดออกเป็น 2 ส่วน ส่วนที่ 1 การวิเคราะห์รายละเอียดของรูปแบบกล่องเอกสาร ส่วนที่ 2 การนำเสนอแนวคิดเพื่อพัฒนาต้นแบบที่สามารถแปลงดัชนี (Index) บนกล่องเอกสารให้เป็นตัวอักษร (Proposed Method) ดังรูปที่ 3.2



รูปที่ 3.2 ขอบเขตขั้นตอนการดำเนินงานวิจัย

3.1 ขั้นตอนการวิเคราะห์และออกแบบระบบ

การจัดทำดัชนีบนกล่องเอกสาร มีปัจจัยที่ต้องทำการวิเคราะห์ด้านต่างๆ ดังต่อไปนี้

3.1.1 การวิเคราะห์คุณลักษณะของกล่องเอกสารของผู้ให้บริการ

ธุรกิจการให้บริการจัดเก็บเอกสารในประเทศไทยจะมีผู้ให้บริการหลักๆ อยู่ 6 บริษัท ซึ่งจะยังไม่รวมถึงกล่องเก็บเอกสารที่หน่วยงานหรือบริษัทฯ ต่างๆ จำผลิตขึ้นมาให้เองภายใน โดยแต่ละบริษัทจะมีลักษณะของกล่องเอกสารที่มีรูปแบบการให้ลูกค้าจัดทำดัชนีที่ข้างกล่องแตกต่างกัน ดังรูปที่ 3.3 ทำให้ขั้นตอนที่จะทำการรู้จำตัวอักษรในภาพกล่องเอกสารจะต้องระบุตำแหน่งของส่วนต่างๆ ให้ชัดเจน เพื่อให้รู้จำแต่ละส่วนได้แม่นยำยิ่งขึ้น เพราะแต่ละตำแหน่งจะมีข้อจำกัดในการเขียน เช่น เขียนเฉพาะตัวเลข ตัวและตัวอักษรภาษาอังกฤษ หรือผสมทั้งตัวเลข ตัวอักษรภาษาไทย และภาษาอังกฤษ เป็นต้น



รูปที่ 3.3 ตัวอย่างรูปแบบของกล่องที่จะจัดทำดัชนี

3.1.2 การวิเคราะห์ประเภทกลุ่มลูกค้า

ลูกค้าที่ทำการจัดเก็บเอกสารจะแบ่งกลุ่มตามรูปแบบของการดำเนินงาน เป็น 2 กลุ่มใหญ่ ได้แก่ กลุ่มที่ 1 หน่วยงานราชการและรัฐวิสาหกิจ กลุ่มที่ 2 หน่วยงานเอกชน โดย 2 กลุ่มนี้จะมีรูปแบบการเขียนเอกสารข้างกล่องของลูกค้า จะมีความแตกต่างกัน หน่วยงานราชการจะนิยมเขียนเป็นตัวอักษรภาษาไทยและตัวเลขอารบิก ส่วนหน่วยงานเอกชนจะนิยมเขียนตัวอักษรภาษาไทย ตัวอักษรภาษาอังกฤษ และตัวเลขอารบิก ซึ่งสามารถแยกกลุ่มของกล่องเอกสารที่เข้ามาในแต่ละครั้ง เพื่อให้ขั้นตอนรู้จำตัวอักษรภาพ สามารถเจาะจงไปยังภาษาใดภาษาหนึ่งได้ทันที นอกจากนี้คำศัพท์ที่ใช้เขียนข้างกล่องของเอกสารประเภทเดียวกันของหน่วยงานราชการและรัฐวิสาหกิจกับหน่วยงานเอกชน จะมีการใช้คำศัพท์ต่างกัน เช่น การเขียนเอกสารประเภทการขอเบิกจ่ายเงิน หน่วยงานราชการจะใช้ คำว่า “ใบฎีกา” ส่วนหน่วยงานเอกชนจะใช้คำว่า “ใบตั้งเบิก” เป็นต้น ในส่วนนี้จะทำให้เวลาที่นำข้อมูลที่ได้จากการรู้จำตัวอักษรไปใช้งาน จะทำให้สามารถแปลข้อมูลได้ถูกต้องมากยิ่งขึ้น

3.1.3 การวิเคราะห์หน่วยงานย่อยภายใน

กล่องเอกสารที่ลูกค้าส่งมาจัดเก็บนั้นสามารถส่งมาจากหลากหลายฝ่ายที่เป็นหน่วยงานภายใน แต่กล่องเอกสารส่วนใหญ่ที่นำมาฝากจะเป็นเอกสารที่ต้องมีการจัดเก็บตามอายุของกฎหมาย โดยจะมาจากฝ่ายการเงิน ฝ่ายกฎหมาย ฝ่ายบุคคล และฝ่ายปฏิบัติการ ซึ่งเป็นหน่วยงานหลักที่จะต้องให้บริการจัดเก็บเอกสาร ซึ่งทำการจัดเก็บตามประเภทเอกสาร (Document Type) และตามเดือนของเอกสารที่เกิดขึ้น เช่น เอกสารใบกำกับภาษี เดือนกันยายน 2558 หรือใบเสร็จรับเงิน เดือนสิงหาคม 2558 เป็นต้น ในแต่ละเดือนที่มีการส่งกล่องเอกสารมาจัดเก็บจะเป็นประเภทเอกสารเดิมๆ ที่จัดเก็บเป็นประจำ จะเห็นว่าถ้ามีการออกแบบให้ขั้นตอนการจับคู่และแปลตัวอักษร สามารถรู้ได้ว่าเป็นเอกสารของฝ่ายใด จะทำให้การรู้จำและการแปลผลสามารถทำได้ถูกต้องมากยิ่งขึ้น

3.1.4 การวิเคราะห์รูปแบบลายมือที่ใช้เขียนกล่องเอกสาร

การเขียนดัชนีข้างกล่องเอกสาร สำหรับใช้ค้นหาเอกสารว่าจัดเก็บอยู่ในกล่องใด จะนิยมเขียนด้วยปากกาเคมีที่มีขนาดใหญ่สามารถมองเห็นได้ชัดเจนในระยะห่าง 2-3 เมตรและเขียนในรูปแบบที่เป็นคำสำคัญ (Keyword) เช่น ใบสั่งซื้อ เดือนกันยายน 2558 จะนิยมเขียนเป็น PO. กย. 2558 หรือ PO.-SEP 2015 เป็นต้น โดยลายมือเขียนที่ปรากฏบนกล่องเอกสารส่วนใหญ่จะเป็นลายมือของเจ้าของเอกสารกล่องนั้นๆ หรือลายมือของผู้ที่ถูกกำหนดให้รับผิดชอบในการจัดทำดัชนีของกล่องเอกสารนั้น ซึ่งจะทำให้ตัวอักษรเดียวกันแต่ละมีคุณลักษณะเฉพาะของลายมือที่แตกต่าง เช่น คนที่ 1 เขียนตัว ข ไข่ แบบมีหัว แต่คนที่ 2 เขียนแบบไม่มีหัว เป็นต้น ซึ่งในขั้นตอนการจับคู่และแปลตัวอักษร จะทำให้มีรูปแบบที่ต้องกำหนดในขั้นตอนนี้เพิ่มขึ้น

3.1.5 การวิเคราะห์ข้อมูลที่ใช้เขียนข้างกล่องเอกสาร

ข้อมูลที่ใช้เขียนข้างกล่องส่วนใหญ่จะเป็นข้อมูลที่ต้องเขียนตามรูปแบบกล่องเอกสารของแต่ละผู้ให้บริการกำหนดมา โดยข้อมูลส่วนใหญ่ที่กำหนดจะเป็นข้อมูลมาตรฐานที่เหมือนกัน ข้อมูลส่วนใหญ่จะเป็นการใส่ข้อมูลในรูปแบบลายมือเขียน แต่จะมีข้อมูลที่เป็นตำแหน่งติดบาร์โค้ดเท่านั้นที่ใช้เป็นสติ๊กเกอร์บาร์โค้ดดังตารางที่ 3.1 ด้านล่าง

ตารางที่ 3.1 ข้อมูลและรูปแบบการบันทึกข้อมูลของดัชนีกล่องเอกสาร

ลำดับ	ข้อมูลดัชนีกล่องเอกสาร	รูปแบบการบันทึกข้อมูล	ความสำคัญ	ประเภทข้อมูลที่กรอก
1	ตำแหน่งติดบาร์โค้ด	สติ๊กเกอร์บาร์โค้ด	บังคับ	ตัวเลขอารบิก
2	รหัสลูกค้า	ลายมือเขียน	บังคับ	ตัวเลขอารบิกและตัวอักษรภาษาอังกฤษ
3	ชื่อกล่องลูกค้า	ลายมือเขียน	บังคับ	ตัวอักษรภาษาอังกฤษ ตัวอักษรภาษาไทยและ ตัวเลขอารบิก
4	วันที่ทำลายเอกสาร	ลายมือเขียน	ไม่บังคับ	ตัวเลขอารบิก
5	รายละเอียดกล่องเอกสาร	ลายมือเขียน	บังคับ	ตัวอักษรภาษาอังกฤษ ตัวอักษรภาษาไทยและ ตัวเลขอารบิก

3.1.6 การวิเคราะห์ภาษาที่ใช้เขียนบนกล่องเอกสาร

การเขียนดัชนีกล่องเอกสารจะมีรูปแบบการใช้ตัวอักษรที่เขียนอยู่ 3 รูปแบบ ดังนี้

แบบที่ 1 ตัวอักษรภาษาไทยและตัวเลขอารบิก

แบบที่ 2 ตัวอักษรภาษาอังกฤษและตัวเลขอารบิก

แบบที่ 3 ตัวอักษรภาษาไทย ตัวอักษรภาษาอังกฤษและตัวเลขอารบิก

การเขียนตัวอักษรไทย สามารถแบ่งระดับบรรทัดของการเขียนแต่ละตัวอักษรออกเป็น 4 ระดับ ตามรูปที่ 3.4

เอกสารแสดงสิทธิการถือหุ้นในบริษัท

ระดับที่ 1
ระดับที่ 2
ระดับที่ 3
ระดับที่ 4

รูปที่ 3.4 บรรทัดของการเขียนตัวอักษรภาษาไทย

ตำแหน่งระดับการเขียนพยัญชนะ ตามตารางที่ 3.2

ตารางที่ 3.2 การแบ่งระดับของการเขียนพยัญชนะภาษาไทย

ตัวอักษร	ระดับที่				ประเภท
	1	2	3	4	
ก			/		ตัวอักษรไทย
ข			/		ตัวอักษรไทย
ฃ			/		ตัวอักษรไทย
ค			/		ตัวอักษรไทย
ฅ			/		ตัวอักษรไทย
ฉ			/		ตัวอักษรไทย
ง			/		ตัวอักษรไทย
จ			/		ตัวอักษรไทย
ฉ			/		ตัวอักษรไทย
ช		/	/		ตัวอักษรไทย
ซ		/	/		ตัวอักษรไทย
ฌ			/		ตัวอักษรไทย
ญ			/		ตัวอักษรไทย
ฎ			/	/	ตัวอักษรไทย
ฏ			/	/	ตัวอักษรไทย
ฐ			/	/	ตัวอักษรไทย
ฑ			/		ตัวอักษรไทย
ฒ			/		ตัวอักษรไทย
ณ			/		ตัวอักษรไทย
ด			/		ตัวอักษรไทย
ต			/		ตัวอักษรไทย
ถ			/		ตัวอักษรไทย
ท			/		ตัวอักษรไทย
ธ			/		ตัวอักษรไทย
น			/		ตัวอักษรไทย
บ			/		ตัวอักษรไทย
ป		/	/		ตัวอักษรไทย
ผ			/		ตัวอักษรไทย
ฝ		/	/		ตัวอักษรไทย
พ			/		ตัวอักษรไทย
ฟ		/	/		ตัวอักษรไทย
ภ			/		ตัวอักษรไทย
ม			/		ตัวอักษรไทย
ย			/		ตัวอักษรไทย
ร			/		ตัวอักษรไทย
ล			/		ตัวอักษรไทย
ว			/		ตัวอักษรไทย
ศ		/	/		ตัวอักษรไทย
ษ			/		ตัวอักษรไทย
ส		/	/		ตัวอักษรไทย
ห			/		ตัวอักษรไทย
ฬ		/	/		ตัวอักษรไทย
อ			/		ตัวอักษรไทย
ฮ		/	/		ตัวอักษรไทย

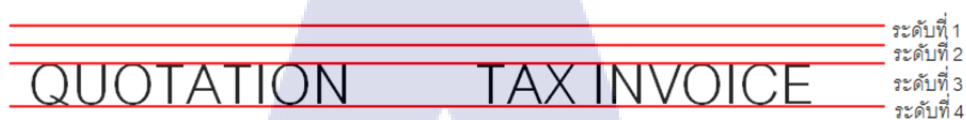
ตำแหน่งระดับการเขียนสระ วรรณยุกต์ และตัวเลขไทย ตามตารางที่ 3.3

ตารางที่ 3.3 การแบ่งระดับของการเขียนสระ วรรณยุกต์และตัวเลขภาษาไทย

ตัวอักษร	ระดับที่				ประเภท
	1	2	3	4	
อะ			/		สระ
อา			/		สระ
อิ		/			สระ
อี		/			สระ
อึ		/			สระ
อื		/			สระ
อุ				/	สระ
อู				/	สระ
เอะ			/		สระ
เอ			/		สระ
แอะ			/		สระ
แอ			/		สระ
เอียะ		/	/		สระ
เอีย		/	/		สระ
เอือะ		/	/		สระ
เอือ		/	/		สระ
อัวะ		/	/		สระ
อัว		/	/		สระ
โอะ		/	/		สระ
โอ		/	/		สระ
เอะ			/		สระ
ออ			/		สระ
เออะ			/		สระ
เออ			/		สระ
อำ		/	/		สระ
ไอ		/	/		สระ
ไอ		/	/		สระ
เอา			/		สระ
ฤ			/	/	สระ
ฦ			/	/	สระ
ฦ			/	/	สระ
ฦ			/	/	สระ
ไม่ได้คู่		/			สระ
ไม้เอก	/	/			วรรณยุกต์
ไม้โท	/	/			วรรณยุกต์
ไม้ตรี	/	/			วรรณยุกต์
ไม้จัตวา	/	/			วรรณยุกต์
การันต์	/	/			
๑			/		ตัวเลขไทย
๒		/	/		ตัวเลขไทย
๓			/		ตัวเลขไทย
๔		/	/		ตัวเลขไทย
๕		/	/		ตัวเลขไทย
๖		/	/		ตัวเลขไทย
๗		/	/		ตัวเลขไทย
๘		/	/		ตัวเลขไทย
๙		/	/		ตัวเลขไทย
๐			/		ตัวเลขไทย

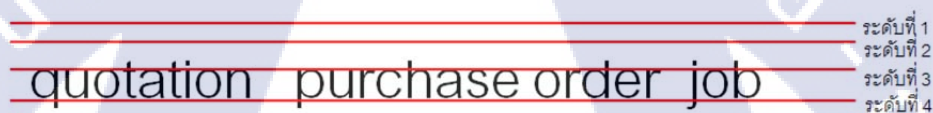
การเขียนตัวอักษรภาษาอังกฤษ สามารถแบ่งระดับบรรทัดของการเขียนแต่ละตัวอักษร
ออกเป็น 2 รูปแบบ ได้แก่

แบบที่ 1 เขียนตัวพิมพ์ใหญ่อย่างเดียว ดังรูปที่ 3.5



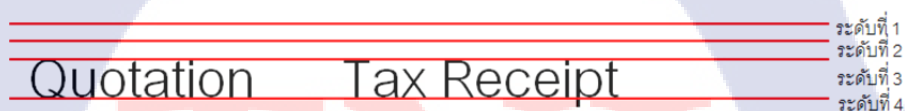
รูปที่ 3.5 ระดับการเขียนตัวอักษรภาษาอังกฤษแบบตัวพิมพ์ใหญ่

แบบที่ 2 เขียนตัวพิมพ์เล็กอย่างเดียว ดังรูปที่ 3.6



รูปที่ 3.6 ระดับการเขียนตัวอักษรภาษาอังกฤษแบบตัวพิมพ์เล็ก

แบบที่ 3 เขียนตัวพิมพ์ใหญ่และตัวพิมพ์เล็ก ดังรูปที่ 3.7



รูปที่ 3.7 ระดับการเขียนตัวอักษรภาษาอังกฤษแบบตัวพิมพ์ใหญ่และตัวพิมพ์เล็ก

ตำแหน่งระดับการเขียนภาษาอังกฤษตัวพิมพ์ใหญ่ ตามตารางที่ 3.4

ตารางที่ 3.4 การแบ่งระดับการเขียนภาษาอังกฤษตัวพิมพ์ใหญ่

ตัวอักษร	ระดับ				ประเภท
	1	2	3	4	
A			/		พยัญชนะพิมพ์ใหญ่
B			/		พยัญชนะพิมพ์ใหญ่
C			/		พยัญชนะพิมพ์ใหญ่
D			/		พยัญชนะพิมพ์ใหญ่
E			/		พยัญชนะพิมพ์ใหญ่
F			/		พยัญชนะพิมพ์ใหญ่
G			/		พยัญชนะพิมพ์ใหญ่
H			/		พยัญชนะพิมพ์ใหญ่
I			/		พยัญชนะพิมพ์ใหญ่
J			/		พยัญชนะพิมพ์ใหญ่
K			/		พยัญชนะพิมพ์ใหญ่
L			/		พยัญชนะพิมพ์ใหญ่
M			/		พยัญชนะพิมพ์ใหญ่

ตัวอักษร	ระดับ				ประเภท
	1	2	3	4	
N			/		พยัญชนะพิมพ์ใหญ่
O			/		พยัญชนะพิมพ์ใหญ่
P			/		พยัญชนะพิมพ์ใหญ่
Q			/		พยัญชนะพิมพ์ใหญ่
R			/		พยัญชนะพิมพ์ใหญ่
S			/		พยัญชนะพิมพ์ใหญ่
T			/		พยัญชนะพิมพ์ใหญ่
U			/		พยัญชนะพิมพ์ใหญ่
V			/		พยัญชนะพิมพ์ใหญ่
W			/		พยัญชนะพิมพ์ใหญ่
X			/		พยัญชนะพิมพ์ใหญ่
Y			/		พยัญชนะพิมพ์ใหญ่
Z			/		พยัญชนะพิมพ์ใหญ่

ตำแหน่งระดับการเขียนภาษาอังกฤษตัวพิมพ์เล็ก ตัวเลขอารบิกและสัญลักษณ์ ตาม
ตารางที่ 3.5

ตารางที่ 3.5 การแบ่งระดับการเขียนภาษาอังกฤษตัวพิมพ์เล็ก ตัวเลขอารบิกและสัญลักษณ์

ตัวอักษร	ระดับ				ประเภท
	1	2	3	4	
a			/		พยัญชนะพิมพ์เล็ก
b		/	/		พยัญชนะพิมพ์เล็ก
c			/		พยัญชนะพิมพ์เล็ก
d		/	/		พยัญชนะพิมพ์เล็ก
e			/		พยัญชนะพิมพ์เล็ก
f		/	/		พยัญชนะพิมพ์เล็ก
g			/	/	พยัญชนะพิมพ์เล็ก
h		/	/		พยัญชนะพิมพ์เล็ก
i		/	/		พยัญชนะพิมพ์เล็ก
j		/	/	/	พยัญชนะพิมพ์เล็ก

ตัวอักษร	ระดับ				ประเภท
	1	2	3	4	
k		/	/		พยัญชนะพิมพ์เล็ก
l		/	/		พยัญชนะพิมพ์เล็ก
m			/		พยัญชนะพิมพ์เล็ก
n			/		พยัญชนะพิมพ์เล็ก
o			/		พยัญชนะพิมพ์เล็ก
p			/	/	พยัญชนะพิมพ์เล็ก
q			/	/	พยัญชนะพิมพ์เล็ก
r			/		พยัญชนะพิมพ์เล็ก
s			/		พยัญชนะพิมพ์เล็ก
t		/	/		พยัญชนะพิมพ์เล็ก

