

การเปรียบเทียบโครงสร้าง CNN สำหรับการจำแนกสำเนียงไทย: VFNet กับโครงสร้าง CNN เชิงอนุกรม

Comparative Analysis of CNN Architectures for Thai Accent Classification: VFNet vs. Sequential CNN

ธีรภัทร เล็บกลิน¹ สุรเดช จิตประไพกุลศาส²

¹สาขาวิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ มหาวิทยาลัยนเรศวร theerapatserb@gmail.com

²สาขาวิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ มหาวิทยาลัยนเรศวร Suradet@nu.ac.th

บทคัดย่อ

งานวิจัยนี้นำเสนอการเปรียบเทียบโครงสร้างของโมเดลโครงข่ายประสาทเทียมแบบคอนโวลูชัน (CNN) สำหรับการจำแนกสำเนียงภาษาไทย โดยเปรียบเทียบระหว่างสถาปัตยกรรมแบบขนานที่ใช้ตัวกรองหลายขนาดพร้อมกันตามแนวทางของ VFNet: A Convolutional Architecture for Accent Classification [1] กับสถาปัตยกรรมแบบอนุกรมที่เรียงตัวกรองขนาดต่างกันในแต่ละเลเยอร์ (เช่น 3→5→7) อินพุตที่ใช้คือคุณลักษณะ MFCC ที่สกัดจากชุดข้อมูลเสียง Thai Dialect Corpus [2] ผลการทดลองแสดงให้เห็นว่า โมเดลทั้งสองแบบให้ค่า Accuracy และ F1-score ใกล้เคียงกัน อย่างไรก็ตาม เมื่อพิจารณาจำนวนพารามิเตอร์และค่า Cross Entropy Loss พบว่าโครงสร้างแบบอนุกรมบางรูปแบบ เช่น 5→5→5 และ 7→5→3 ให้ผลลัพธ์ที่ดีกว่าทั้งในด้านประสิทธิภาพและความประหยัดทรัพยากร ทั้งนี้ การวิเคราะห์ขนาด Receptive Field (RF) แบบ 2 มิติ ยังช่วยอธิบายได้ว่า โครงสร้างที่มี RF ขนาดกลางจะให้ผลการจำแนกที่ดีกว่าโครงสร้างที่มี RF เล็กหรือใหญ่เกินไป สะท้อนให้เห็นถึงข้อได้เปรียบของการออกแบบสถาปัตยกรรมโมเดลที่เหมาะสม สำหรับการใช้งานจริงในสภาพแวดล้อมที่มีข้อจำกัดด้านหน่วยประมวลผลและหน่วยความจำ

คำสำคัญ: สำเนียงภาษาไทย, การจำแนกสำเนียง, CNN, VFNet, MFCC

Abstract

This research presents a comparative study of Convolutional Neural Network (CNN) architectures for Thai accent classification. It contrasts a parallel architecture based on [1] VFNet: A Convolutional Architecture for Accent Classification, which uses multi-size filters simultaneously, with a sequential architecture that stacks different kernel sizes across layers (e.g., 3→5→7). The input features are Mel-Frequency Cepstral Coefficients (MFCCs) extracted from the Thai Dialect Corpus [2]. Experimental results show that both models achieve comparable accuracy and F1-scores. However, further analysis reveals that sequential models such as 5→5→5 and 7→5→3 outperform the VFNet-based parallel architecture in terms of lower parameter count and cross-entropy loss. A

detailed 2D receptive field (RF) analysis also indicates that architectures with moderate RF sizes tend to deliver better classification performance compared to those with very small or excessively large RFs. These findings emphasize the practical advantages of well-structured sequential CNNs for real-world deployment under computational and memory constraints.