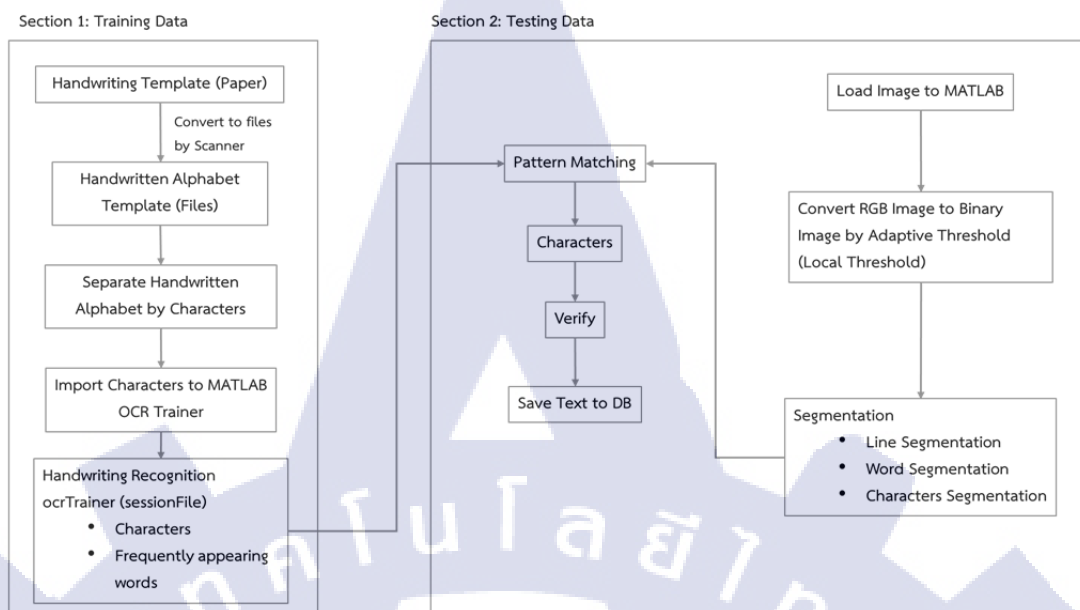
ตารางที่ 3.5 การแบ่งระดับการเขียนภาษาอังกฤษตัวพิมพ์เล็ก ตัวเลขอารบิกและสัญลักษณ์ (ต่อ)

ตัวอักษร	ระดับ				ประเภท
	1	2	3	4	
u			/		พยัญชนะพิมพ์เล็ก
v			/		พยัญชนะพิมพ์เล็ก
w			/		พยัญชนะพิมพ์เล็ก
x			/		พยัญชนะพิมพ์เล็ก
y			/	/	พยัญชนะพิมพ์เล็ก
z			/		พยัญชนะพิมพ์เล็ก
1			/		ตัวเลขอารบิก
2			/		ตัวเลขอารบิก
3			/		ตัวเลขอารบิก
4			/		ตัวเลขอารบิก
5			/		ตัวเลขอารบิก
6			/		ตัวเลขอารบิก
7			/		ตัวเลขอารบิก
8			/		ตัวเลขอารบิก

ตัวอักษร	ระดับ				ประเภท
	1	2	3	4	
9			/		ตัวเลขอารบิก
0			/		ตัวเลขอารบิก
‘ , “			/		สัญลักษณ์
: , ; .			/		สัญลักษณ์
() , [] , { }			/		สัญลักษณ์
*			/		สัญลักษณ์
+			/		สัญลักษณ์
-			/		สัญลักษณ์
/ , \			/		สัญลักษณ์
#			/		สัญลักษณ์
%			/		สัญลักษณ์
&			/		สัญลักษณ์
?			/		สัญลักษณ์
=			/		สัญลักษณ์

3.2 การนำเสนอแนวคิดเพื่อพัฒนาต้นแบบที่สามารถแปลงดัชนี (Index) บนกล่องเอกสารให้เป็นตัวอักษร (Proposed Method)

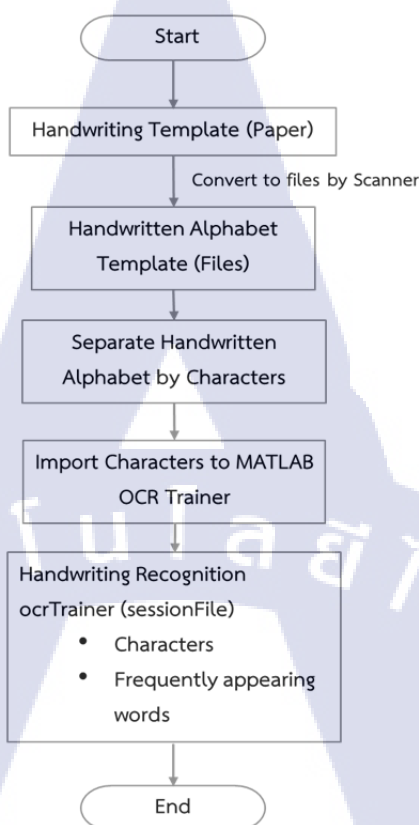
กรอบแนวคิดและการพัฒนาต้นแบบที่สามารถแปลงดัชนี (Index) ที่เป็นข้อความที่อยู่บนกล่องเอกสารให้เป็นตัวอักษรที่สามารถใช้งานได้บนคอมพิวเตอร์ จะแบ่งเป็น 2 Sections ได้แก่ Training Data และ Testing Data ดังรูปที่ 3.8



รูปที่ 3.8 แผนภาพการนำเสนอแนวคิดเพื่อพัฒนาต้นแบบ

3.2.1 ขั้นตอนการเตรียมตัวอักษรต้นแบบ (Training Data)

ขั้นตอนนี้จะเป็นการเตรียมตัวแบบอักษรที่จะใช้ในประมวลจับคู่รูปแบบของตัวอักษรที่ปรากฏบนกล่องเอกสารให้เป็นตัวอักษรที่ใช้งานบนคอมพิวเตอร์ได้ ในการเตรียมตัวแบบอักษรนี้จะให้ผู้ปฏิบัติงานทำการเขียนตัวอย่างตัวอักษรและคำที่พบเป็นประจำลงกระดาษ เหตุผลที่ต้องเขียนคำที่พบเป็นประจำเข้าไป เพราะผู้ปฏิบัติงานส่วนใหญ่จะเขียนด้วยลายมือของคนเดิมและแบบเดิม ทำให้จุดนี้เป็นจุดที่เพิ่มความแม่นยำของการแปลงผลของตัวอักษรได้ดียิ่งขึ้น จากนั้นจะนำลายมือที่เขียนตัวอักษรมาสแกนเป็นไฟล์และตัดตัวอักษรแต่ละตัว จากนั้นจะนำเข้า OCR Trainer ใน MATLAB โดยมีการปรับปรุงคุณภาพโดยใช้ Global Thresholding และกำหนด Segmentation แบบอัตโนมัติหรือทำการกำหนดเองโดยผู้ใช้งานของตัวอักษรแต่ละตัวที่จะให้ OCR Trainer ทำการรู้จำว่าตัวอักษรแต่ละตัวคือตัวอักษรอะไร สำหรับเตรียมไว้ใช้สำหรับตอนที่ทำการประมวลจับคู่รูปแบบของตัวอักษร ดังรูปที่ 3.9



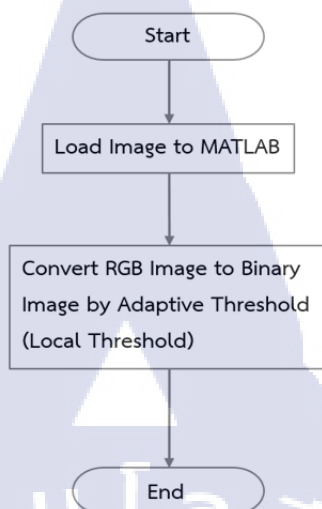
รูปที่ 3.9 ขั้นตอนการเตรียมตัวอักษรต้นแบบ (Training Data)

3.2.2 ขั้นตอนการทดสอบภาพตัวอย่าง (Testing Data)

ขั้นตอนนี้จะมีขั้นตอนย่อย 4 ขั้นตอน ได้แก่

ขั้นตอนที่ 1 การปรับปรุงคุณภาพของภาพ (Pre- Processing)

ขั้นตอนของการปรับปรุงคุณภาพของภาพ (Pre- Processing) ทำเพื่อจัดเตรียมภาพที่อยู่ในคุณลักษณะที่จะนำไปเข้าสู่กระบวนการประมวลผล ได้อย่างมีประสิทธิภาพ ขั้นตอนจะเริ่มจากนำรูปภาพสีที่ได้จากกล้องถ่ายรูปที่ถ่ายไว้แล้ว ซึ่งรูปภาพที่นำเข้ามาจะเป็นการถ่ายภาพกล่องเอกสารโดยมีระยะห่าง 30 เซนติเมตรระหว่างกล่องกับกล่องเอกสาร เพื่อที่จะสามารถเก็บข้อมูลรายละเอียดของข้างกล่องที่ได้จัดทำดัชนีไว้ จากนั้นนำภาพสีไปแปลงให้เป็นภาพขาวดำ (Binary Image) โดยทำการปรับปรุงคุณภาพใช้ Local Thresholding เพื่อกำจัดแสงเงาที่เข้ามาในรูปภาพ เพื่อให้ในขั้นตอนถัดไปสามารถแปลงตัวอักษรได้ง่ายขึ้น จากนั้นจะทำการแยกส่วนของภาพ (Segmentation) เพื่อเตรียมนำเข้ากระบวนการรู้จำตัวอักษร ดังรูปที่ 3.10



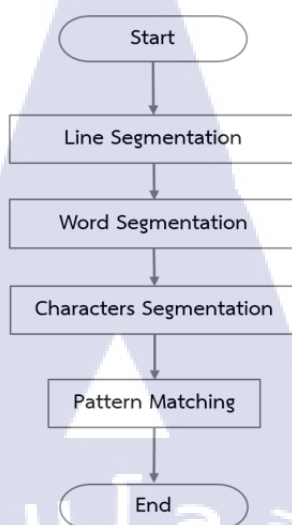
รูปที่ 3.10 ขั้นตอนการทำงานการปรับปรุงคุณภาพของภาพ (Pre-Processing)

ขั้นตอนที่ 2 การหา Segmentation ของตัวอักษรที่จะทดลอง

ขั้นตอนการหา Segmentation ของตัวอักษรที่จะทดลองจะเกิดขึ้นหลังจากที่ได้ทำการปรับปรุงคุณภาพของภาพเสร็จเรียบร้อยแล้ว ทำให้ได้ภาพกล่องเอกสารที่เป็นภาพขาวดำ (Binary Image) ก่อนจะนำภาพเอกสารที่ได้ไปทำการแปลงตัวอักษรแต่ละตัวให้อยู่ในรูปแบบของตัวอักษรภาษาไทย ตัวอักษรภาษาอังกฤษ สัญลักษณ์หรือตัวเลขออกมา จะหา Segmentation ของตัวอักษร โดยเริ่มจากการหาบรรทัด คำ และตัวอักษรที่จะทำการแปลความหมาย

ขั้นตอนที่ 3 การรู้จำตัวอักษรของภาพ (Optical Character Recognition)

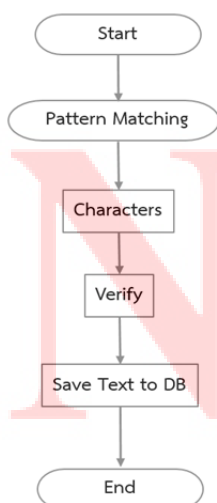
จากขั้นตอนที่ 2 จะนำตัวอักษรภาษาไทยและภาษาอังกฤษที่ได้ มาทำการรู้จำตัวอักษรของภาพ (Optical Character Recognition) และแปลความหมายของตัวอักษรที่ได้เทียบกับพจนานุกรมคำศัพท์ที่สร้างขึ้นมา เพื่อให้ได้ความหมายของคำศัพท์ที่เป็นชื่อเต็มของตัวอักษรย่อ นั้นๆ เพื่อสามารถนำไปใช้งานในคอมพิวเตอร์ได้ ดังรูปที่ 3.11



รูปที่ 3.11 การหา Segmentation ของตัวอักษรที่จะทดลอง

ขั้นตอนที่ 4 การตรวจสอบความถูกต้องโดยมนุษย์ (Human Intervention)

เมื่อกระบวนการรู้จำตัวอักษรของภาพ (Optical Character Recognition) ได้ทำงานเสร็จสิ้น จะให้ผู้ปฏิบัติงานเข้ามาตรวจสอบความถูกต้องของตัวอักษรที่ได้แปลงข้อมูล ก่อนที่จะทำการบันทึกข้อมูลเก็บในฐานข้อมูลต่อไป ดังรูปที่ 3.12



รูปที่ 3.12 ขั้นตอนการตรวจสอบความถูกต้องโดยมนุษย์ (Human Intervention)

3.3 การวัดประสิทธิภาพและความถูกต้อง

สำหรับการประสิทธิภาพและความถูกต้องของงานวิจัยนี้จะทำการวัด อยู่ 3 ส่วน ได้แก่ ส่วนที่ 1 วัดเป็นค่าร้อยละของความถูกต้องของชุดข้อมูลทดสอบ (Testing Data) กับชุดข้อมูลฝึกสอน (Training Data)

ส่วนที่ 2 วัดค่าเป็นร้อยละของความถูกต้องของชุดข้อมูลทดสอบ (Testing Data) ในแต่ละช่องที่ปรากฏอยู่บนดัชนีกล่องเอกสารจากจำนวนข้อมูลทดสอบทั้งหมด

ส่วนที่ 3 วัดเป็นค่าร้อยละของประสิทธิภาพของเวลาที่ใช้ในการบันทึกข้อมูลดัชนี (Index) ด้วยตัวต้นแบบกับการบันทึกข้อมูลดัชนี (Index) ด้วยวิธีเดิม

3.4 เครื่องมือที่ใช้พัฒนา

3.4.1 โปรแกรม MATLAB เวอร์ชัน R2018b โดยมี Toolbox ที่ใช้งาน ดังนี้

- Image Processing Toolbox
- Image Acquisition Toolbox
- Computer Vision System Toolbox

3.4.2 เครื่องคอมพิวเตอร์ มีคุณสมบัติขั้นต่ำ ดังนี้

- CPU: Intel® Core™ i3-7130U CPU @ 2.70GHz 2.71 GHz
- System Type: 64-bit Operating System, x64 based processor
- OS: Windows version 10
- Memory (Ram): 8 GB
- HDD: 120 GB

3.4.3 กล้องถ่ายภาพดิจิทัล ที่มีความละเอียด 5 ล้านพิกเซลขึ้นไป

บทที่ 4

ผลของการวิจัย

การดำเนินการวิจัยมีวัตถุประสงค์การวิจัย 3 ข้อ ดังนี้

1. เพื่อจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้
2. เพื่อสามารถแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้มีความถูกต้อง (Accuracy) อย่างน้อย 80%
3. เพื่อลดระยะเวลาที่ใช้ในขั้นตอนการบันทึกข้อมูลดัชนี (Index) เป็น 30% ของกระบวนการเดิม

จากที่ได้ทำการทดลองได้ผลการดำเนินงานวิจัยเรื่อง การเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือ ดังนี้

1. ผลการสังเคราะห์รูปแบบการจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้ (วัตถุประสงค์ข้อ 1)
2. ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารมีความถูกต้อง (Accuracy) 87.30% (วัตถุประสงค์ข้อ 2)
3. ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเฉพาะกล่องของบริษัทลูกค้ามีความถูกต้อง (Accuracy) 85.49% (วัตถุประสงค์ข้อ 2) โดยมีรายละเอียดความถูกต้อง (Accuracy) การแปลงข้อมูลเฉพาะช่อง ดังนี้

ช่องสำนักงาน	เท่ากับ	80.00%
ช่องกลุ่มงาน	เท่ากับ	95.00%
ช่องกล่องเลขที่	เท่ากับ	82.23%
ช่องเลขที่อ้างอิง	เท่ากับ	87.74%
ช่องสถานที่เก็บของ	เท่ากับ	100.00%
ช่องวันที่ทำลาย	เท่ากับ	82.67%

4. ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเฉพาะกล่องของบริษัทที่ให้บริการจัดเก็บเอกสารมีความถูกต้อง (Accuracy) 88.52% (วัตถุประสงค์ข้อ 2) โดยมีรายละเอียดการแปลงข้อมูลเฉพาะช่อง ดังนี้

ช่องชื่อกล่องลูกค้า	เท่ากับ	93.89%
ช่องรายละเอียดกล่อง	เท่ากับ	85.09%

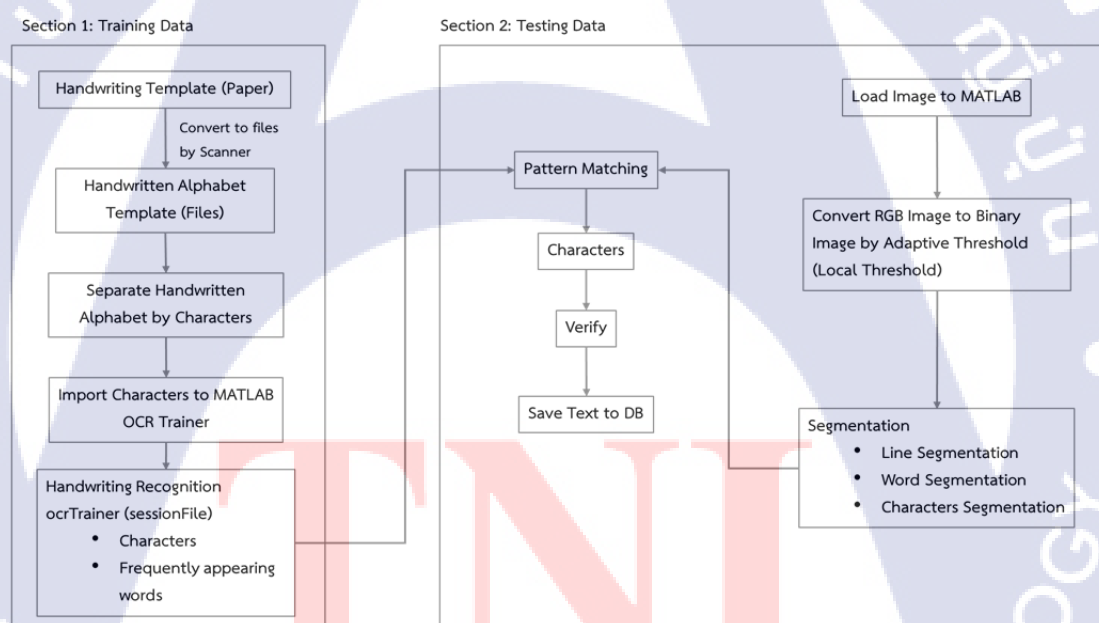
5. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงานด้านระยะเวลาลดลง 68.33% (วัตถุประสงค์ข้อ 3)

6. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงานด้านจำนวนคนลดลง 50.00% (วัตถุประสงค์ข้อ 3)

7. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงานด้านต้นทุนลดลง 50.00% (วัตถุประสงค์ข้อ 3)

4.1 ผลการสังเคราะห์รูปแบบการจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร

กำหนดกรอบแนวคิดเกี่ยวกับการจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร ประกอบด้วย 2 ส่วน คือ ส่วนที่ 1 Training Data ส่วนที่ 2 Testing Data โดยแต่ละส่วนจะมีความทำงานให้สอดคล้องกัน ดังรูปที่ 4.1



รูปที่ 4.1 แนวคิดการจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรม

โดยแต่ละส่วนมีรายละเอียดการทำงาน ดังนี้

ส่วนที่ 1 Training Data เป็นส่วนที่ออกแบบให้รองรับการนำเข้าตัวอย่างการเขียนตัวอักษรด้วยลายมือของผู้ใช้งาน เพื่อที่จะสอนให้ตัวโปรแกรมสามารถรู้จำว่าตัวอักษรแต่ละตัวคืออะไร โดย

ตัวอักษรแต่ละตัวที่ได้จะนำไปเตรียมไว้ใช้สำหรับตอนที่ทำการประมวลจับคู่รูปแบบของตัวอักษรที่ปรากฏบนกล่องเอกสารที่นำเข้ามาทดลอง

ส่วนที่ 2 Testing Data เป็นส่วนที่ออกแบบรองรับการนำเข้าภาพถ่ายกล่องเอกสารที่จะใช้ในการทดลอง โดยจะมีการปรับปรุงภาพให้อยู่ในรูปแบบภาพขาวดำ และทำการแยกส่วนของตัวอักษรเพื่อเตรียมไว้สำหรับการประมวลผล จากนั้นจะทำการประมวลผลจับคู่รูปแบบของตัวอักษรที่เตรียมไว้กับตัวอักษรที่ปรากฏบนกล่องเอกสารที่นำเข้ามาทดลอง โดยขั้นตอนนี้จะให้มนุษย์เข้ามาช่วยตรวจสอบหลังจากที่โปรแกรมทำการแปลงตัวอักษรออก เพื่อให้สามารถบันทึกข้อมูลลงฐานข้อมูลได้ถูกต้อง

4.2 ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร

จากการทดลองของกล่องตัวอย่างกล่องเอกสารที่นำมาทดลองจะแบ่งกล่องเอกสารออกเป็น 2 กลุ่ม ได้แก่ กลุ่มที่ 1 กล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า จำนวน 202 กล่อง กลุ่มที่ 2 กล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร จำนวน 234 กล่อง ซึ่ง ทั้ง 2 กลุ่มจะมีลักษณะเฉพาะของข้อมูลที่ปรากฏอยู่ข้างกล่องแตกต่างกัน

กลุ่มที่ 1 กล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า มีประมวลผลข้อมูล จำนวน 6 ช่อง ได้แก่ สำนักงาน กลุ่มงาน กล่องเลขที่ เลขที่อ้างอิง สถานที่เก็บของ วันที่ทำลาย ดังรูปที่ 4.2