Keywords: Thai accent, accent classification, CNN, VFNet, MFCC

1. บทนำ

ในยุคที่เทคโนโลยีรู้จำเสียงพูด (Automatic Speech Recognition: ASR) มีบทบาทสำคัญต่อการสื่อสารระหว่างมนุษย์กับเครื่องจักร เช่น ระบบผู้ช่วยอัจฉริยะ (Voice Assistant), ระบบบริการลูกค้าอัตโนมัติ (Chatbot) และระบบสั่งงานด้วยเสียง เทคโนโลยีเหล่านี้ต้องอาศัยความแม่นยำในการเข้าใจเสียงพูดของผู้ใช้จากหลากหลายภูมิภาค ความท้าทายในปัญหาที่สำคัญในการพัฒนาระบบ ASR ภาษาไทยคือความหลากหลายของสำเนียง ซึ่งกระตุ้นให้เกิดงานวิจัยที่มุ่งเน้นการสร้างชุดข้อมูลเสียงที่ครอบคลุมความแตกต่างของสำเนียง [2] และการพัฒนาเทคนิคเพื่อรองรับความแปรปรวนของสำเนียง

ดังนั้น งานวิจัยนี้จึงนำเสนอการเปรียบเทียบประสิทธิภาพของโมเดล CNN สองแบบหลัก ได้แก่ สถาปัตยกรรมแบบ VFNet [1] ที่ใช้ฟิลเตอร์หลายขนาดพร้อมกัน และสถาปัตยกรรมแบบอนุกรมที่เรียงลำดับฟิลเตอร์ขนาดต่าง ๆ เพื่อประเมินว่าโครงสร้างแบบใดเหมาะสมกับการใช้งานจริงมากที่สุด ทั้งในแง่ประสิทธิภาพและความประหยัดทรัพยากร

2. การแก้ปัญหา

การแก้ปัญหาสามารถทำได้โดยการเพิ่มความสามารถของระบบในการจำแนกสำเนียงก่อนการรู้จำคำพูด ซึ่งจะช่วยให้ระบบ ASR ปรับตัวเลือกโมเดลให้เหมาะสมกับสำเนียงผู้พูดแต่ละคนได้อย่างชาญฉลาด หรือนำไปใช้ในการทำ label อัตโนมัติ สำหรับงาน ASR งานวิจัยนี้จึงมีเป้าหมายเพื่อพัฒนาโมเดลการจำแนกสำเนียงภาษาไทยที่มีประสิทธิภาพสูง พร้อมกับเปรียบเทียบโครงสร้างโมเดลโครงข่ายประสาทแบบคอนโวลูชัน (CNN) ที่มีแนวทางการออกแบบต่างกัน ได้แก่

โครงสร้างที่อิงจากงานวิจัย VFNet [1] ซึ่งออกแบบให้ใช้ตัวกรองความถี่หลายขนาดพร้อมกัน เพื่อดึงคุณลักษณะเสียงในหลายช่วงสเกลอย่างมีประสิทธิภาพ

โครงสร้างแบบอนุกรม ซึ่งเรียงลำดับการใช้ kernel ขนาดต่าง ๆ คล้ายกับ VFNet [1] แต่นำมาจัดเรียงแบบลำดับ (เช่น $3 \times 3 \rightarrow 5 \times 3 \rightarrow 7 \times 3$) เพื่อสกัดข้อมูลจากรายละเอียดระดับล่างไปสู่บริบทกว้างอย่างเป็นลำดับขั้น รวมถึง โครงสร้างแบบอนุกรมที่ใช้ ตัวกรองขนาดเท่ากันในทุก layer ลำดับ (เช่น $3 \times 3 \rightarrow 3 \times 3 \rightarrow 3 \times 3$)

ทั้งสองแนวทางได้รับการประเมินด้วยข้อมูลจากชุด Thai Dialect Corpus [2] ซึ่งรวบรวมเสียงพูดจากผู้ใช้นามกลุ่มสำเนียงหลัก ได้แก่ สำเนียงไทยกลาง สำเนียงไทยคำเมือง และสำเนียงไทยโคราช โดยใช้อินพุตเป็นคุณลักษณะ MFCC ซึ่งเป็นหนึ่งในฟีเจอร์ที่นิยมในงานรู้จำเสียง

แม้สถาปัตยกรรม VFNet จะมีจุดเด่นด้านการรวบรวมฟีเจอร์จากหลายขนาด แต่ก็ต้องแลกกับจำนวนพารามิเตอร์ที่มากขึ้น ในขณะที่สถาปัตยกรรมแบบอนุกรมอาจมีศักยภาพในการลดจำนวนพารามิเตอร์ได้มากโดยไม่ลดทอนประสิทธิภาพ ซึ่งจะเป็ข้อได้เปรียบในบริบทของการนำโมเดลไปใช้งานจริงที่มีข้อจำกัดด้านทรัพยากรและอีกทั้งถ้าโมเดลมีการนำไปใช้ร่วมกับงาน ASR ซึ่งต้องใช้ทรัพยากรมากในการประมวลผล ยิ่งต้องคำนึงถึงจำนวนพารามิเตอร์ที่ใช้

นอกจากการเปรียบเทียบประสิทธิภาพของโมเดลในแง่ของ Accuracy และ F1-score แล้ว งานวิจัยนี้ยังได้วิเคราะห์ความสัมพันธ์กับขนาดของ Receptive Field (RF) และจำนวนพารามิเตอร์ เพื่อค้นหาความสัมพันธ์ระหว่างคุณภาพของโมเดลและต้นทุนการประมวลผล งานวิจัยนี้จึงมีคุณค่าทั้งในเชิงวิชาการและการประยุกต์ โดยสามารถนำผลการศึกษาไปต่อยอดในการสร้างระบบ ASR ภาษาไทยที่สามารถเข้าใจเสียงผู้ใช้งานหลากหลายสำเนียงได้อย่างมีประสิทธิภาพและคุ้มค่าต่อการใช้งานจริง

3. วิธีการทดลอง

3.1 ชุดข้อมูล

งานวิจัยนี้ใช้ชุดข้อมูลเสียงพูดจากโครงการ Thai Dialect Corpus ซึ่งเผยแพร่โดย SLSCU (Speech and Language Processing Laboratory, Chulalongkorn University) โดยเลือกมาเฉพาะ 3 กลุ่มสำเนียง ได้แก่ สำเนียงไทยกลาง คัดเลือกมา แบ่งเป็น train 50 ชั่วโมง validation 1.71 ชั่วโมง สำเนียงไทยโคราช (ตะวันออกเฉียงเหนือ) train 46.63 validation ชั่วโมง 1.50 ชั่วโมง สำเนียงไทยคำเมือง (ภาคเหนือ) train 32.43 ชั่วโมง validation 1.50 ชั่วโมง

3.2 การเตรียมข้อมูลและการสกัดคุณลักษณะ

การสกัดคุณลักษณะเสียงในงานวิจัยนี้ใช้เทคนิค Mel Frequency Cepstral Coefficients (MFCC) โดยกำหนดค่า Sampling Rate ที่ 22,050 เฮิรตซ์ และสกัดคุณลักษณะ MFCC ออกมาทั้งหมดจำนวน 64 ค่าจากสเปกตรัมที่ได้จากการประมวลผลผ่านกรองเสียง Mel จำนวน 64 แถบความถี่ (Mel filter banks)

กระบวนการประมวลผลนี้อาศัยการใช้หน้าต่างฟูเรียร์แบบ FFT (Fast Fourier Transform) ที่มีขนาด 512 จุด และกำหนดค่าการเลื่อนหน้าต่าง (hop length) ที่ 256 จุด เพื่อควบคุมความละเอียดเชิงเวลา โดย