กล่องเก็บเอกสาร (20๒๒)	
สำนักงาน <u>บัญชี</u>	กลุ่มงาน <u>POS</u>
กล่องเลขที่ <u>ดิจิทัล ๑๑</u>	เลขที่อ้างอิง <u>กันยายน/55</u>
สถานที่เก็บของ _____	วันที่ทำลาย _____

รูปที่ 4.2 ตัวอย่างกล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า

กลุ่มที่ 2 กล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร มีประมวลผลข้อมูล จำนวน 6 ช่อง ได้แก่ รหัสลูกค้า ชื่อกล่องลูกค้า ปี/เดือนทำลาย รายละเอียดของเอกสาร ดังรูปที่ 4.3

001008X 50000243233	Customer Code รหัสลูกค้า
Reference Code ชื่อกล่องลูกค้า กล่องที่ 5	Destruction Date ปี(ค.ศ.)/เดือน/ทำลาย
Document Description รายละเอียดของเอกสาร บันทึกเรียนจบชั้นมัธยมศึกษา(ศก๑)	

รูปที่ 4.3 ตัวอย่างกล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร

การประเมินผลการแปลงข้อมูลที่เขียนด้วยลายมือ ได้นำข้อมูลที่เป็นเฉพาะลายมือมาประมวลผลเท่านั้น โดยมีการกำหนดวิธีการประเมินผล ดังนี้

4.2.1 ประเมินผลความถูกต้องของข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่อง

ส่วนที่ 1 การประเมินผลการแปลงตัวอักษรที่แปลงได้ทั้งหมด

การประเมินผลภาพรวมของข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่องได้ผลดังนี้ จำนวนตัวอักษรที่ปรากฏบนกล่องทั้งหมดมี 13,160 ตัวอักษร มีจำนวนความถูกต้องของข้อมูลที่ทำ การแปลงออกมาเท่ากับ 87.30% คิดเป็น 11,489 ตัวอักษร จำนวนตัวอักษรที่ทำการแปลงผิดเท่ากับ 12.70% คิดเป็น 1,671 ตัวอักษร ดังตารางที่ 4.1

ตารางที่ 4.1 การประเมินผลภาพรวมของข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่อง

รายละเอียด	จำนวนอักษร	เปอร์เซ็นต์
ตัวอักษรทั้งหมด	13,160	100.00%
ตัวอักษรถูกต้อง	11,489	87.30%
ตัวอักษรผิด	1,671	12.70%

ส่วนที่ 2 การประเมินผลการแปลงตัวอักษรที่แปลงได้ โดยจำแนกตามกลุ่มรูปแบบของกล่องเอกสารได้ ดังนี้

กลุ่มที่ 1 กล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า มีจำนวนตัวอักษรที่ปรากฏบนกล่องทั้งหมดมี 7,215 ตัวอักษร พบว่ามีจำนวนความถูกต้องของข้อมูลที่ทำ การแปลงออก

มาเท่ากับ 85.49% คิดเป็น 6,168 ตัวอักษร จำนวนตัวอักษรที่ทำการแปลงผิดเท่ากับ 14.51% คิดเป็น 1,047 ตัวอักษร ดังตารางที่ 4.2

ตารางที่ 4.2 ผลข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า

รายละเอียด	จำนวนอักษร	เปอร์เซ็นต์
ตัวอักษรทั้งหมด	7,215	100.00%
ตัวอักษรถูกต้อง	6,168	85.49%
ตัวอักษรผิด	1,047	14.51%

กลุ่มที่ 2 กล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร จำนวนตัวอักษรที่ปรากฏบนกล่องทั้งหมดมี 5,945 ตัวอักษร พบว่ามีจำนวนความถูกต้องของข้อมูลที่ทำกรแปลงออกมาเท่ากับ 88.52% คิดเป็น 5,262 ตัวอักษร จำนวนตัวอักษรที่ทำการแปลงผิดเท่ากับ 11.48% คิดเป็น 683 ตัวอักษร ดังตารางที่ 4.3

ตารางที่ 4.3 ผลข้อมูลลายมือทั้งหมดที่ปรากฏบนกล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร

รายละเอียด	จำนวนอักษร	เปอร์เซ็นต์
ตัวอักษรทั้งหมด	5,945	100.00%
ตัวอักษรถูกต้อง	5,262	88.52%
ตัวอักษรผิด	683	11.48%

4.2.2 ประเมินผลความถูกต้องของข้อมูลลายมือที่ปรากฏในแต่ละช่องของกล่อง

กลุ่มที่ 1 กล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า จำนวนตัวอักษรในแต่ละช่องของกล่องเอกสาร มีทั้งหมด 6 ช่อง ดังตารางที่ 4.4 พบว่า

ช่องที่ 1 สำนักงาน มีจำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 975 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 80% คิดเป็น 780 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดเท่ากับ 20% คิดเป็น 195 ตัวอักษร

ช่องที่ 2 กลุ่มงาน จำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 1,352 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 95% คิดเป็น 1,284 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดเท่ากับ 5% คิดเป็น 68 ตัวอักษร

ช่องที่ 3 กล่องเลขที่ จำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 2,305 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 82.23% คิดเป็น 1,895 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดเท่ากับ 17.77% คิดเป็น 410 ตัวอักษร

ช่องที่ 4 เลขที่อ้างอิง จำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 2,367 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 87.74% คิดเป็น 2,077 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดเท่ากับ 12.16% คิดเป็น 290 ตัวอักษร

ช่องที่ 5 สถานที่เก็บของ จำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 66 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 100% คิดเป็น 66 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดไม่มี

ช่องที่ 6 วันที่ทำลาย จำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 150 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 82.67% คิดเป็น 124 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดเท่ากับ 17.33% คิดเป็น 26 ตัวอักษร

ตารางที่ 4.4 ผลความถูกต้องของข้อมูลลายมือที่ปรากฏในแต่ละช่องของกล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า

ช่องที่	รายละเอียด	จำนวนอักษร	เปอร์เซ็นต์
1. สำนักงาน	ตัวอักษรทั้งหมด	975	100.00%
	ตัวอักษรถูกต้อง	780	80.00%
	ตัวอักษรผิด	195	20.00%
2. กลุ่มงาน	ตัวอักษรทั้งหมด	1,352	100.00%
	ตัวอักษรถูกต้อง	1,284	95.00%
	ตัวอักษรผิด	68	5.00%
3. กล่องเลขที่	ตัวอักษรทั้งหมด	2,305	100.00%
	ตัวอักษรถูกต้อง	1,895	82.23%
	ตัวอักษรผิด	410	17.77%
4. เลขที่อ้างอิง	ตัวอักษรทั้งหมด	2,367	100.00%
	ตัวอักษรถูกต้อง	2,077	87.74%
	ตัวอักษรผิด	290	12.16%
5. สถานที่เก็บของ	ตัวอักษรทั้งหมด	66	100.00%
	ตัวอักษรถูกต้อง	66	100.00%
	ตัวอักษรผิด	0	0%

ตารางที่ 4.4 ผลความถูกต้องของข้อมูลลายมือที่ปรากฏในแต่ละช่องของกล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า (ต่อ)

ช่องที่	รายละเอียด	จำนวนอักษร	เปอร์เซ็นต์
6. วันที่ทำลาย	ตัวอักษรทั้งหมด	150	100.00%
	ตัวอักษรถูกต้อง	124	82.67%
	ตัวอักษรผิด	26	17.33%

กลุ่มที่ 2 กล่องของบริษัทที่ให้บริการจัดเก็บเอกสาร จำนวนตัวอักษรในแต่ละช่องของกล่องเอกสาร มีทั้งหมด 4 ช่อง ดังตารางที่ 4.5 พบว่า

ช่องที่ 1 Customer Code กล่องตัวอย่างที่นำมาวิจัยไม่ได้มีการเขียนข้อมูลในช่องนี้

ช่องที่ 2 ชื่อกล่องลูกค้า จำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 1,638 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 93.89% คิดเป็น 1,538 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดเท่ากับ 6.11% คิดเป็น 100 ตัวอักษร

ช่องที่ 3 ปีที่ทำลาย กล่องตัวอย่างที่นำมาวิจัยไม่ได้มีการเขียนข้อมูลในช่องนี้

ช่องที่ 4 รายละเอียดกล่อง จำนวนตัวอักษรที่ปรากฏบนกล่องเท่ากับ 4,307 ตัวอักษร แปลงตัวอักษรถูกต้องเท่ากับ 85.09% คิดเป็น 3,665 ตัวอักษร จำนวนตัวอักษรที่แปลงผิดเท่ากับ 14.91% คิดเป็น 642 ตัวอักษร

ตารางที่ 4.5 ผลความถูกต้องของข้อมูลลายมือที่ปรากฏในแต่ละช่องของกล่องบริษัทที่ให้บริการจัดเก็บเอกสาร

ช่องที่	รายละเอียด	จำนวนอักษร	เปอร์เซ็นต์
1. Customer Code	ตัวอักษรทั้งหมด	N/A	N/A
	ตัวอักษรถูกต้อง	N/A	N/A
	ตัวอักษรผิด	N/A	N/A
2. ชื่อกล่องลูกค้า	ตัวอักษรทั้งหมด	1,638	100.00%
	ตัวอักษรถูกต้อง	1,538	93.89%
	ตัวอักษรผิด	100	6.11%
3. ปีที่ทำลาย	ตัวอักษรทั้งหมด	N/A	N/A
	ตัวอักษรถูกต้อง	N/A	N/A
	ตัวอักษรผิด	N/A	N/A

ตารางที่ 4.5 ผลความถูกต้องของข้อมูลลายมือที่ปรากฏในแต่ละช่องของกล่องบริษัทที่ให้บริการ
จัดเก็บเอกสาร (ต่อ)

ช่องที่	รายละเอียด	จำนวนอักษร	เปอร์เซ็นต์
4. รายละเอียดกล่อง	ตัวอักษรทั้งหมด	4,307	100.00%
	ตัวอักษรถูกต้อง	3,665	85.09%
	ตัวอักษรผิด	642	14.91%

4.3 ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงาน

การประเมินผลในส่วนนี้จะเป็นการประเมินผลในเชิงความคุ้มค่าทางธุรกิจ ซึ่งเป็นการเปรียบเทียบวิธีการเดิมกับวิธีการใหม่ที่นำเข้ามาช่วยปรับปรุงการทำงานให้มีประสิทธิภาพมากขึ้น ขั้นตอนประเมินเฉพาะที่ทำการบันทึกข้อมูลเท่านั้น โดยปัจจัยสำหรับการประเมินผลอยู่ 3 ปัจจัย ได้แก่ จำนวนคน ระยะเวลาแต่ละขั้นตอนการทำงาน และค่าใช้จ่าย ดังที่แสดงผลการเปรียบเทียบในตารางที่ 4.6 ด้านล่าง

ตารางที่ 4.6 ผลการทำงานแบบเดิมเทียบกับการใช้โปรแกรมเข้ามาช่วยในการทำงาน

ปัจจัย	วิธีการเดิม	วิธีการใหม่	ผลการเปรียบเทียบ
ระยะเวลาแต่ละขั้นตอนการทำงาน	พนักงานคนที่ 1 - นำกล่องมาวาง (10 วินาที) - บันทึกข้อมูลลายมือข้างกล่อง (30 วินาที) พนักงานคนที่ 2 - ตรวจสอบความถูกต้อง (14 วินาที) - บันทึกข้อมูล (1 วินาที) - ย้ายกล่องไปเก็บ (5 วินาที) รวมเวลาที่ใช้ 60 วินาที	พนักงานคนที่ 1 - ดูรูปจากไฟล์ภาพ (3 วินาที) - โปรแกรมช่วยประมวลผลภาพ (5 วินาที) พนักงานคนที่ 2 - ตรวจสอบความถูกต้อง (10 วินาที) - บันทึกข้อมูล (1 วินาที) รวมเวลาที่ใช้ 19 วินาที	ระยะเวลาที่ใช้ลดลง 41 วินาที คิดเป็น 68.33% จากวิธีการเดิม ขั้นตอนการทำงานลดลง 1 ขั้นตอน แต่มี 2 ขั้นตอนที่เป็นนัยทางประสิทธิภาพการทำงานคือลดการขนย้ายกล่อง แล้วใช้วิธีการดูจากไฟล์ภาพแทน ซึ่งจะช่วยลดพื้นที่การทำงานด้วย

ตารางที่ 4.6 ผลการทำงานแบบเดิมเทียบกับการใช้โปรแกรมเข้ามาช่วยในการทำงาน (ต่อ)

ปัจจัย	วิธีการเดิม	วิธีการใหม่	ผลการเปรียบเทียบ
จำนวนคน	2 คน	1 คน	จำนวนคนลดลง 1 คน คิดเป็น 50% ของ จำนวนคนทั้งหมดที่ใช้ ในวิธีการเดิม
ต้นทุน (ค่าแรงขั้นต่ำ 340 บาทต่อคน)	ใช้ 2 คน เท่ากับ 680 บาทต่อวัน	ใช้ 1 คน เท่ากับ 340 บาทต่อวัน	ค่าใช้จ่ายลดลง 340 บาทต่อคนต่อวัน คิด เป็น 50% ของต้นทุนที่ ใช้ในวิธีการเดิม

จากการประเมินผลพบว่าระยะเวลาแต่ละขั้นตอนทำงาน การบันทึกข้อมูลดัชนีกล่องเอกสารใช้เวลาลดลง 41 วินาที คิดเป็น 68.33% จากระยะเวลาของวิธีการเดิม ในขั้นตอนการทำงานที่เป็น Cost ต่อเกี่ยวเนื่องคือการขนย้ายกล่องมายังจุดบันทึกข้อมูลจะไม่เกิดขึ้นและพื้นที่ในการจัดวางกล่อง ณ บริเวณที่ทำการบันทึกข้อมูลลดลง โดยใช้ภาพถ่ายที่เป็นไฟล์ดิจิทัลมาทำการบันทึกแทนในด้านจำนวนคนจะใช้คนลดลงจากเดิม 1 คน คิดเป็น 50% ของจำนวนคนจากวิธีการเดิม โดยวิธีการใหม่จะคน 1 คน เพื่อทำการตรวจสอบไฟล์ที่ได้จากการแปลงข้อมูลจากโปรแกรมเท่านั้น ซึ่งส่งผลให้ต้นทุนในการบันทึกดัชนีกล่องเอกสารลดลง 340 วันต่อวัน

บทที่ 5

สรุปและอภิปรายผลการวิจัย ข้อเสนอแนะ

การศึกษาวิจัยเรื่อง การเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือโดยงานวิจัยนี้มีวัตถุประสงค์เพื่อจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้ จะเน้นการแปลงข้อมูลลายมือที่ปรากฏอยู่บนกล่องเอกสารให้เป็นตัวอักษรที่สามารถใช้งานในคอมพิวเตอร์ได้ รวมไปถึงเพื่อเพิ่มประสิทธิภาพการทำงานของขั้นตอนการบันทึกข้อมูลให้สามารถลดต้นทุนด้านจำนวนคน ระยะเวลาและค่าใช้จ่ายอื่นที่เกี่ยวข้องได้ โดยมีบทสรุปผลการวิจัยดังต่อไปนี้

1. สรุปและอภิปรายผลการวิจัย
2. ปัญหาและอุปสรรคในการทำวิจัย
3. ข้อเสนอแนะ

5.1 สรุปและอภิปรายผลการวิจัย

ในการสรุปและอภิปรายผลการวิจัย ผู้วิจัยได้ทำการสรุปผลและอภิปรายตามวัตถุประสงค์ของงานวิจัยในแต่ละข้อ ดังนี้

1. การจัดทำต้นแบบสำหรับใช้พัฒนาโปรแกรมที่สามารถลดขั้นตอนการบันทึกข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร จะมีแยกขั้นตอนออกเป็น 2 ขั้นตอน ดังนี้

ขั้นตอนที่ 1 Training Data จะประกอบไปด้วยกระบวนการย่อยๆ ภายใน จะเริ่มจากการจัดเก็บข้อมูลลายมือจากกลุ่มตัวอย่างและทำการแปลงข้อมูลลายมือตัวอย่างให้อยู่ในรูปแบบไฟล์อิเล็กทรอนิกส์ จากนั้นจะแยกตัวอักษรแต่ละตัว โดยใช้ OCR Trainer สอนให้โปรแกรมนำรูปแบบตัวอักษรแต่ละตัว เพื่อเตรียมไว้สำหรับการทำ Pattern Matching ในขั้นตอน Human Intervention

ขั้นตอนที่ 2 Testing Data จะประกอบไปด้วยกระบวนการย่อยๆ ภายใน โดยเริ่มจากการนำภาพเข้ามาประมวลผลในโปรแกรม จากนั้นจะทำการปรับปรุงคุณภาพโดยเปลี่ยนข้อมูลภาพสีให้เป็นขาวดำ โดยใช้ Adaptive Threshold ในการลดแสงเงาที่ปรากฏบนกล่องเอกสาร จากนั้นจะทำการแยกส่วน (Segmentation) ของชุดข้อมูลที่นำเข้า โดยจะทำการแยกบรรทัด คำ และตัวอักษรตามลำดับ เพื่อเตรียมไว้สำหรับการทำ Pattern Matching กับข้อมูล โดยนำมาจับคู่และแปลงผลภาพตัวอักษรบนกล่องเอกสารให้เป็นตัวอักษรตามที่ได้จัดทำรูปแบบไว้ จากนั้นจะให้พนักงานทำหน้าที่ตรวจสอบความถูกต้องของข้อมูลที่แปลงออกมาได้ก่อน จากนั้นจึงจะบันทึกข้อมูลเก็บไว้ในระบบต่อไป

2. ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร มีเป้าหมายจะมีความถูกต้องของข้อมูลที่แปลงออกมาอย่างน้อย 80% จากการทดลองปรากฏว่าตัวอักษรที่ปรากฏบนกล่องทั้งหมดมี 13,160 ตัวอักษร สามารถแปลงข้อมูลได้ถูกต้อง 11,489 ตัวอักษร คิดเป็น 87.30% ซึ่งได้ผลลัพธ์สูงกว่าที่ตั้งไว้ 7.30%

3. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงาน โดยทำการประเมิน 3 ปัจจัย ได้แก่ ระยะเวลาแต่ละขั้นตอนการทำงาน จำนวนคนและค่าใช้จ่าย พบว่า ระยะเวลาที่ใช้ลดลง 68.33% จากวิธีการเดิม จำนวนคนในขั้นตอนการทำงานลดลง 50% ขั้นตอนการทำงานลดลง 1 ขั้นตอน คิดเป็น 20% ของขั้นตอนในวิธีการเดิม แต่มี 2 ขั้นตอนที่เป็นนัยทางประสิทธิภาพการทำงานคือลดขั้นตอนการขนย้ายกล่องมาวางเพื่อบันทึกข้อมูลและยกไปเก็บเมื่อบันทึกข้อมูลเสร็จ โดย 2 ขั้นตอนดังกล่าวถูกเปลี่ยนเป็นการใช้วิธีการดูจากไฟล์ภาพแทน ซึ่งจะช่วยลดพื้นที่การทำงานด้วย ในส่วนของค่าใช้จ่ายจะแปรผันตรงจากจำนวนคนที่ลดลงทำให้ลดค่าใช้จ่ายลง 50% ต่อวัน คิดเป็นเงิน 340 บาทต่อคนต่อวัน

การดำเนินงานวิจัยการเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือได้มีการจัดทำต้นแบบ (Prototype) เพื่อใช้ทดสอบตามข้อสมมติฐานที่ว่าสามารถแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสารได้มีความถูกต้อง (Accuracy) 80% และลดระยะเวลาที่ใช้ในขั้นตอนการบันทึกข้อมูลดัชนี (Index) เป็น 30% ของกระบวนการเดิม โดยผลการทดลองพบว่าสามารถแปลงข้อมูลได้ถูกต้อง 87.30% และสามารถลดระยะเวลาการบันทึกข้อมูลลง 68.33% จากกระบวนการเดิม เหตุผลที่ทำให้ผลการทดลองเป็นไปตามที่ตั้งสมมติฐานไว้ เนื่องจากตัวต้นแบบ (Prototype) ได้มีการจัดเตรียมข้อมูลของลายมือที่เป็นเฉพาะของผู้ที่ทำการบันทึกข้อมูลและได้มีการจัดทำคลังข้อมูลคำศัพท์ที่พบบ่อย (Dictionary) เพื่อให้ตัวต้นแบบ (Prototype) ทำการแปลงตัวอักษรได้มีความแม่นยำมากยิ่งขึ้น

5.2 ปัญหาและอุปสรรคในการวิจัย

ปัญหาและอุปสรรคในการวิจัยของการเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือมีดังนี้