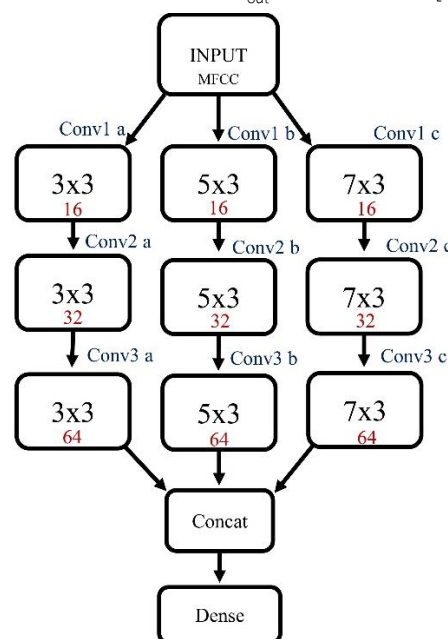
ขั้นตอนการสกัดทั้งหมดดำเนินการผ่านฟังก์ชัน torchaudio.transforms.MFCC จากไลบรารี Torchaudio

3.3 สถาปัตยกรรมของโมเดล

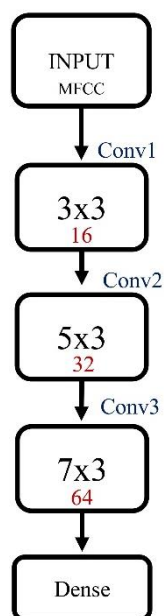
ในบทความในการทดลอง ได้เปรียบเทียบโมเดลหลายรูปแบบซึ่งแบ่งออกเป็น 3 กลุ่มหลัก ได้แก่ โมเดลแบบอนุกรม (Sequential CNN) ที่ใช้ kernel สามขนาดที่เรียงลำดับกันในสามเลเยอร์ เช่น $3 \times 3 \rightarrow 5 \times 3 \rightarrow 7 \times 3$ โดยใช้โดเมนเวลาเท่ากันคือ 3, โมเดลแบบขนาน VFNet [1] ที่ใช้ kernel หลายขนาดพร้อมกันโดยเลือกมา 3 ขนาดได้แก่ 3×3 , 5×3 และ 7×3 และโมเดลแบบหนึ่งเลเยอร์ (1-layer CNN) ที่ใช้ kernel เพียงขนาดเดียว เช่น 3×3 หรือ 5×5 ทั้งนี้ การเลือกใช้ตัวกรองสามขนาดดังกล่าวมีเหตุผลหลักคือเพื่ออ้างอิงตามสถาปัตยกรรมของ VFNet [1] ซึ่งเป็นโมเดลต้นแบบสำหรับการเปรียบเทียบและเพื่อจับคุณลักษณะเสียงในหลายระดับสเกล (multi-scale) [1] ตั้งแต่รายละเอียดระดับต่ำด้วยฟิลเตอร์ขนาดเล็ก (3×3) ไปจนถึงบริบทที่กว้างขึ้นด้วยฟิลเตอร์ขนาดใหญ่ (7×3) นอกจากนี้ การเลือกใช้ขนาดเป็นเลขคี่ยังเป็นหลักปฏิบัติทั่วไปในการออกแบบ CNN เพื่อให้ฟิลเตอร์มีพิกเซลกลางที่ชัดเจนซึ่งช่วยในการรักษาลำดับชั้นเชิงพื้นที่ของคุณลักษณะได้ดี [3-4] โดยทุกรูปแบบของโมเดลมีการใช้ ReLU, Batch Normalization และ MaxPooling ในแต่ละชั้น Convolution เพื่อเพิ่มประสิทธิภาพในการเรียนรู้ โมเดลแบบ 3 เลเยอร์ใช้จำนวนฟิลเตอร์ในแต่ละเลเยอร์เท่ากับ 16, 32 และ 64 ตามลำดับ ส่วนโมเดลแบบ 1 เลเยอร์ใช้จำนวนฟิลเตอร์คงที่เท่ากับ 32 ฟิลเตอร์ ทุกโมเดล CNN นำมาต่อด้วย 128 Fully Connected Layers แต่ละโมเดลมีจำนวนพารามิเตอร์ของ CNN ตามสูตรนี้

$$\text{Parameters} = (C_{in} \times K_h \times K_w + 1) \times C_{out} \quad (1)$$