5.2.1 ภาพกล่องเอกสารติดแสงและเงา

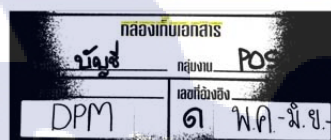
การที่ภาพถ่ายกล่องเอกสารมีแสงและเงาติดมา ดังรูปที่ 5.1 สาเหตุเกิดจากการถ่ายภาพในช่วงเวลาที่ต่างกันทำให้มีความแสงสว่างในแต่ละช่วงเวลาไม่เท่ากัน รวมไปถึงตำแหน่งเงาของคนที่ย้ายภาพไปบังแสงและทำให้เกิดเงาบนภาพกล่องเอกสาร จากสาเหตุข้างต้นสามารถทำการแก้ไขได้โดย Adaptive Threshold เข้ามาช่วยปรับปรุงคุณภาพของภาพเพื่อให้สามารถนำภาพไปใช้ใน

ขั้นตอนถัดไปได้ โดยจำนวนของกล่องที่นำเข้ามาทำการทดลองทั้งหมด 436 กล่อง โดยมีกล่องที่มีภาพติดแสงและเงาทั้งหมด 22 กล่อง กล่องเหล่านี้เมื่อได้ทำการปรับปรุงคุณภาพแล้วสามารถแปลงข้อมูลได้ถูกต้อง 81.25%

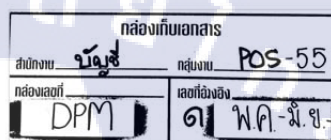
RGB Image (Color Image)



Black and White Image
(Auto Threshold)



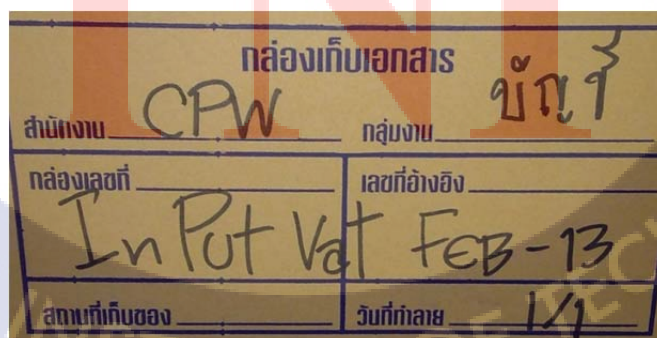
Black and White Image
(Adaptive Threshold)



รูปที่ 5.1 ภาพเปรียบเทียบกล่องเอกสารที่มีแสงเงาติดมาในภาพ

5.2.2 การเขียนตัวอักษรในแต่ละช่องไม่เป็นไปตามประเภทของช่องข้อมูลที่ระบุไว้

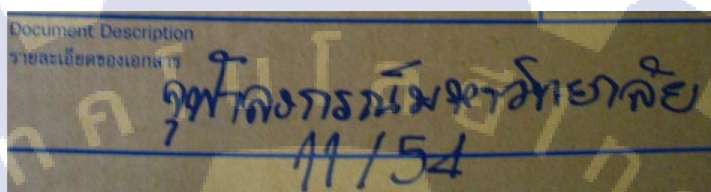
การเขียนตัวอักษรที่กล่องเอกสารไม่เป็นไปตามประเภทของช่องข้อมูลที่กำหนด ทำให้การแปลงข้อมูลเป็นตัวอักษรทำได้ยากขึ้น ดังรูปที่ 5.2 เพราะทำให้โปรแกรมไม่สามารถกำหนดรูปแบบที่เป็นตัวอักษรเฉพาะที่ต้องใช้ในตำแหน่งช่องนั้นๆ ได้ โดยเฉพาะช่องที่เป็นรายละเอียดของเอกสารจะมีการเขียนรายละเอียดเกินออกมาจากช่องที่กำหนดไว้ ทำให้เขียนเลยไปยังบริเวณพื้นที่ของช่องข้อมูลอื่น รวมไปถึงเขียนทับขอบของช่องข้อมูล วิธีการแก้ไขที่เป็นมาตรฐานทั่วไปจะต้องทำการกำหนดรหัสของประเภทเอกสารนั้นๆ ที่จะบรรจุลงกล่อง จะทำให้การแปลงผลทำได้ง่ายขึ้น



รูปที่ 5.2 ภาพกล่องเอกสารที่เขียนตัวอักษรในแต่ละช่องไม่เป็นไปตามประเภทที่ระบุไว้

5.2.3 เขียนตัวอักษรต่อเนื่องกัน

ตัวอย่างที่นำมาดำเนินการวิจัย พบว่ามีการเขียนตัวอักษรที่มีการลากเส้นปากกาเคมีต่อเนื่องกันของตัวอักษร ดังรูปที่ 5.3 ทำให้การแปลงผลตัวอักษรผิดไปจากรูปแบบที่กำหนดไว้ สาเหตุอาจเกิดจากผู้เขียนอาจเร่งรีบหรือเป็นลายมือปกติของผู้เขียนเอง วิธีการแก้ไข ควรให้ผู้เขียนเขียนตัวอักษรแต่ละตัวให้มีอิสระจากกัน เพื่อให้สามารถแปลงผลได้ง่ายขึ้น ในส่วนของโปรแกรมจะต้องนำข้อมูลลายมือที่มีลักษณะดังกล่าวไปทำการสอนในระบบให้สามารถรู้จำรูปแบบที่เป็นลักษณะแบบนี้ได้



รูปที่ 5.3 ภาพกล่องเอกสารที่เขียนตัวอักษรต่อเนื่องกัน

5.2.4 ตัวอักษรบางตัวมีการเขียนใกล้เคียงกัน

ลายมือของผู้เขียนแต่ละคนจะมีเอกลักษณ์เฉพาะตัวบุคคลทำให้ตัวอักษรบางตัวมีความคล้ายคลึงกัน เช่น ข กับ บ เป็นต้น ทำให้การแปลงตัวอักษรได้ผลไม่ตรงกับข้อมูลที่ปรากฏบนกล่องเอกสาร

5.3 ข้อเสนอแนะ

1. การจัดทำดัชนีบนกล่องเอกสารด้วยเทคนิคการรู้จำอักขระภาพลายมือได้ทำการวิจัยเป็นข้อมูลที่มีความเฉพาะและข้อจำกัดในเรื่องของลายมือที่นำมาทำตัวต้นแบบการรู้จำลายมือ (Training Data) ทำให้กล่องเอกสารใหม่ๆ ที่มีลายมือแตกต่างจากตัวต้นแบบที่ได้จัดทำไว้ จะทำให้การแปลงตัวอักษรมีความผิดพลาดได้ ดังนั้นโปรแกรมจะต้องมีการทำการรู้จำแบบของลายมือของแต่ละแบบไว้เพื่อใช้เป็นข้อมูลในการเปรียบเทียบเพื่อใช้แปลงผลข้อมูลให้มีความแม่นยำมากขึ้น

2. ขั้นตอนทำการเลือกบริเวณที่จะให้รู้จำตัวอักษรในแต่ละช่อง ควรจะต้องให้โปรแกรมสามารถกำหนดบริเวณที่จะทำการรู้จำตัวอักษรได้โดยอัตโนมัติ เพราะจะทำให้กระบวนการทำได้รวดเร็วมากยิ่งขึ้นมากกว่าการใช้คนเลือกบริเวณที่จะทำการรู้จำตัวอักษรซึ่งมีข้อผิดพลาดเกิดขึ้น แต่อย่างไรก็ตามยังต้องใช้คนในการตรวจเช็คข้อมูลที่ได้จากการรู้จำตัวอักษรได้โดยอัตโนมัติก่อนบันทึกลงฐานข้อมูล

3. ขั้นตอนหลังจากโปรแกรมได้ทำการรู้จำแล้ว ควรจะต้องทำให้การโปรแกรมสามารถทำการเรียนรู้ข้อมูลลายมือที่แตกต่างจากข้อมูล Training Data ที่ทำได้

4. การเขียนข้อมูลดัชนีข้างกล่องเอกสารควรจะต้องเขียนให้อยู่ภายในช่องที่กำหนดไว้โดยไม่เขียนทับบริเวณกรอบของแต่ละช่องและควรต้องเขียนเป็นลักษณะคำสำคัญหรือรหัส เพื่อที่จะให้สามารถลบกรอบออกได้ โดยตัวอักษรที่เขียนแต่ละตัวจะต้องเขียนให้แยกจากกันชัดเจนไม่ติดและต้องไม่ลากเส้นของตัวอักษรแต่ละตัวต่อเนื่องกัน เพื่อให้การรู้จำตัวอักษรทำได้อย่างมีประสิทธิภาพมากขึ้น

5. การถ่ายรูปตัวกล่องเอกสารที่จะนำมาใช้เป็น Testing Data ควรจะต้องมีการควบคุมเรื่องต่างๆ ดังนี้ แสงเงา ตำแหน่งการวางกล่องเอกสารต้องคงที่ ระยะห่างของกล่องกับกล้องถ่ายรูปต้องคงที่ องศาที่ถ่ายและตำแหน่งของกล้องถ่ายรูป เพื่อให้ทำให้ในการรู้จำตัวอักษรที่ปรากฏอยู่บนกล่องมีประสิทธิภาพ เพราะถ้าหากไม่มีการควบคุมจะส่งผลทำให้กระบวนการประมวลผลยากขึ้น

6. ข้อมูลที่ได้ทำการวิจัยเป็นเพียงต้นแบบเพื่อให้สามารถนำไปจัดทำโปรแกรมได้จริง จึงควรมีการนำข้อมูลที่ได้จากการวิจัยไปดำเนินการจัดทำเป็นโปรแกรมเพื่อทดสอบการแปลงข้อมูลแบบอัตโนมัติได้

TNI

THAI - NICHI INSTITUTE OF TECHNOLOGY

บรรณานุกรม

TNI

บรรณานุกรม

- [1] วิเชษฐ์รจน์ เอี่ยมสำอางค์ และรัชฎา คงคะจันทร์, “การแยกภาพตัวอักษรลายมือเขียนภาษาไทยแบบอัตโนมัติ โดยใช้การวิเคราะห์ถดถอยเชิงเส้นและการเรียนรู้ด้วยต้นไม้ตัดสินใจ,” *Thai Journal of Science and Technology*, ปีที่ 2, ฉบับที่ 2, หน้า 166-174, พฤษภาคม-สิงหาคม 2556.
- [2] พาฝัน ดวงไพศาล และคณะ, “การจดจำการเขียนลายมือตัวเลขไทยด้วยโครงข่ายการส่งข้อมูลแบบไม่ย้อนกลับ,” *Joint Conference on Applied Computer Technology and Information Systems & The National Conference on Business Administrator 2015*, ACTIS & NCOBA 2015, Nakorn Phanom, Thailand, January 30-31, 2015, pp. 293-297.
- [3] จีระเดช ราชไพบูลย์ และนุชนาฏ สัตยาภิ, “การรู้จำตัวเลขอารบิกจากลายมือสำหรับบรรทัดผู้สอบบนกระดาษคำตอบแบบปรนัย,” *วิทยานิพนธ์ วศ.ม. (วิศวกรรมคอมพิวเตอร์)*, มหาวิทยาลัยเกษตรศาสตร์, กรุงเทพฯ, 2556.
- [4] อมรเทพ วิจิตรเจริญ และพฤษดี ศิริแสงตระกูล, “การสกัดคุณลักษณะเด่นแบบผสมเพื่อรู้จำตัวเขียนอักษรธรรมอีสานด้วยโครงข่ายประสาทเทียม,” *4th Conference on Knowledge and Smart Technology, KST*, ชลบุรี, 7-8 กรกฎาคม 2555, หน้า 45-52.
- [5] V. S. Tapkir and S. D. Shelke, “OCR for handwritten marathi script,” *International Journal of Scientific & Engineering Research*, vol. 3, no. 8, pp. 1-6, August 2012.
- [6] M. Nayak and A. K. Nayak, “Odia characters recognition by training tesseract OCR engine,” *International Conference in Distributed in Distributed Computing & Internet Technology, ICDCIT-2014*, Bhubaneswar, India, February 6-9, 2014, pp. 25-30.
- [7] R. K. Bawa and G. K. Sethi, “A binarization technique for extraction of devanagari text from camera based images,” *Signal & Image Processing : An International Journal (SIPIJ)*, vol. 5, no. 2, pp. 29-37, April 2014.
- [8] อุษณีย์ สังฆธรรม, “การรู้จำตัวอักษรพิมพ์ไทยโดยใช้วิธีคอนดิชันนัลแรนดอมฟิวส์และระยะเซนทรอยด์แบบลำดับชั้น,” *วิทยานิพนธ์ วท.ม. (สถิติประยุกต์และเทคโนโลยีสารสนเทศ)*, สถาบันบัณฑิตพัฒนบริหารศาสตร์, กรุงเทพฯ, 2556.
- [9] ประสิทธิ์ บุญเอนก และอาทิตย์ ศรีแก้ว, “การจดจำลายมือเขียนตัวอักษรไทยด้วยแผนผังคุณลักษณะจัดการ,” *วิทยานิพนธ์ วศ.ม. (วิศวกรรมแมคคาทรอนิกส์)*, มหาวิทยาลัยเทคโนโลยีสุรนารี, นครราชสีมา, 2551.

- [10] S. P. Godara and P. S. Patwal, "Latin script detection and removal from devanagari document image for OCR," *International Journal of Computer & Organization Trends*, vol. 6, no. 1, pp. 33-36, March 2014.
- [11] ศุภรัตน์ อภิวงค์โสภณ, *การรู้จำตัวอักษรภาษาไทยที่เขียนด้วยลายมือในแบบฟอร์ม*, กรุงเทพฯ : สำนักพิมพ์จุฬาลงกรณ์มหาวิทยาลัย, 2550.
- [12] บุญธรรม ภัทราจารกุล, *การประมวลผลภาพดิจิทัลเบื้องต้น Fundamentals of Digital Image Processing*, กรุงเทพฯ : ซีเอ็ดดูเคชั่น, 2556.
- [13] Wikipedia, "การรู้จำแบบ," [Online]. Available from : <http://th.wikipedia.org/wiki/การรู้จำแบบ>. [Accessed : May 9, 2019].
- [14] Wikipedia, "การรู้จำอักขระด้วยแสง," [Online]. Available from : <https://th.wikipedia.org/wiki/การรู้จำอักขระด้วยแสง>. [Accessed : May 9, 2019].
- [15] Google Open Source, "Tesseract OCR," [Online]. Available from : <https://opensource.google.com/projects/tesseract>. [Accessed : May 9, 2019].
- [16] MathWorks, "Train Optical Character Recognition for Custom Fonts," [Online]. Available from : <https://www.mathworks.cn/help/vision/ug/train-optical-character-recognition-for-custom-fonts.html>. [Accessed : May 9, 2019].

ภาคผนวก

TNI

ภาคผนวก ก.

ตัวอย่างข้อมูลลายมือที่จัดเก็บเพื่อทำ Training Data

TNI

ตัวอย่างข้อมูลลายมือที่จัดเก็บเพื่อทำ Training Data

การจัดเก็บข้อมูลลายมือเพื่อใช้สำหรับทำ Training Data ได้ดำเนินการจัดเก็บข้อมูลจากกลุ่มตัวอย่างผู้ปฏิบัติงานจำนวน 15 คน โดยให้ผู้ปฏิบัติงานเขียนด้วยลายมือของตนด้วยปากกาเคมีสีน้ำเงิน โดยมีรูปแบบของอักษรที่ทำการจัดเก็บ จำนวน 7 รูปแบบ ดังนี้

1. ตัวอักษรภาษาไทยที่เป็นพยัญชนะ จำนวน 44 ตัวอักษร ดังรูปที่ ก.1

ตัวอักษร	ลายมือ
ก	ก
ข	ข
ฃ	ฃ
ค	ค
ฅ	ฅ
ฉ	ฉ

รูปที่ ก.1 ตัวอย่างการเขียนตัวอักษรภาษาไทยที่เป็นพยัญชนะ

2. ตัวอักษรภาษาไทยที่เป็นสระและวรรณยุกต์ จำนวน 25 ตัวอักษร ดังรูปที่ ก.2

ตัวอักษร	ลายมือ	ตัวอักษร	ลายมือ
อ	อ	ฤ	ฤ
อิ	อิ	ฤา	ฤา
อุ	อุ	ฤ	ฤ
อู	อู	ฤา	ฤา
เอะ	เอะ	ไม้เอก	เ
เอ	เอ	ไม้โท	๗

รูปที่ ก.2 ตัวอย่างการเขียนตัวอักษรภาษาไทยที่เป็นสระและวรรณยุกต์

3. ตัวอักษรภาษาไทยที่เป็นตัวเลขไทย จำนวน 10 ตัวอักษร ดังรูปที่ ก.3

ตัวอักษร	ลายมือ
๒	๒
๓	๓
๔	๔
๕	๕
๖	๖
๗	๗
๘	๘

รูปที่ ก.3 ตัวอย่างการเขียนตัวอักษรภาษาไทยที่เป็นตัวเลขไทย

4. ตัวอักษรภาษาอังกฤษตัวพิมพ์ใหญ่ จำนวน 26 ตัวอักษร ดังรูปที่ ก.4

ตัวอักษร	ลายมือ
A	A
B	B
C	C
D	D
E	E
F	F

รูปที่ ก.4 ตัวอย่างการเขียนตัวอักษรภาษาอังกฤษตัวพิมพ์ใหญ่

5. ตัวอักษรภาษาอังกฤษตัวพิมพ์เล็ก จำนวน 26 ตัวอักษร ดังรูปที่ ก.5

ตัวอักษร	ลายมือ
a	<i>a</i>
b	<i>b</i>
c	<i>c</i>
d	<i>d</i>
e	<i>e</i>
f	<i>f</i>

รูปที่ ก.5 ตัวอย่างการเขียนตัวอักษรภาษาอังกฤษตัวพิมพ์เล็ก

6. ตัวอักษรที่เป็นเครื่องหมาย จำนวน 24 ตัวอักษร ดังรูปที่ ก.6

ตัวอักษร	ลายมือ
[	<i>[</i>
>	<i>></i>
<	<i><</i>
{	<i>{</i>
}	<i>}</i>
/	<i>/</i>

รูปที่ ก.6 ตัวอย่างการเขียนตัวอักษรที่เป็นเครื่องหมาย

7. ตัวอักษรที่เป็นตัวเลขอารบิก จำนวน 10 ตัวอักษร ดังรูปที่ ก.7

ตัวอักษร	ลายมือ
5	5
6	6
7	7
8	8
9	9
0	0

ดังรูปที่ ก.7 ตัวอักษรที่เป็นตัวเลขอารบิก

ภาคผนวก ข.

ตัวอย่างข้อมูลกล่องเอกสารที่นำมาใช้ทำเป็น Testing Data

TNI

ตัวอย่างข้อมูลกล่องเอกสารที่นำมาใช้ทำเป็น Testing Data

การจัดเก็บตัวอย่างข้อมูลกล่องเอกสารที่นำมาใช้ทำเป็น Testing Data จะมาจากการถ่ายรูปด้วยกล้องดิจิทัล โดยมีจำนวน 436 กล่อง แบ่งเป็นกล่องที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า จำนวน 202 กล่อง และกล่องบริษัทที่ให้บริการจัดเก็บเอกสาร จำนวน 234 กล่อง ซึ่งกล่องแต่ละประเภทมีช่องรายละเอียดสำหรับเขียนข้อมูลแตกต่างกัน ดังนี้

กล่องที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า จำนวน 202 กล่อง มีช่องรายละเอียดสำหรับเขียนข้อมูล จำนวน 6 ช่อง ได้แก่ ช่องที่ 1 สำนักงาน ช่องที่ 2 กลุ่มงาน ช่องที่ 3 กล่องเลขที่ ช่องที่ 4 เลขที่อ้างอิง ช่องที่ 5 สถานที่เก็บของ และช่องที่ 6 วันที่ทำลาย ดังรูปที่ ข.1

กล่องเก็บเอกสาร (20๒๒)	
สำนักงาน <u>บัญชี</u>	กลุ่มงาน <u>POS</u>
กล่องเลขที่ <u>ดิจิทัล ๑๑</u>	เลขที่อ้างอิง <u>กันยายน/๕๕</u>
สถานที่เก็บของ	วันที่ทำลาย

รูปที่ ข.1 ภาพแสดงกล่องที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า

กล่องบริษัทที่ให้บริการจัดเก็บเอกสาร จำนวน 234 กล่อง มีช่องรายละเอียดสำหรับเขียนข้อมูล จำนวน 4 ช่อง ได้แก่ ช่องที่ 1 รหัสลูกค้า ช่องที่ 2 ชื่อกล่องลูกค้า ช่องที่ 3 ปี/เดือนทำลาย ช่องที่ 4 รายละเอียดกล่องของเอกสาร ดังรูปที่ ข.2

Barcode สำหรับ 000180G 10000001734	Customer Code รหัสลูกค้า 000180G
Reference Code ชื่อกล่องลูกค้า ส้มอ 07	Destruction Date ปี(ค.ศ.)/เดือน/ทำลาย ____/____/____
Document Description รายละเอียดของกล่องเอกสาร กล่อง 1	

รูปที่ ข.2 ภาพแสดงกล่องเอกสารของบริษัทที่ให้บริการจัดเก็บเอกสาร

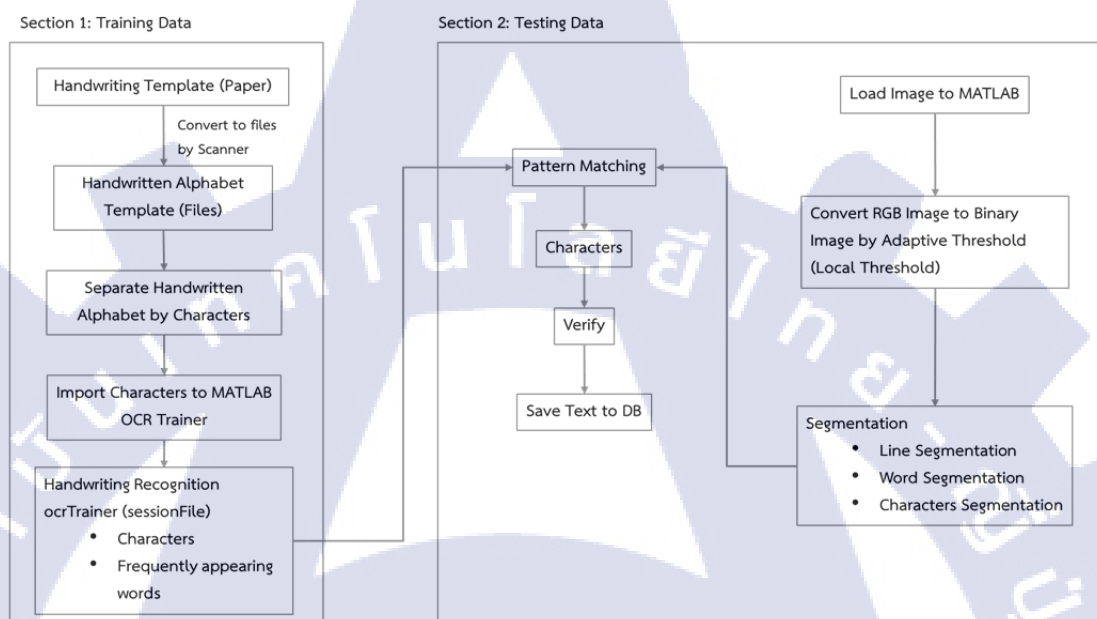
ภาคผนวก ค.

การเตรียมข้อมูลและวิธีการทดลอง

TNI

การเตรียมข้อมูลและวิธีการทดลอง

การเตรียมข้อมูลและวิธีการทดลองการเพิ่มประสิทธิภาพการจัดทำดัชนีบนกล่องเอกสาร ด้วยเทคนิคการรู้จำอักขระภาพลายมือจะแบ่งออกเป็น 2 ส่วน ได้แก่ ส่วนที่ 1 Training Data ส่วนที่ 2 Testing Data ดังรูปที่ ค.1



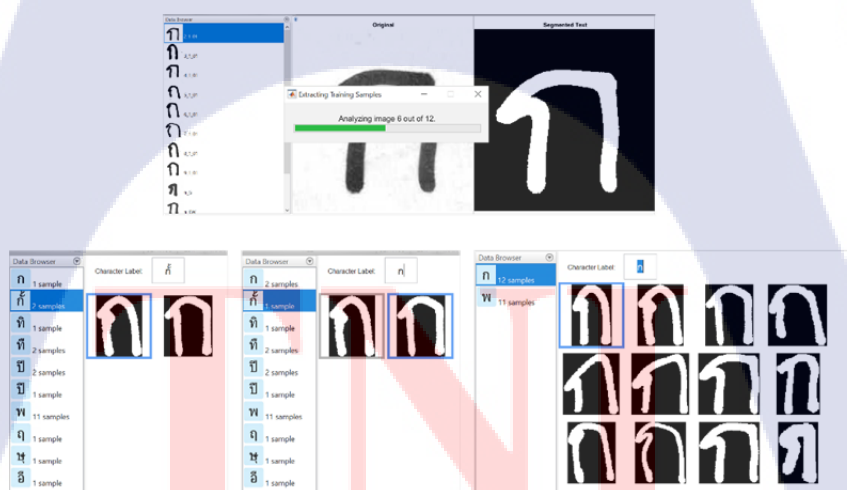
รูปที่ ค.1 ขั้นตอนการทดลองส่วนที่ 1 Training Data และส่วนที่ 2 Testing Data

ส่วนที่ 1 Training Data เป็นส่วนที่ไว้สำหรับเตรียมข้อมูลที่เป็นลายมือกับคำที่พบบ่อยบนกล่องเอกสาร กระบวนการจะเริ่มจากการให้ผู้ปฏิบัติงานเขียนด้วยลายมือของตนด้วยปากกาเคมีสีน้ำเงินลงบนแบบฟอร์มที่ได้จัดเตรียมไว้ จากนั้นจะนำลายมือที่เขียนไว้ไปสแกนด้วยเครื่องสแกนเนอร์และไปดำเนินการตัดภาพให้อยู่ในรูปแบบของตัวอักษรต่างๆ แล้วนำไปจัดเป็นกลุ่มของตัวอักษรนั้นๆ ไว้ ดังรูปที่ ค.2



รูปที่ ค.2 ตัวอย่างการจัดกลุ่มไฟล์ภาพตัวอักษร

เมื่อทำการจัดกลุ่มภาพตัวอักษรเสร็จเรียบร้อยแล้ว นำตัวอักษรเข้า MATLAB OCR Trainer เพื่อทำการ Train ตัวอักษรต้นแบบไว้ ส่วนของคำศัพท์ที่พบบ่อย จะใช้จากรูปกล่องเอกสาร ที่ทำการถ่ายไว้มานำเข้า MATLAB OCR Trainer แล้วจึงเลือกเฉพาะบริเวณของคำศัพท์ที่พบบ่อย นั้นๆ เพื่อทำการ Train เพื่อเตรียมไว้สำหรับการทำ Pattern Matching ดังรูปที่ ค.3 และรูปที่ ค.4

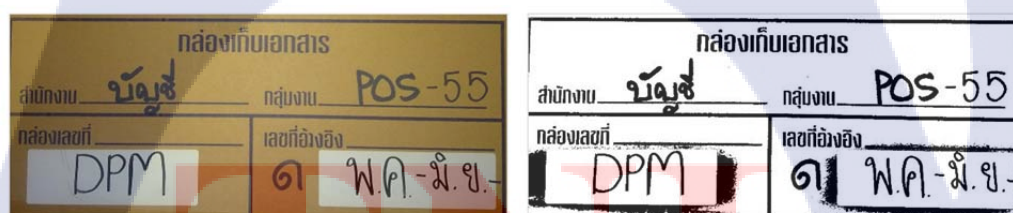


รูปที่ ค.3 การนำตัวอักษรเข้า MATLAB OCR Trainer



รูปที่ ค.4 การนำภาพของคำที่พบบ่อยเข้า MATLAB OCR Trainer

ส่วนที่ 2 Testing Data เป็นส่วนที่ใช้สำหรับการนำรูปภาพกล่องเอกสารเข้ามาทำการแปลงอักษรที่เขียนด้วยลายมือที่ปรากฏอยู่บนกล่องเอกสารให้อยู่ในรูปแบบตัวอักษรที่สามารถใช้งานบนคอมพิวเตอร์ได้ กระบวนการเริ่มจากนำรูปภาพสีกล่องเอกสารที่ได้จากกล้องถ่ายภาพดิจิทัล มาแปลงให้เป็นรูปภาพขาวดำโดยใช้ Adaptive Thresholding เพื่อทำการลดอิทธิพลของแสงเงาที่อยู่ในรูปภาพ ดังรูปที่ ค.5



รูปที่ ค.5 การแปลงรูปภาพสีเป็นรูปภาพขาวดำ โดยใช้ Adaptive Thresholding

หลังจากนั้นในขั้นตอนการทำ Segmentation จะใช้วิธีเลือก Region of Interest (roi) เป็นการเลือกบริเวณที่จะทำการแปลงข้อมูล ดังรูปที่ ค.6 เนื่องจากรูปภาพที่นำมาทำการวิจัยมี Perspective ที่ต่างกัน จากการทดลองใช้วิธีกำหนดบริเวณที่จะทำการแปลงแบบคงที่ จะพบว่ารูปภาพที่มี Perspective ไม่ตรงกับค่าที่ตั้งไว้จะมีผลการแปลงตัวอักษรผิดไปจากตัวอักษรที่ปรากฏอยู่บนกล่องเอกสาร เมื่อเลือก Region of Interest (roi) เรียบร้อยแล้ว จะทำการหา Boundary ของ

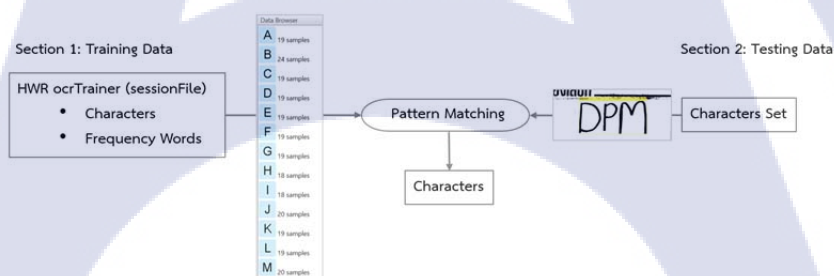
ตัวอักษรแต่ละตัวที่ทำการเลือกจาก Region of Interest (roi) เพื่อเตรียมนำตัวอักษรที่ได้ไปทำ Pattern Matching ในขั้นตอนถัดไป

กล่องเก็บเอกสาร	
สำเนาแบบ บัญชี	กลุ่มงาน POS-55
กล่องเลขที่ DPM	เลขที่อ้างอิง ๑ พ.ค.-ม.ย.

roi =
 250 584 459 208
 roi = Region of Interest

รูปที่ ค.6 การทำ Segmentation จะใช้วิธีเลือก Region of Interest (roi)

การทำ Pattern Matching จะเป็นการนำข้อมูลตัวอักษร Training Data ที่ได้เตรียมไว้มาเปรียบเทียบกับข้อมูลของตัวอักษร Testing Data ที่ได้เลือกไว้ในขั้นตอน Region of Interest (roi) เพื่อทำการแปลงข้อมูลตัวอักษรบนกล่องเอกสารให้เป็นตัวอักษรที่ใช้งานในคอมพิวเตอร์ ดังรูปที่ ค.7 และรูปที่ ค.8



รูปที่ ค.7 การทำ Pattern Matching

กล่องเก็บเอกสาร	
สำเนาแบบ บัญชี	กลุ่มงาน POS-55
กล่องเลขที่ DPM	เลขที่อ้างอิง ๑ พ.ค.-ม.ย.

รูปที่ ค.8 ผลการแปลงตัวอักษรที่ได้จากการทำ Pattern Matching

ภาคผนวก ง.

เอกสารที่นำเสนอวันสอบป้องกันสารนิพนธ์

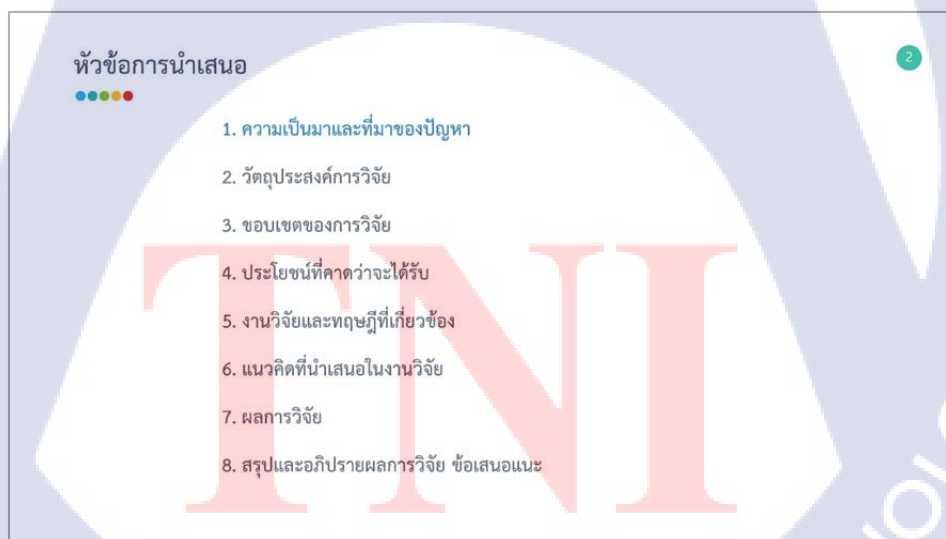
TNI

เอกสารที่นำเสนอวันสอบป้องกันสารนิพนธ์

การสอบป้องกันสารนิพนธ์ ได้จัดทำเอกสารนำเสนอสารนิพนธ์ โดยมีรายละเอียดตามรูปที่ ง.1 ถึง ง.55



รูปที่ ง.1 เอกสารนำเสนอสารนิพนธ์หน้าที่ 1



รูปที่ ง.2 เอกสารนำเสนอสารนิพนธ์หน้าที่ 2

ความเป็นมาและที่มาของปัญหา



ภาพการจัดเก็บเอกสารก่อนที่จะนำมาจัดเก็บกับบริษัท



ภาพการจัดเก็บเอกสารลงกล่อง

ขั้นตอนการจัดเก็บเอกสาร

- บรรจุนเอกสารลงกล่อง
- ติด Barcode กล่องเอกสาร
- เขียนดัชนี (Index) ข้างกล่อง

ที่มา บริษัท เซอร์วิส จำกัด

รูปที่ 3 เอกสารนำเสนอสารนิพนธ์หน้าที่ 3

ความเป็นมาและที่มาของปัญหา



ภาพการบันทึกข้อมูลกล่องเอกสาร




ภาพสถานที่จัดเก็บกล่องเอกสาร

เลข Barcode กล่อง	รหัสลูกค้า	ชื่อกล่องลูกค้า	ปี (ค.ศ.)/ เดือนทำสาย	รายละเอียดของ กล่องเอกสาร
10000001734	00018DG	สมอ 07		กล่อง 1

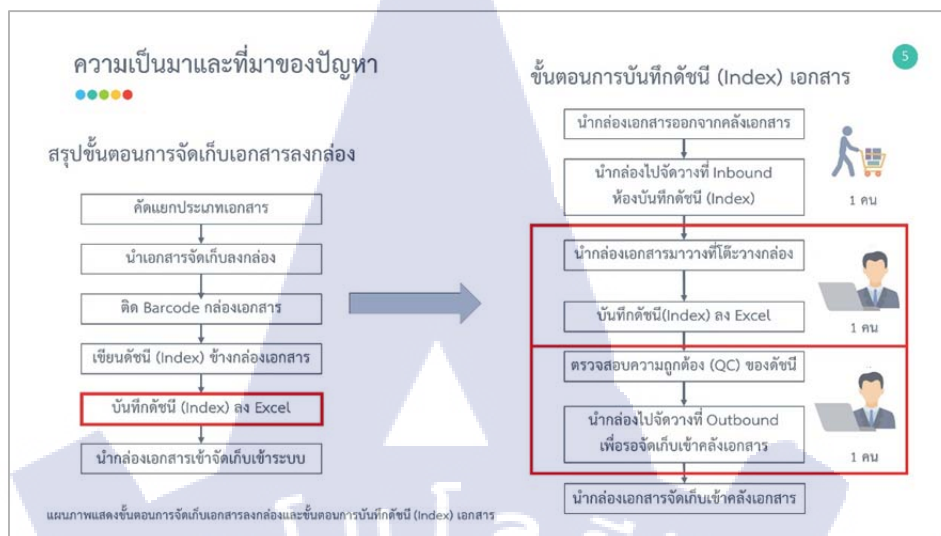
ตารางแสดงข้อมูลที่ทำการบันทึกจาก Index ข้างกล่องเอกสาร

ที่มา บริษัท เซอร์วิส จำกัด และบริษัท บลู พิค โซลูชั่น จำกัด

จัดเก็บกล่องเอกสารเข้าชั้นวางเอกสาร

นำข้อมูลที่บันทึกดัชนีเข้าระบบเพื่อใช้ค้นหา

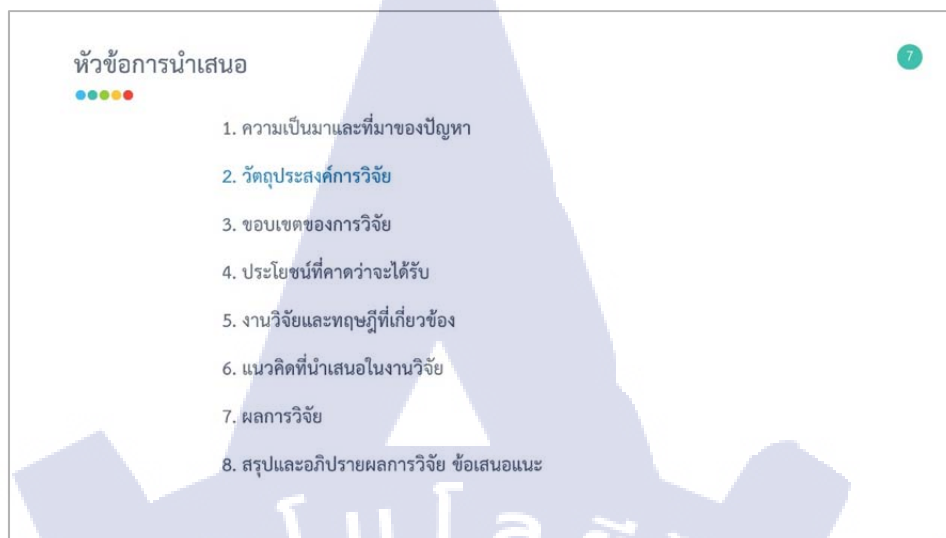
รูปที่ 4 เอกสารนำเสนอสารนิพนธ์หน้าที่ 4



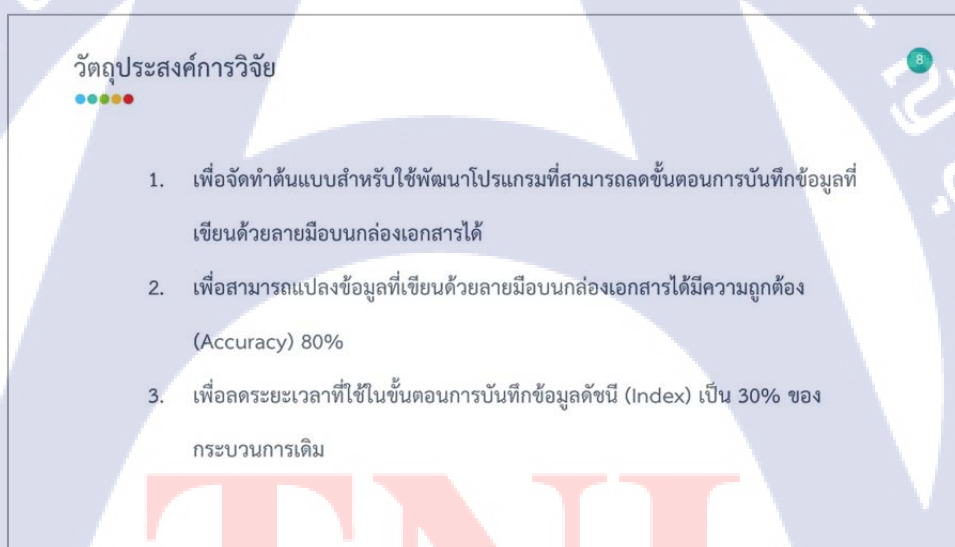
รูปที่ ง.5 เอกสารนำเสนอสารนิพนธ์หน้าที่ 5



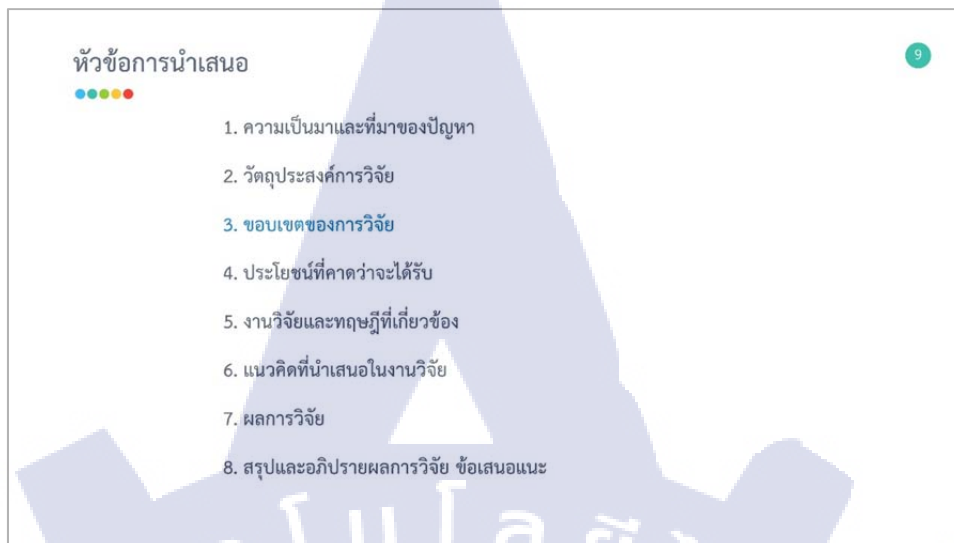
รูปที่ ง.6 เอกสารนำเสนอสารนิพนธ์หน้าที่ 6