โดยที่ C_{in} คือจำนวนช่องสัญญาณเข้า $K_h \times K_w$ คือคือขนาดของ kernel +1 คือ bias มีค่า 1 ต่อ filter และ C_{out} คือจำนวนฟิลเตอร์ [4]



รูปที่ 1 โครงสร้าง VFNet แบบ 3 layer ที่ใช้หลาย kernel ขนาดพร้อมกัน



รูปที่ 2 ตัวอย่างโครงสร้าง โมเดลเชิงอนุกรมที่ใช้ kernel หลายขนาดแบบเรียงลำดับ

3.4 การป้องกัน Overfitting และการจัดการความไม่สมดุลของข้อมูล

ในการฝึกโมเดลโครงข่ายประสาทเทียมแบบคอนโวลูชัน (CNN) ในงานวิจัยนี้ ได้มีการกำหนดขั้นตอนและพารามิเตอร์ที่เหมาะสมเพื่อเพิ่มประสิทธิภาพการเรียนรู้และลดปัญหา Overfitting โดยใช้ตัวปรับค่าพารามิเตอร์ (Optimizer) คือ Adam ซึ่งตั้งค่าอัตราการเรียนรู้เริ่มต้นไว้ที่ 0.001

จำนวนรอบการฝึก (Epoch) สูงสุดตั้งไว้ที่ 1000 แต่มีการใช้กลไก Early Stopping เพื่อหยุดการฝึกหากโมเดลไม่มีแนวโน้มพัฒนา โดยกำหนดค่า Patience เท่ากับ 25 หมายความว่า หาก Validation Loss ไม่ลดลงต่อเนื่องใน 25 รอบ จะหยุดการฝึกทันที และใช้ค่า Min Delta เท่ากับ 0.001 เป็นเกณฑ์พิจารณาว่าการลดลงของ Validation Loss ถือว่ามีความหมายเพียงพอ การบันทึกโมเดลจะกระทำเฉพาะกรณีที่ Validation Loss ต่ำที่สุดตลอดการฝึก

เพื่อเสริมความสามารถในการควบคุมอัตราการเรียนรู้ให้เหมาะสมกับช่วงเวลาต่าง ๆ ของการฝึก ได้มีการใช้กลยุทธ์ปรับค่าอัตราการเรียนรู้แบบ ReduceLROnPlateau โดยหาก Validation Loss ไม่มีการปรับลดลงในช่วง 5 Epoch ระบบจะลดค่า learning rate ลงครึ่งหนึ่งโดยอัตโนมัติ

เพื่อลดความเสี่ยงของการเกิด Overfitting ได้มีการใช้ Dropout กับแต่ละชั้นของโมเดล โดยกำหนด Dropout Rate เท่ากับ 0.3 สำหรับ Convolutional Layers และ 0.5 สำหรับ Fully Connected Layers ซึ่ง

ช่วยป้องกันการเรียนรู้ที่จำเพาะเจาะจงเกินไป และส่งเสริมให้โมเดลสามารถ generalize ได้ดีขึ้น [3]

สำหรับการแก้ปัญหาความไม่สมดุลของข้อมูลในแต่ละคลาส ได้ใช้เทคนิค Random Oversampling โดยการสุ่มเพิ่มจำนวนตัวอย่างในคลาสที่มีจำนวนน้อยให้เทียบเท่ากับคลาสที่มีจำนวนมากที่สุด ซึ่งช่วยให้โมเดลไม่เกิด Bias ต่อคลาสที่มีข้อมูลจำนวนมากกว่า

แนวทางทั้งหมดนี้มีจุดมุ่งหมายเพื่อให้โมเดลมีประสิทธิภาพสูงสุดภายใต้เงื่อนไขของทรัพยากรที่จำกัด และลดความซับซ้อนของการฝึกที่อาจไม่จำเป็น

4. ตารางพารามิเตอร์ของโมเดลและผลการทดลอง

ตารางที่ 1 จำนวนพารามิเตอร์โมเดลแบบ 1 เลเยอร์

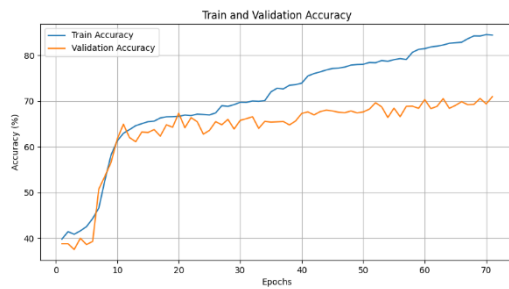
โครงสร้างโมเดล	Conv1	รวมทั้งหมด
1 Layer (3x3)	320	320
1 Layer (5x3)	512	512
1 Layer (7x3)	704	704
1 Layer (5x5)	832	832
1 Layer (7x7)	1600	1600
VFNet 1 Layer (3x3,5x3,7x3)	1536	1536

ตารางที่ 2 โมเดลแบบ 1 เลเยอร์

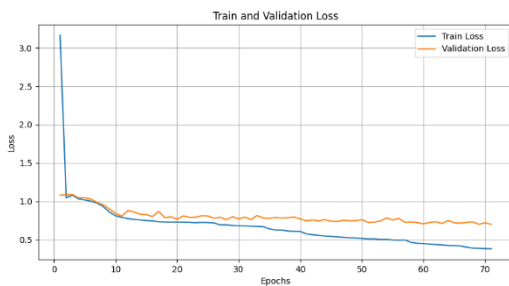
Model	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)	Best Cross entropy loss
VFNet	69.81	70.09	69.75	69.71	0.7144
3 x 3	67.54	67.69	67.51	67.52	0.7369
5 x 3	70.96	70.93	70.68	70.71	0.7000
7 x 3	69.30	69.17	68.86	68.96	0.7085
5 x 5	68.22	69.32	68.70	68.30	0.7055
7 x 7	69.63	70.17	69.49	69.54	0.6962

ตารางที่ 3 ตัวอย่าง classification report ของขนาด 5x3 หนึ่งในโมเดลแบบ 1 เลเยอร์

Class	Precision (%)	Recall (%)	F1-score (%)
คำเมือง	65.89	69.90	67.83
โคราช	71.88	74.54	73.18
กลาง	75.03	67.59	71.12
Accuracy	-	-	70.96
Macro Avg	70.93	70.68	70.71
Weighted Avg	71.16	70.96	70.97



รูปที่ 3 ตัวอย่าง learning curve ค่า Accuracy ของขนาด 5x3 หนึ่งในโมเดลแบบ 1 เลเยอร์



รูปที่ 4 ตัวอย่าง learning curve ค่า Cross entropy loss ของขนาด 5x3 หนึ่งในโมเดลแบบ 1 เลเยอร์

ตารางที่ 4 จำนวนพารามิเตอร์โมเดลแบบ 3 เลเยอร์

โครงสร้างโมเดล	Conv1	Conv2	Conv3	รวมทั้งหมด
3→3→3	160	4640	18496	23296
3→5→7	160	7712	43072	50944
3→7→5	160	10784	30784	41728
5→5→5	256	7712	30784	38752
5→3→7	256	4640	43072	47968
5→7→3	256	10784	18496	29536
7→7→7	352	10,784	43,072	54,208
7→3→5	352	4640	30784	35776
7→5→3	352	7712	18496	26560
VFNet 3L	768	23136	92352	116256

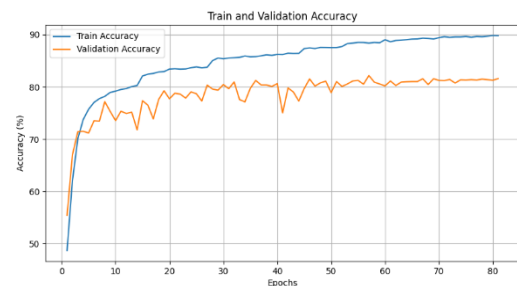
ตารางที่ 5 โมเดลแบบ 3 เลเยอร์

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)	Best Cross entropy loss
3→3→3	80.34	80.57	80.67	80.58	0.5581
3→5→7	80.66	81.86	80.00	80.46	0.5693
3→7→5	82.57	83.58	82.15	82.60	0.5333
5→5→5	82.71	82.74	82.86	82.70	0.5212