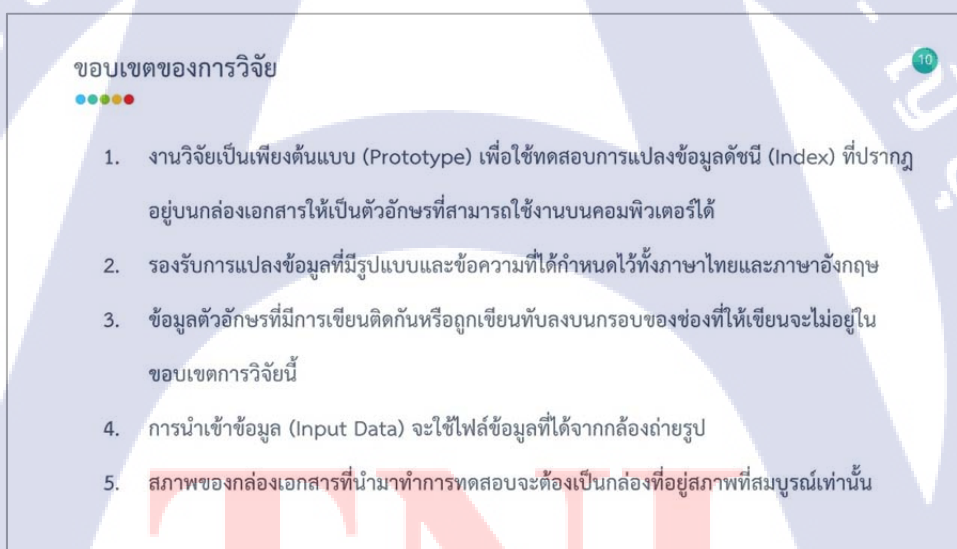
รูปที่ ง.7 เอกสารนำเสนอสารนิพนธ์หน้าที่ 7



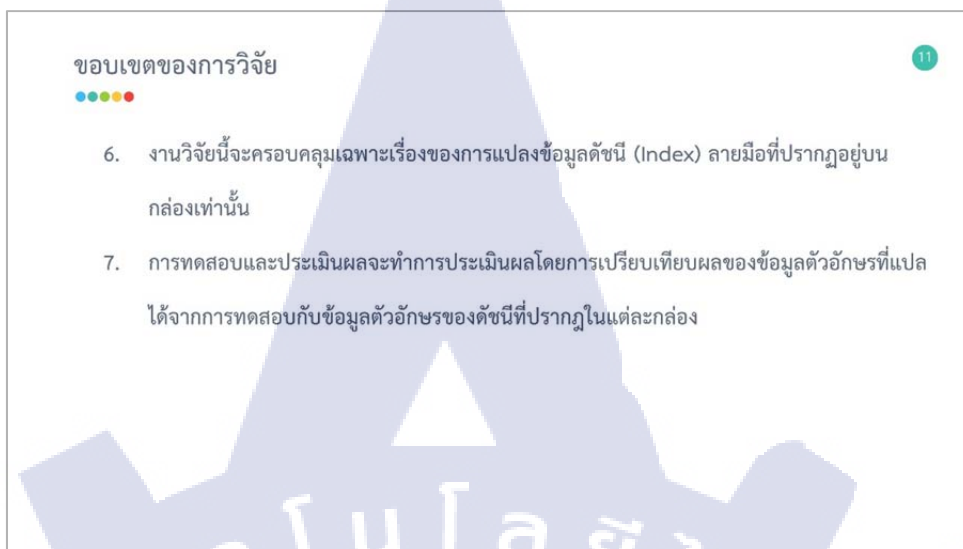
รูปที่ ง.8 เอกสารนำเสนอสารนิพนธ์หน้าที่ 8



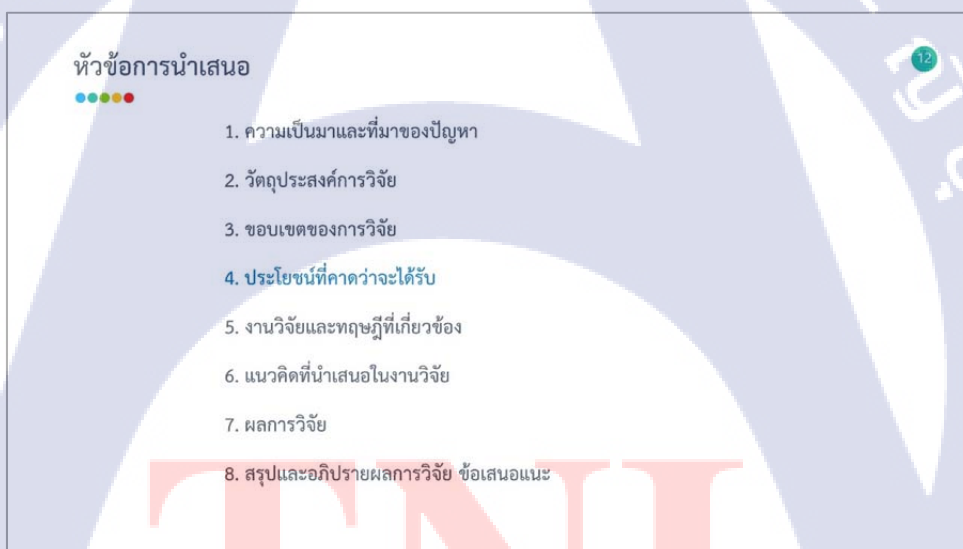
รูปที่ 9.9 เอกสารนำเสนอสารนิพนธ์หน้าที่ 9



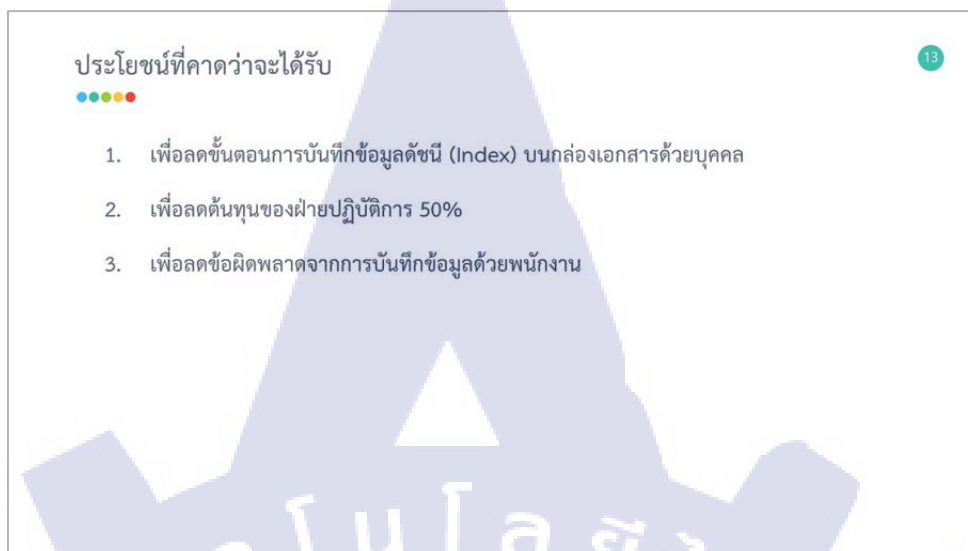
รูปที่ 9.10 เอกสารนำเสนอสารนิพนธ์หน้าที่ 10



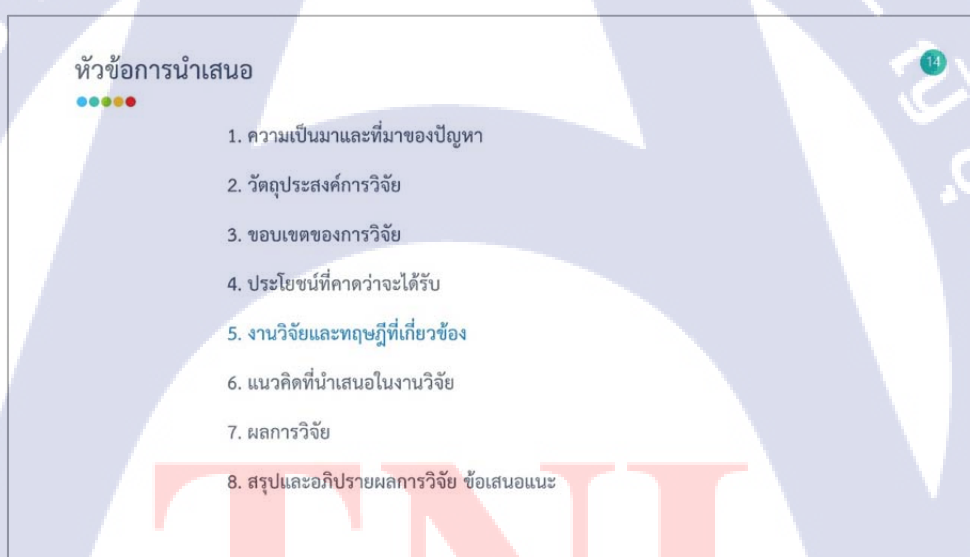
รูปที่ ง.11 เอกสารนำเสนอสารนิพนธ์หน้าที่ 11



รูปที่ ง.12 เอกสารนำเสนอสารนิพนธ์หน้าที่ 12



รูปที่ ง.13 เอกสารนำเสนอสารนิพนธ์หน้าที่ 13



รูปที่ ง.14 เอกสารนำเสนอสารนิพนธ์หน้าที่ 14

งานวิจัยที่เกี่ยวข้อง

15

ตารางสรุปงานวิจัยที่เกี่ยวข้อง

กลุ่ม	เรื่อง	วัตถุประสงค์	ภาษา	รูปแบบอักษรที่ทำ OCR	ความแม่นยำ
1 เน้นการทำ Pre- Processing	การแยกภาพตัวอักษรลายมือเขียนภาษาไทยแบบอัตโนมัติ โดยใช้การวิเคราะห์หาคออยเจินเส้นและการเรียนรู้ด้วยต้นไม้ตัดสินใจ (วิเศษฐรณ์ เอี่ยมสำอางค์ และรัชฎา คงกะจิตันท์)	- สามารถคิดแยกให้เป็นตัวอักษรเดี่ยวและตัวอักษรติดกัน - ในแนวนอนและแนวตั้งใช้เส้นการวิเคราะห์การถดถอย - สำหรับตัดแบ่งระดับกับระยะกับสระ - วิเคราะห์ตัวอักษรด้วยคุณลักษณะต่างๆ ของตัวอักษรไทย	ไทย	ลายมือ	90.44%
	การจดจำการเขียนลายมือตัวไทยด้วยโครงข่ายการส่งข้อมูลแบบไม่ย้อนกลับ (พวัฒน์ ดวงไพศาล, พนิดา พลอวงศ์ตระกูล, พงษ์ มีสัง)	- รู้จักตัวเลขภาษาไทยที่เขียนด้วยลายมือที่แตกต่างกันจากเดิมได้ - เรียนรู้ลักษณะ	ไทย	ลายมือ	87.86%
	การรู้จำตัวเลขอารบิกจากลายมือสำหรับระบบรหัสผู้สอบบนกระดาษคำตอบแบบปรนัย (ธีรเดช ราชไพบุลย์, นุชนาฏ สัตยาการ)	รู้จักตัวเลขของช่องรหัสผู้สอบในกระดาษคำตอบ	เลขอารบิก	ลายมือ	77%
	การสกัดคุณลักษณะเด่นแบบผสมเพื่อรู้จำตัวเขียนอักษรธรรมอีสานด้วยโครงข่ายประสาทเทียม (อมรเทพ วิจิตรเจริญ และ พงษ์ศักดิ์ ศิริแสงตระกูล)	รู้จักตัวเขียนอักษรธรรมอีสาน	อักษรอีสาน	ลายมือ	52-71%

รูปที่ ง.15 เอกสารนำเสนอสารนิพนธ์หน้าที่ 15

งานวิจัยที่เกี่ยวข้อง

ตารางสรุปงานวิจัยที่เกี่ยวข้อง

กลุ่ม	เรื่อง	วัตถุประสงค์	ภาษา	รูปแบบอักษรที่ทำ OCR	ความแม่นยำ
1 เน้นการทำ Pre-Processing	OCR For Handwritten Marathi Script (Mrs.Vinaya. S. Tapkir, Mrs.Sushma.D.Shelke)	• สกัดตัวอักษร (Features Extraction) ออกมาจากภาพโดยใช้เทคนิค Projection Histogram และ Zoning density features	Marathi	ลายมือ	Line Segmentation 100% Character Segmentation = 98%
	Latin Script Detection and Removal from Devanagari Document Image for OCR (Savita Pal Godara, Pratap Singh Patwal)	• เน้นการทำ Pre-Processing (Binarization, Dilation and Erosion, Skew Correction) ก่อนที่จะทำการสกัดเอาตัวอักษรที่อยู่ในรูปภาพออกมา	Latin Devanagari	ลายมือ	99.35%-99.80%
	A Binarization Technique For Extraction Of Devanagari Text From Camera Based Images (Rajesh K.Bawa and Ganesh K.Sethi)	• ใช้เทคนิค Binarization ในการสกัดตัวอักษรเพนนาหรือออกมาจากภาพ	Devanagari	ตัวพิมพ์	N/A
	การรู้จำตัวอักษรพิมพ์ไทยโดยใช้คอนดิชันนัลแนตอมพีวี่สและระยะเซนทรอยด์แบบลำดับชั้น (อุษณีย์ สังขธรรม)	• สกัดคุณลักษณะเด่น (Features Extraction) • การจำแนกตัวอักษร (classification) • เทคนิคการกระจายทิศทาง	ไทย	ลายมือ	96.96%

รูปที่ ง.16 เอกสารนำเสนอสารนิพนธ์หน้าที่ 16

งานวิจัยที่เกี่ยวข้อง

งานวิจัยแบ่งออกเป็น 2 กลุ่ม

กลุ่ม	เรื่อง	วัตถุประสงค์	ภาษา	รูปแบบอักษรที่ทำ OCR	ความแม่นยำ
2 Processing & Post- Processing	การจดจำลายมือเขียนตัวอักษรไทยด้วยแผนผังคุณลักษณะจัดการตัวเอง (ประสิทธิ์ บุญเอก และ มศ.ดร.อาทิตย์ ศรีแก้ว)	จดจำลายมือเขียนตัวอักษรไทยที่สามารถจดจำลายมือเขียนภาษาไทยที่มีลักษณะรูปร่างคล้ายกันไป	ไทย	ลายมือ	91.09%
	การรู้จำตัวอักษรภาษาไทยที่เขียนด้วยลายมือในแบบฟอร์ม (ศุภรัตน์ อภิวงค์โสภณ)	จำแนกลายมือจากการกรอกแบบฟอร์มโดยใช้ทำการแบ่งระดับข้อความ หาขอบภาพของแต่ละตัวอักษรเพื่อนำมาแบ่งครึ่งตัวอักษรแล้วมาเปรียบเทียบกับค่ารหัสลูกโซ่ 8 บิต	ไทย	ลายมือ	91%
	Odia Characters Recognition by Training Tesseract OCR Engine (Mamata Nayak, Ajit Kumar Nayak)	ใช้ Tesseract OCR Engine เพื่อรู้จำตัวอักษร	Odia	ลายมือ	98-100%

รูปที่ ง.17 เอกสารนำเสนอสารนิพนธ์หน้าที่ 17

สรุปจุดเด่นของงานวิจัยที่ได้ศึกษา

1. สามารถแยกตัวอักษรที่เขียนติดกันได้
2. สามารถจดจำลายมือใหม่ๆ หรือเขียนแตกต่างจากเดิมได้
3. ใช้เทคนิค Projection Profile ในการหากลุ่มของข้อมูลเพื่อสกัดออกเป็นตัวอักษรได้
4. สามารถแยกข้อมูลเป็นรายบรรทัด คำ และตัวอักษร
5. สามารถแปลงภาษาเดียวที่ตรงกับตัวอักษรที่กำหนดได้
6. มีการใช้เทคนิค Binarization ทำให้ตัวอักษรชัดเจนขึ้น
7. มีการรู้จำข้อมูลแล้วนำไปกรอกแบบฟอร์มให้อัตโนมัติ
8. ใช้ Tesseract ในการ Training Data ตัวอักษรก่อนที่จะทำการรู้จำ

รูปที่ ง.18 เอกสารนำเสนอสารนิพนธ์หน้าที่ 18

สรุปข้อจำกัดของงานวิจัยที่ได้ศึกษา



19

1. การเขียนตัวอักษรไม่มีหัวทำให้เกิดความคลาดเคลื่อน
2. เขียนตัวอักษรลากเส้นติดกันทำให้ประมวลผลผิดพลาด
3. มีสระปนอยู่ในตัวอักษรทำให้เกิดข้อผิดพลาดในการตัดตัวอักษร
4. ข้อมูลที่นำมาประมวลผลมีขนาดใหญ่ ทำให้ประมวลผลช้า
5. เอกสารที่มี 2 ภาษา จะสามารถรู้จำได้เฉพาะที่มีการกำหนดไว้
6. เอกสารต้นฉบับมีคุณภาพไม่ดีทำประสิทธิภาพลดลง

รูปที่ ง.19 เอกสารนำเสนอสารนิพนธ์หน้าที่ 19

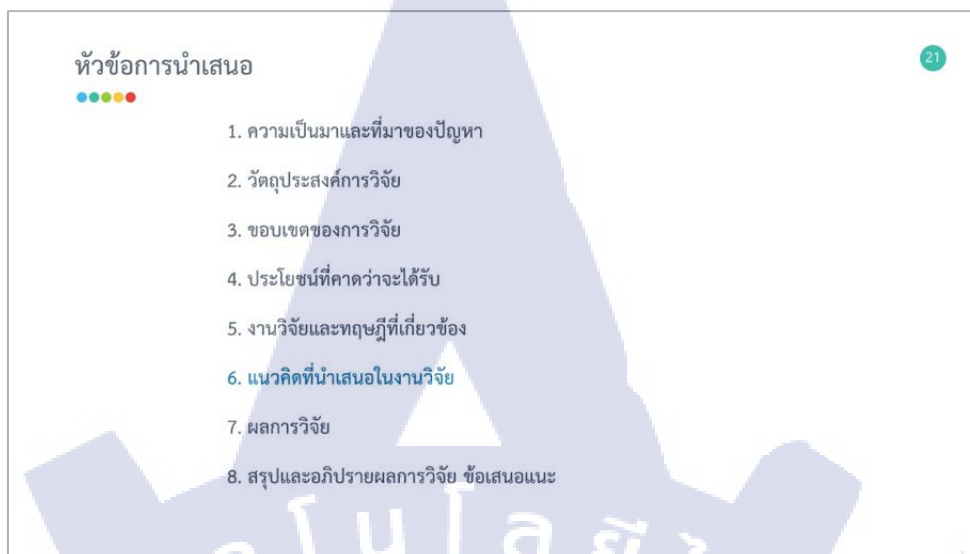
ทฤษฎีที่เกี่ยวข้องกับงานวิจัย



20

1. การประมวลผลภาพ (Digital Image Processing)
 - Binarization
 - Segmentation
2. Pattern Recognition
3. Optical Character Recognition
4. Tesseract OCR Engine
5. MATLAB OCR

รูปที่ ง.20 เอกสารนำเสนอสารนิพนธ์หน้าที่ 20



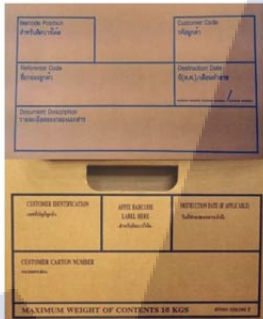
รูปที่ ง.21 เอกสารนำเสนอสารนิพนธ์หน้าที่ 21



รูปที่ ง.22 เอกสารนำเสนอสารนิพนธ์หน้าที่ 22

การวิเคราะห์รายละเอียดของกล่องเอกสาร

23



ตัวอย่างรูปแบบของกล่องที่จะจัดทำขึ้น

ที่มา บริษัท เจริญชัย จำกัด

ตารางแสดงดัชนีที่สื่อบันทึกข้อมูลของกล่องเอกสาร

ลำดับ	ข้อมูลดัชนีกล่องเอกสาร
1	ตำแหน่งติดบาร์โค้ด
2	รหัสลูกค้า
3	ชื่อกล่องลูกค้า
4	วันที่ทำลายเอกสาร
5	รายละเอียดกล่องเอกสาร