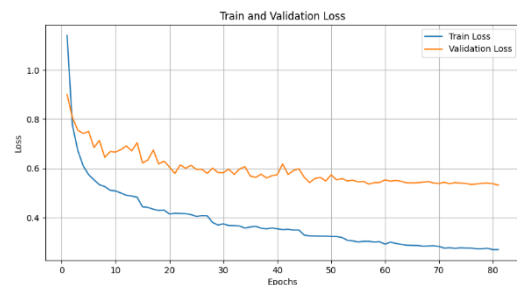
5→3→7	80.27	80.69	80.26	80.40	0.5424
5→7→3	81.52	81.70	81.58	81.62	0.5196
7→7→7	80.66	81.37	80.46	80.71	0.5485
7→3→5	80.66	80.93	81.36	80.75	0.5276
7→5→3	82.17	83.05	81.93	82.23	0.5325
VFNet	79.76	81.09	79.39	79.70	0.5572

ตารางที่ 6 ตัวอย่าง classification report ของโมเดล 7→5→3 หนึ่งในโมเดลแบบ 3 เลเยอร์

Class	Precision (%)	Recall (%)	F1-score (%)
คำเมือง	82.09	83.83	83.36
โคราช	77.53	86.57	81.80
กลาง	88.73	75.39	81.52
Accuracy	-	-	82.17
Macro Avg	83.05	81.93	82.23
Weighted Avg	82.70	82.17	82.16



รูปที่ 5 ตัวอย่าง learning curve ค่า Accuracy ของโมเดล 7→5→3 หนึ่งในโมเดลแบบ 3 เลเยอร์



รูปที่ 6 ตัวอย่าง learning curve ค่า Cross entropy loss ของโมเดล 7→5→3 หนึ่งในโมเดลแบบ 3 เลเยอร์

ตารางที่ 7 Receptive Field (RF) 2D สำหรับแต่ละโครงสร้าง CNN

โครงสร้างโมเดล	RF ชั้น Layer 1	RF ชั้น Layer 2	RF ชั้น Layer 3	RF สุดท้าย
3→3→3	4×4	10×10	22×22	22×22
3→5→7	4×4	14×10	42×22	42×22

3→7→5	4×4	18×10	38×22	38×22
5→3→7	6×4	12×10	40×22	40×22
5→5→5	6×4	16×10	36×22	36×22
5→7→3	6×4	20×10	32×22	32×22
7→3→5	8×4	14×10	34×22	34×22
7→5→3	8×4	18×10	30×22	30×22
7→7→7	8×4	22×10	50×22	50×22

ตารางที่ 8 เปรียบเทียบ RF F1-score และ Cross Entropy Loss

โครงสร้าง	RF	F1-score (%)	Loss	พารามิเตอร์	ข้อสังเกต
3→3→3	22×22	80.58	0.5581	23,296	baseline, param ต่ำ แต่ loss สูง
3→5→7	42×22	80.46	0.5693	50,944	RF ใหญ่ → param สูง แต่ผลไม่ดี
3→7→5	38×22	82.60	0.5333	41,728	F1 สูง, param ปานกลาง
5→5→5	36×22	82.70	0.5212	38,752	สมดุลที่สุด
5→3→7	40×22	80.40	0.5424	47,968	param สูง แต่ไม่คุ้ม
5→7→3	32×22	81.62	0.5196	29,536	ดีมาก loss ต่ำสุดและ parameter ต่ำ
7→7→7	50×22	80.71	0.5485	54,208	RF ใหญ่มาก, param สูง
7→3→5	34×22	80.75	0.5276	35,776	สมดุลดี, F1 ปานกลาง
7→5→3	30×22	82.23	0.5325	26,560	F1 สูง, param น้อยสุดในกลุ่มที่แม่นยำเกิน 82
VFNet (3L)	-	79.70	0.5572	116,256	แม้ใช้ multi-scale แต่แพ้ทุกตัวด้าน efficiency