รูปที่ ง.23 เอกสารนำเสนอสารนิพนธ์หน้าที่ 23

การวิเคราะห์รายละเอียดของกล่องเอกสาร

24



ตัวอย่างลายมือที่เขียนข้างกล่องเอกสาร

ที่มา บริษัท เจริญชัย จำกัด

- เขียนด้วยปากกาเคมีที่มีขนาดใหญ่
- มองเห็นได้ชัดเจนในระยะห่าง 2-3 เมตร
- เขียนในรูปแบบที่เป็นคำสำคัญ (Keyword)

ใบสั่งซื้อ เดือนกันยายน 2558 จะนิยมเขียนเป็น
PO. กย. 2558 หรือ PO.-SEP 2015

- ลายมือจะเป็นลายมือของเจ้าหน้าที่คนเดิม
- ตัวอักษรที่เขียนมีคุณลักษณะต่างกัน

รูปที่ ง.24 เอกสารนำเสนอสารนิพนธ์หน้าที่ 24

การวิเคราะห์รายละเอียดของกล่องเอกสาร

ตัวอย่างลายมือที่เขียนข้างกล่องเอกสาร

ตัวอย่างลายมือที่เขียนข้างกล่องเอกสาร

ที่มา บริษัท เซลล์สัน จำกัด

การวางแสดงข้อมูลและรูปแบบการบันทึกข้อมูลของดัชนีกล่องเอกสาร

ลำดับ	ข้อมูลดัชนีกล่องเอกสาร	รูปแบบการบันทึกข้อมูล	ประเภทข้อมูล		
			ไทย	อังกฤษ	ตัวเลขและสัญลักษณ์
1	ตำแหน่งติดบาร์โค้ด	สติ๊กเกอร์บาร์โค้ด	/	/	/
2	รหัสลูกค้า	ลายมือเขียน	/	/	/
3	ชื่อกล่องลูกค้า	ลายมือเขียน	/	/	/
4	วันที่ทำลายเอกสาร*	ลายมือเขียน			/
5	รายละเอียดกล่องเอกสาร	ลายมือเขียน	/	/	/

หมายเหตุ
* เป็นข้อมูลที่ไม่มีบังคับ

รูปที่ ง.25 เอกสารนำเสนอสารนิพนธ์หน้าที่ 25

การวิเคราะห์รายละเอียดของกล่องเอกสาร

ตัวอย่างลายมือที่เขียนข้างกล่องเอกสาร

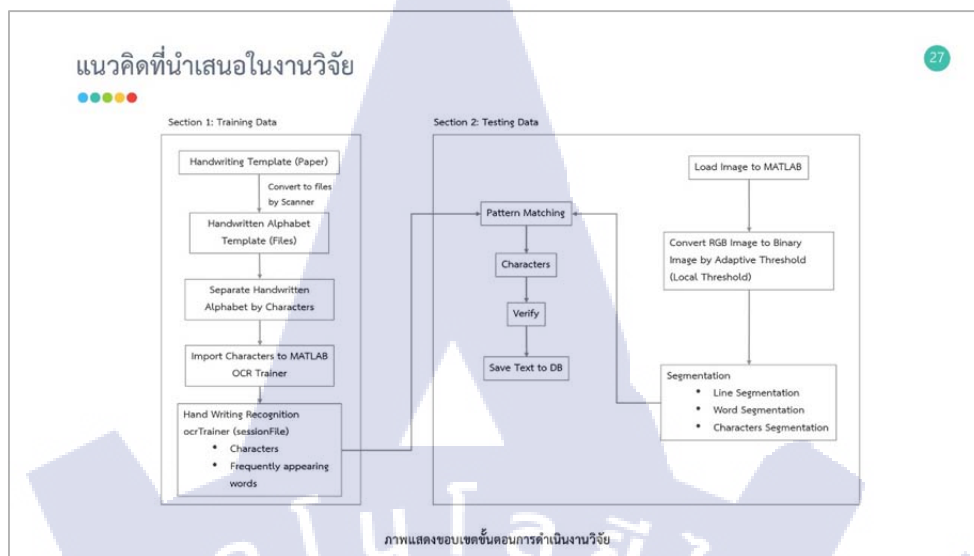
ตัวอย่างลายมือที่เขียนข้างกล่องเอกสาร

ที่มา บริษัท เซลล์สัน จำกัด

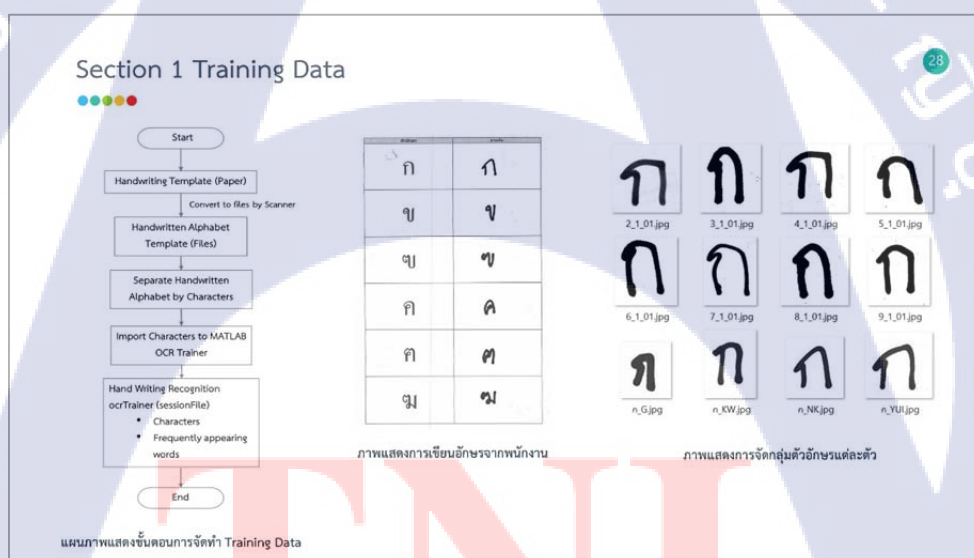
การวางแสดงข้อมูลและรูปแบบอักษรที่ปรากฏบนกล่องเอกสาร

แบบที่	รูปแบบข้อมูลที่ปรากฏบนกล่องเอกสาร	ประเภทข้อมูล		
		ไทย	อังกฤษ	ตัวเลขและสัญลักษณ์
1	ตัวอักษรภาษาไทย ตัวเลขอารบิกและสัญลักษณ์	/	/	/
2	ตัวอักษรภาษาอังกฤษ ตัวเลขอารบิกและสัญลักษณ์	/	/	/
3	ตัวอักษรภาษาไทย ตัวอักษรภาษาอังกฤษ ตัวเลขอารบิกและสัญลักษณ์	/	/	/

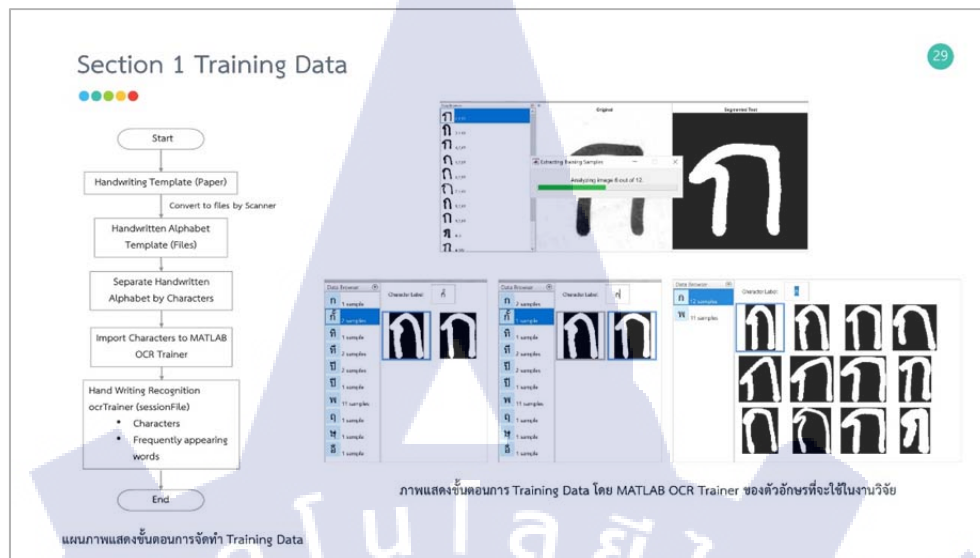
รูปที่ ง.26 เอกสารนำเสนอสารนิพนธ์หน้าที่ 26



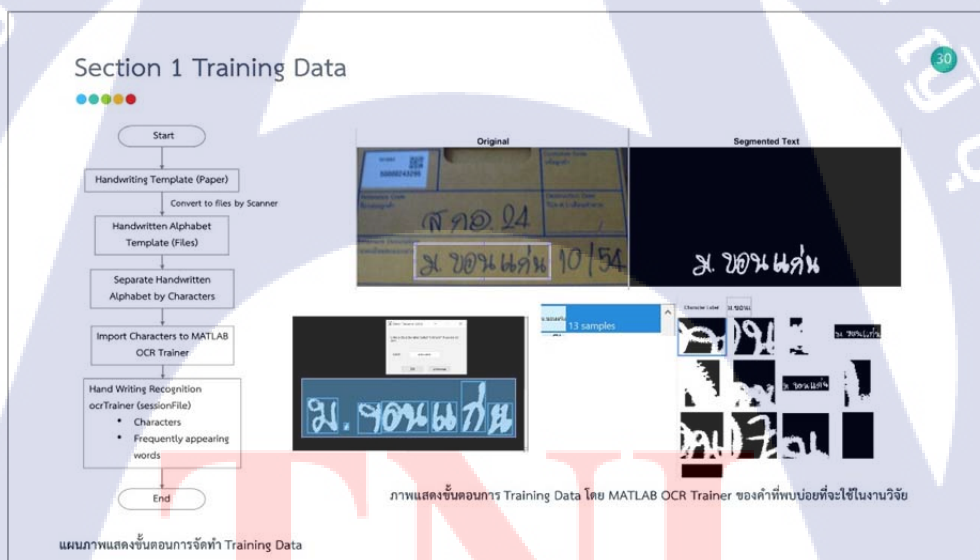
รูปที่ ง.27 เอกสารนำเสนอสารนิพนธ์หน้าที่ 27



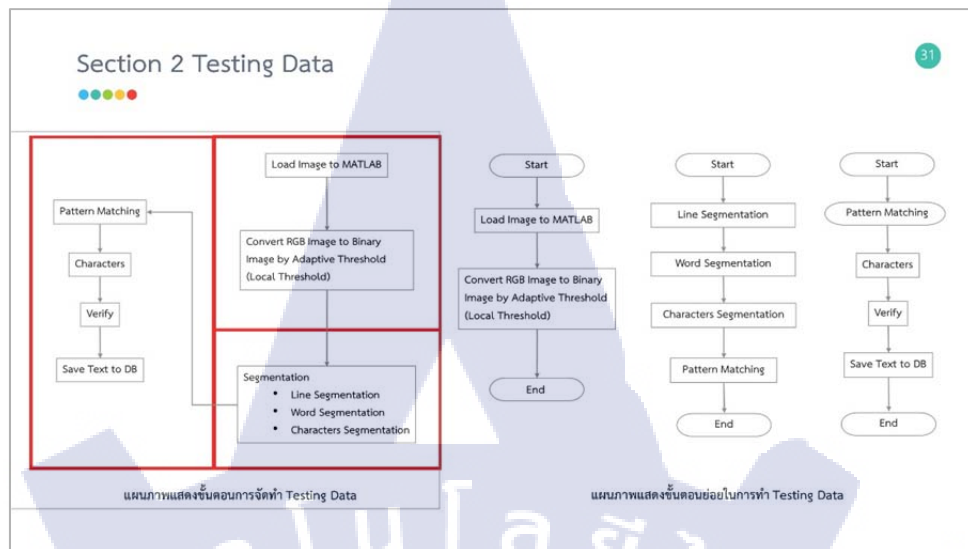
รูปที่ ง.28 เอกสารนำเสนอสารนิพนธ์หน้าที่ 28



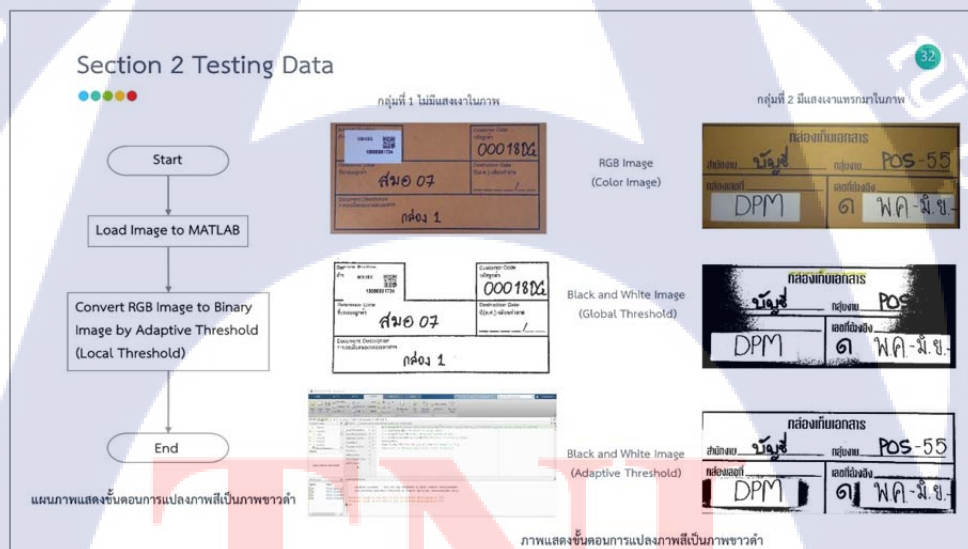
รูปที่ ง.29 เอกสารนำเสนอสารนิพนธ์หน้าที่ 29



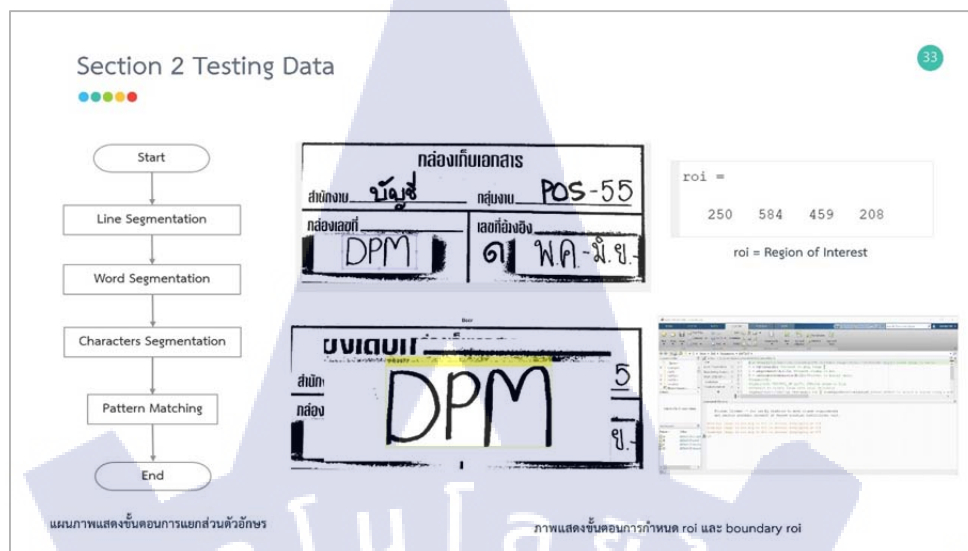
รูปที่ ง.30 เอกสารนำเสนอสารนิพนธ์หน้าที่ 30



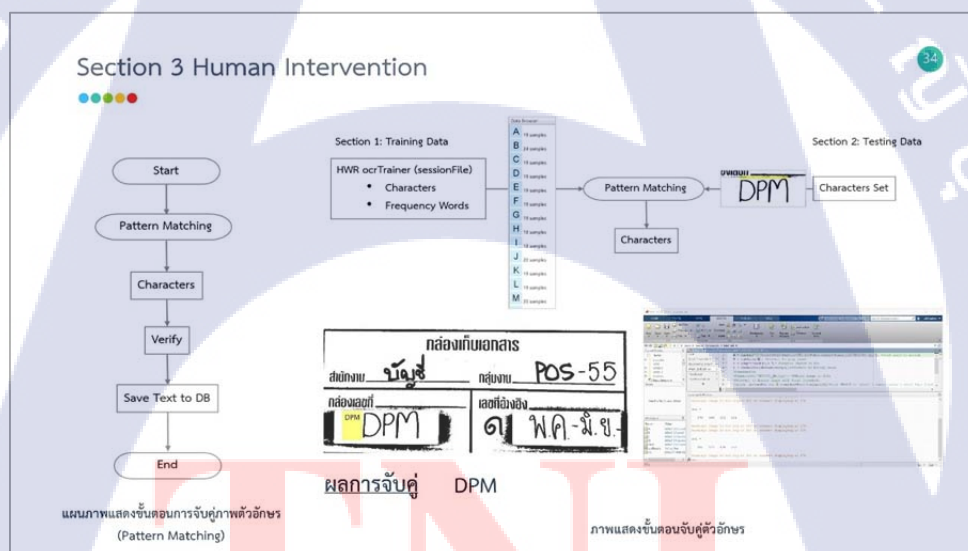
รูปที่ ง.31 เอกสารนำเสนอสารนิพนธ์หน้าที่ 31



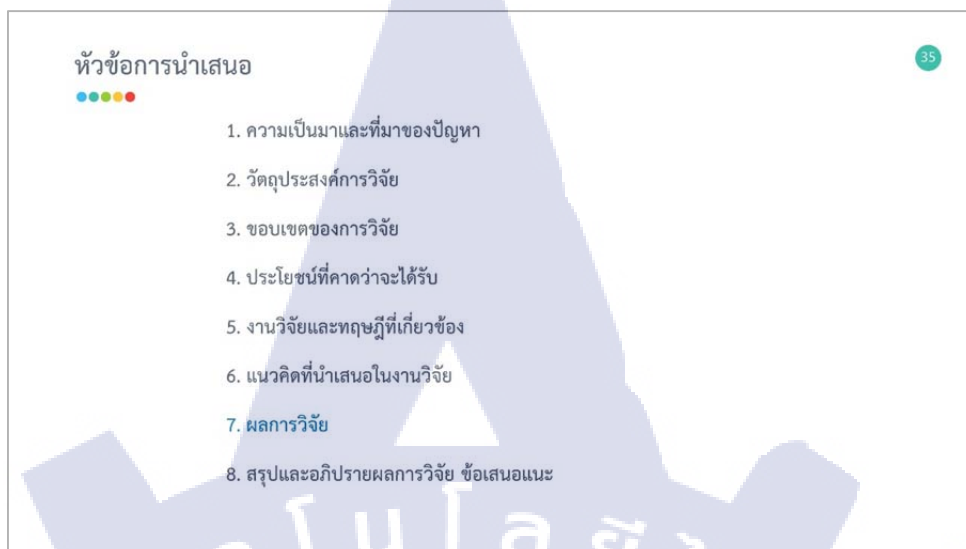
รูปที่ ง.32 เอกสารนำเสนอสารนิพนธ์หน้าที่ 32



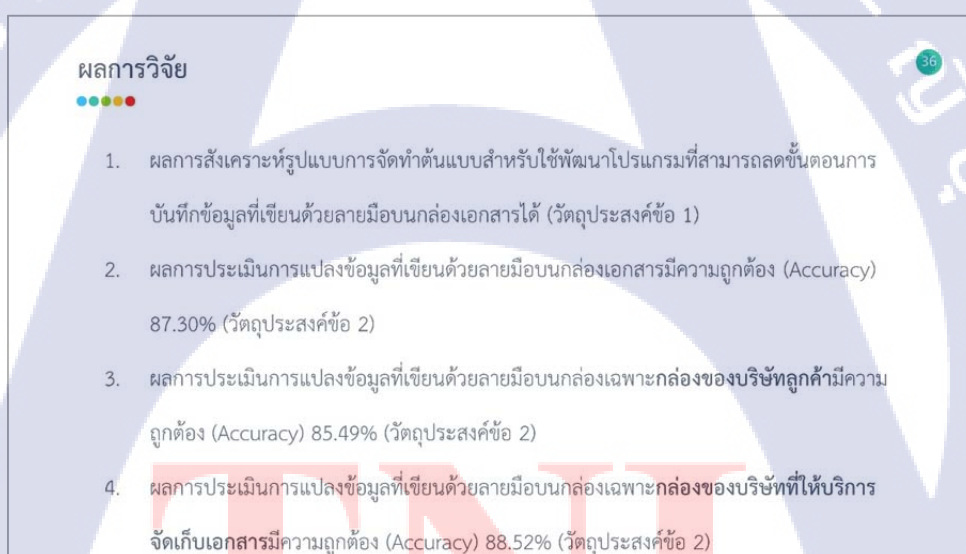
รูปที่ ง.33 เอกสารนำเสนอสารนิพนธ์หน้าที่ 33



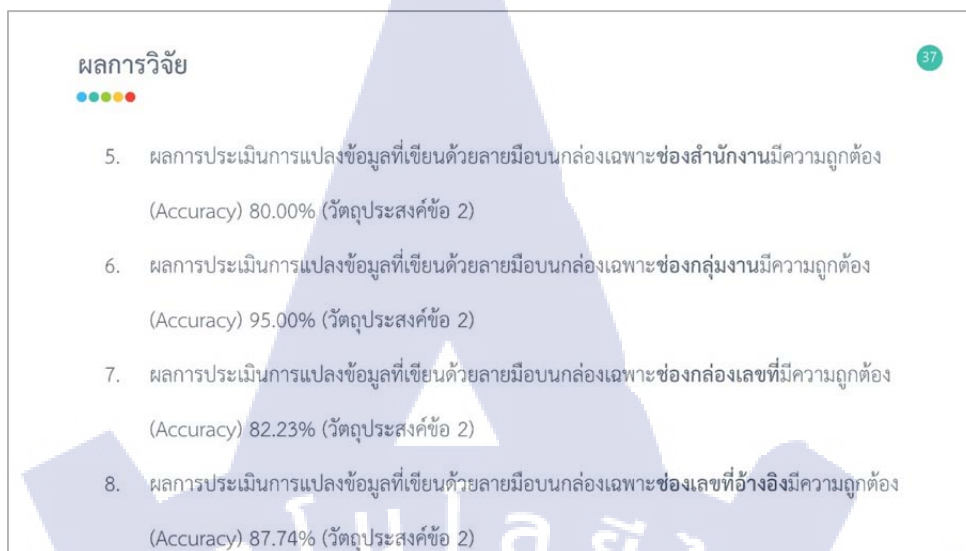
รูปที่ ง.34 เอกสารนำเสนอสารนิพนธ์หน้าที่ 34



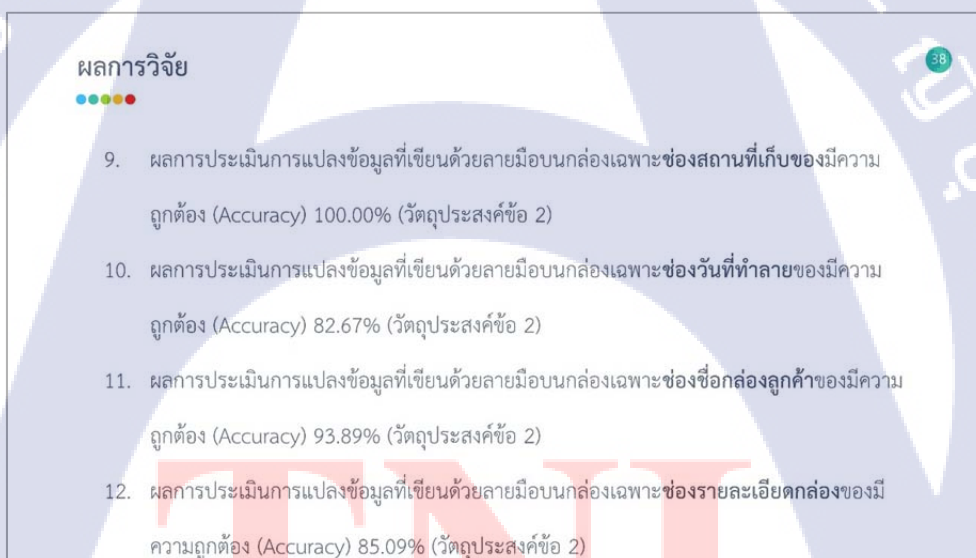
รูปที่ 3.35 เอกสารนำเสนอสารนิพนธ์หน้าที่ 35



รูปที่ 3.36 เอกสารนำเสนอสารนิพนธ์หน้าที่ 36



รูปที่ ง.37 เอกสารนำเสนอสารนิพนธ์หน้าที่ 37



รูปที่ ง.38 เอกสารนำเสนอสารนิพนธ์หน้าที่ 38

ผลการวิจัย



39

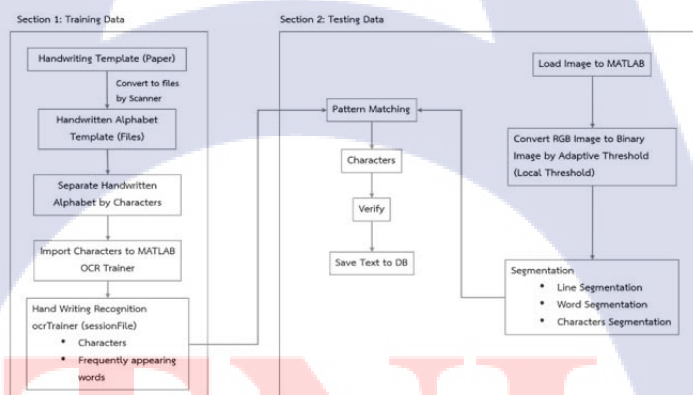
13. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงานด้าน
ระยะเวลาลดลง 68.33% (วัดอุปสงค์ข้อ 3)
14. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงานด้าน
จำนวนคนลดลง 50.00% (วัดอุปสงค์ข้อ 3)
15. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่เข้ามาช่วยในขั้นตอนการทำงานด้าน
ต้นทุนลดลง 50.00% (วัดอุปสงค์ข้อ 3)

รูปที่ ง.39 เอกสารนำเสนอสารนิพนธ์หน้าที่ 39

1. ผลการสังเคราะห์รูปแบบการจัดทำต้นแบบ



40



ภาพแสดงผลการสังเคราะห์รูปแบบการจัดทำต้นแบบ

รูปที่ ง.40 เอกสารนำเสนอสารนิพนธ์หน้าที่ 40

2. ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร



41

วิธีการวัดประสิทธิภาพและความถูกต้อง

1. วัดเป็นค่าร้อยละของความถูกต้องของชุดข้อมูลทดสอบ (Testing Data) กับชุดข้อมูลฝึกสอน (Training Data)
2. วัดค่าเป็นร้อยละของความถูกต้องของชุดข้อมูลทดสอบ (Testing Data) ในแต่ละช่องที่ปรากฏอยู่บนดัชนีกล่องเอกสารจากจำนวนข้อมูลทดสอบทั้งหมด
3. วัดเป็นค่าร้อยละของประสิทธิภาพของเวลาที่ใช้ในการบันทึกข้อมูลดัชนี (Index) ด้วยตัวต้นแบบกับการบันทึกข้อมูลดัชนี (Index) ด้วยวิธีเดิม

รูปที่ ง.41 เอกสารนำเสนอสารนิพนธ์หน้าที่ 41

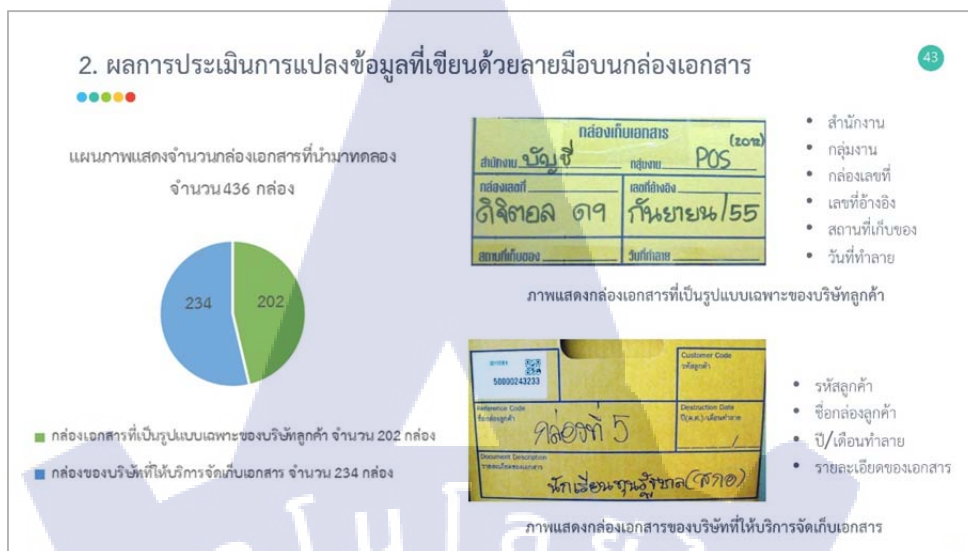
2. ผลการประเมินการแปลงข้อมูลที่เขียนด้วยลายมือบนกล่องเอกสาร



42

1. การประเมินผลความถูกต้องของตัวอักษรที่แปลงได้ทั้งหมด
 - กล่องเอกสารทั้งหมด
 - แยกตามประเภทกล่องเอกสาร
2. การประเมินผลความถูกต้องของแต่ละช่องของกล่องเอกสาร
 - กล่องเอกสารที่เป็นรูปแบบเฉพาะของบริษัทลูกค้า
 - กล่องเอกสารของบริษัทผู้ให้บริการ

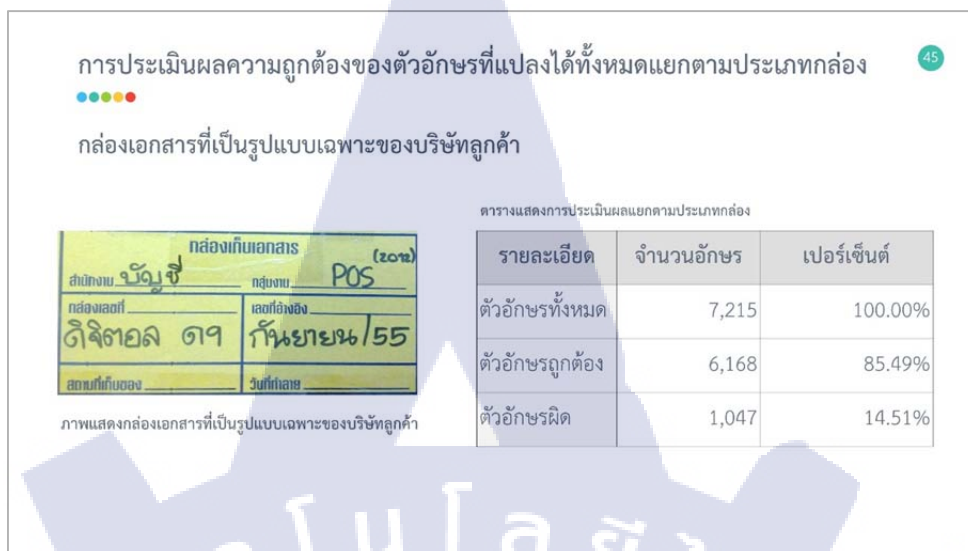
รูปที่ ง.42 เอกสารนำเสนอสารนิพนธ์หน้าที่ 42



รูปที่ ง.43 เอกสารนำเสนอสารนิพนธ์หน้าที่ 43



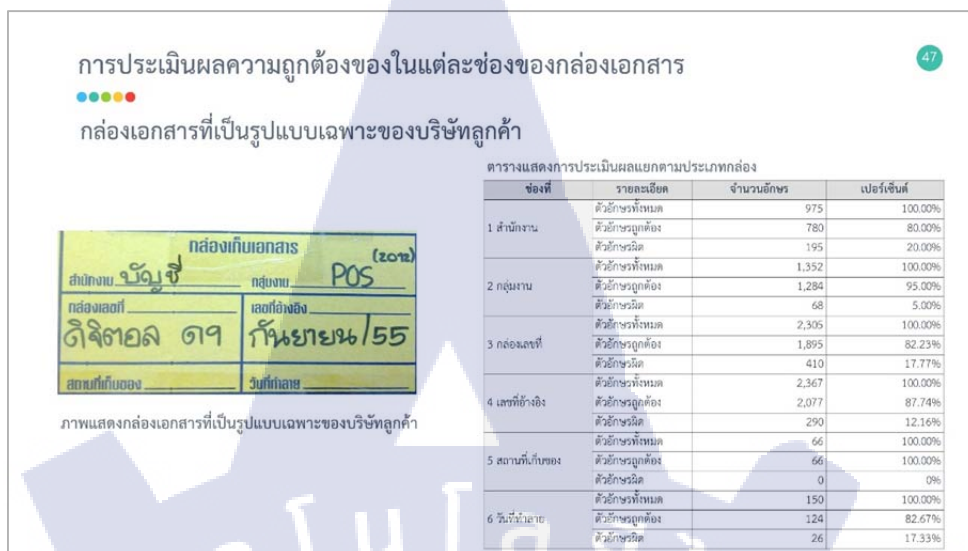
รูปที่ ง.44 เอกสารนำเสนอสารนิพนธ์หน้าที่ 44



รูปที่ ง.45 เอกสารนำเสนอสารนิพนธ์หน้าที่ 45



รูปที่ ง.46 เอกสารนำเสนอสารนิพนธ์หน้าที่ 46



รูปที่ ง.47 เอกสารนำเสนอสารนิพนธ์หน้าที่ 47



รูปที่ ง.48 เอกสารนำเสนอสารนิพนธ์หน้าที่ 48

3. ผลการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่



49

ตารางแสดงการประเมินการทำงานวิธีการเดิมเทียบกับวิธีการใหม่

ปัจจัย	วิธีการเดิม	วิธีการใหม่	ผลการเปรียบเทียบ
ระยะเวลาแต่ละขั้นตอนการทำงาน	พนักงานคนที่ 1 - นำกล้องมาวาง (10 วินาที) - บันทึกข้อมูลลายมือข้างล่าง (30 วินาที) พนักงานคนที่ 2 - ตรวจสอบความถูกต้อง (14 วินาที) - บันทึกข้อมูล (1 วินาที) - ย้ายกล้องไปเก็บ (5 วินาที) รวมเวลาที่ใช้ 60 วินาที	พนักงานคนที่ 1 - ดูรูปจากไฟล์ภาพ (3 วินาที) - โปรแกรมช่วยประมวลผลภาพ (5 วินาที) - ตรวจสอบความถูกต้อง (10 วินาที) - บันทึกข้อมูล (1 วินาที) รวมเวลาที่ใช้ 19 วินาที	ระยะเวลาที่ใช้ลดลง 41 วินาที คิดเป็น 68.33% จากวิธีการเดิม ขั้นตอนการทำงานลดลง 1 ขั้นตอน และมี 2 ขั้นตอนที่เป็นขั้นตอนประสิทธิภาพการทำงานคือลดการขนย้ายกล้อง แล้วใช้วิธีการดูจากไฟล์ภาพแทน ซึ่งจะช่วยลดพื้นที่การทำงานด้วย
จำนวนคน	2 คน	1 คน	จำนวนคนลดลง 1 คน คิดเป็น 50% ของจำนวนคนทั้งหมดที่ใช้ในวิธีการเดิม
ต้นทุน (ค่าแรงขั้นต่ำ 340 บาท ต่อคนต่อวัน)	ใช้ 2 คน เท่ากับ 680 บาทต่อวัน	ใช้ 1 คน เท่ากับ 340 บาทต่อวัน	ค่าใช้จ่ายลดลง 340 บาทต่อคนต่อวัน คิดเป็น 50% ของต้นทุนที่ใช้ในวิธีการเดิม

รูปที่ ง.49 เอกสารนำเสนอสารนิพนธ์หน้าที่ 49

หัวข้อการนำเสนอ



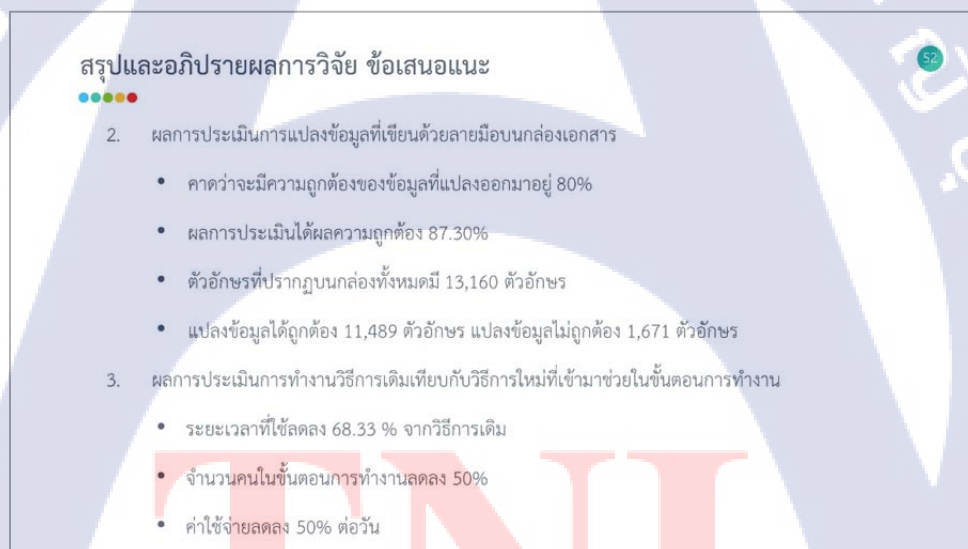
50

1. ความเป็นมาและที่มาของปัญหา
2. วัตถุประสงค์การวิจัย
3. ขอบเขตของการวิจัย
4. ประโยชน์ที่คาดว่าจะได้รับ
5. งานวิจัยและทฤษฎีที่เกี่ยวข้อง
6. แนวคิดที่นำเสนอในงานวิจัย
7. ผลการวิจัย
8. สรุปและอภิปรายผลการวิจัย ข้อเสนอแนะ

รูปที่ ง.50 เอกสารนำเสนอสารนิพนธ์หน้าที่ 50



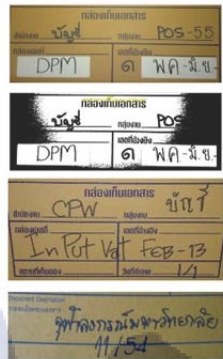
รูปที่ ง.51 เอกสารนำเสนอสารนิพนธ์หน้าที่ 51



รูปที่ ง.52 เอกสารนำเสนอสารนิพนธ์หน้าที่ 52

ปัญหาและอุปสรรคในการวิจัย

1. ภาพกล่องเอกสารติดแสงและเงา
2. การเขียนตัวอักษรในแต่ละช่องไม่เป็นไปตามประเภทของช่อง
3. เขียนตัวอักษรต่อเนื่องกัน
4. ตัวอักษรบางตัวมีการรูปแบบการเขียนใกล้เคียงกัน เช่น ข กับ บ เป็นต้น



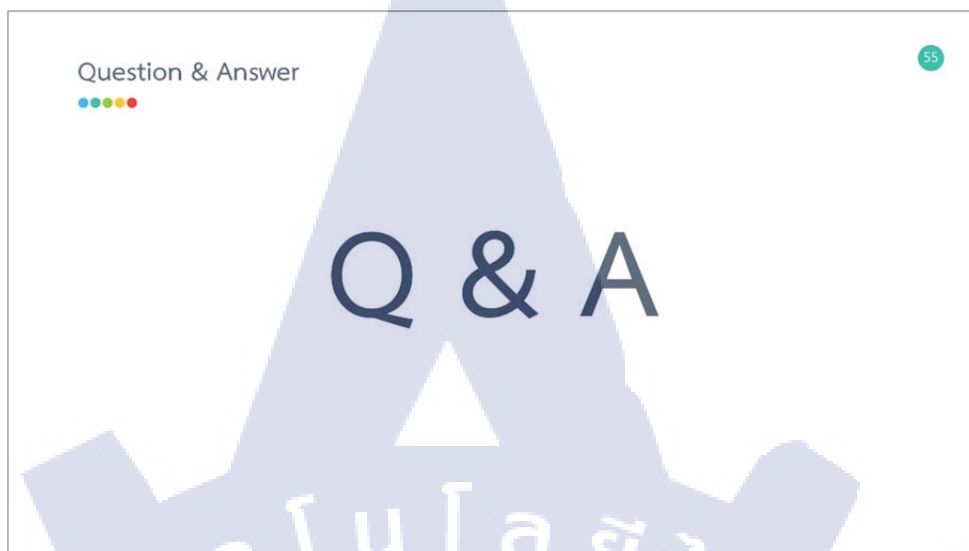
ภาพตัวอย่างกล่องเอกสารที่เป็นอุปสรรคในการวิจัย

รูปที่ ง.53 เอกสารนำเสนอสารนิพนธ์หน้าที่ 53

ข้อเสนอแนะ

- ข้อมูลที่มีความเฉพาะและข้อจำกัดในเรื่องของลายมือที่นำมาทำตัวต้นแบบการรู้จำลายมือ (Training Data)
- กล่องเอกสารใหม่ๆ (Testing Data) ที่มีลายมือแตกต่างจากตัวต้นแบบที่ได้จัดทำไว้ จะทำให้การแปลงตัวอักษรมีความผิดพลาดได้
- ควรจะต้องมีการทำการรู้จำรูปแบบของลายมือของแต่ละแบบไว้เพื่อใช้เป็นข้อมูลในการเปรียบเทียบเพื่อใช้แปลงผลข้อมูลให้มีความแม่นยำมากขึ้น

รูปที่ ง.54 เอกสารนำเสนอสารนิพนธ์หน้าที่ 54



รูปที่ ง.55 เอกสารนำเสนอสารนิพนธ์หน้าที่ 55