5. อภิปรายผล

จากผลการทดลองและตารางเปรียบเทียบโครงสร้างโมเดล CNN สำหรับการจำแนกสำเนียงไทย พบว่าโครงสร้างแบบ อนุกรม (sequential) เช่น 5→5→5 และ 3→7→5 ให้ผลลัพธ์ที่ดียิ่งขึ้น

ด้าน Accuracy, F1-score และ Loss โดยไม่ต้องใช้จำนวนพารามิเตอร์สูงเกินจำเป็น

โดยเฉพาะโครงสร้าง 5→5→5 ซึ่งให้ F1-score สูงที่สุด (82.70%), Loss ต่ำที่ (0.5212) และใช้พารามิเตอร์รวมเพียง 38,752 แสดงให้เห็นถึงสมดุลระหว่างประสิทธิภาพและต้นทุนโมเดล

นอกจากนี้ยังพบว่าโครงสร้าง 7→5→3 ซึ่งเรียง kernel จากใหญ่ไปเล็ก ให้ผลลัพธ์ F1-score สูงถึง 82.23% โดยใช้พารามิเตอร์เพียง 26,560 ต่ำที่สุดในกลุ่มโมเดลที่ให้ผลลัพธ์ระดับสูง ซึ่งเห็นว่าการออกแบบโครงสร้างให้เรียงลำดับ kernel อย่างมีแบบแผนสามารถลดต้นทุนโมเดลได้โดยไม่ลดประสิทธิภาพ

ในด้านของ Receptive Field (RF) พบว่าโครงสร้างที่มี RF ในช่วง 30–38 (แกน Frequency) และ 22 (แกน Time) มักจะให้ผลลัพธ์ดีที่สุด เนื่องจากสามารถจับบริบทได้กว้างพอ โดยไม่สูญเสียรายละเอียดในพื้นที่เล็ก

โครงสร้างที่มี RF ใหญ่เกินไป เช่น 7→7→7 (RF = 50×22) หรือ เล็กเกินไป เช่น 3→3→3 (RF = 22×22) มีแนวโน้มให้ผลลัพธ์ด้อยลงในด้าน Loss และ F1-score

สำหรับโครงสร้าง VFNet แบบ 3 ชั้น แม้ออกแบบให้ดึงฟีเจอร์หลายขนาดพร้อมกัน แต่กลับใช้พารามิเตอร์สูงถึง 116,256 และยังให้ F1-score เพียง 79.70% ต่ำกว่าหลายโมเดลแบบอนุกรม

โมเดล 1 เลเยอร์แบบ VFNet ที่ใช้ฟิลเตอร์ 3×3, 5×3 และ 7×3 พร้อมกัน ให้ F1-score สูงกว่า kernel เดี่ยวเพียงเล็กน้อย แม้จะใช้พารามิเตอร์เพิ่มขึ้นจาก 320 → 1,536 ซึ่งยังถือว่าไม่คุ้มค่าในการใช้งานจริง

6. สรุป

การเปรียบเทียบโครงสร้าง Convolutional Neural Network สำหรับการจำแนกสำเนียงไทยในครั้งนี้ แสดงให้เห็นว่าโครงสร้างแบบ อนุกรมที่เลือกใช้ kernel ขนาดต่างกันในแต่ละชั้นอย่างมีแบบแผน เช่น 5→7→3, 3→7→5 และ 7→5→3 สามารถให้ประสิทธิภาพสูงเทียบเท่าหรือเหนือกว่า VFNet แบบขนาน แม้ใช้พารามิเตอร์น้อยกว่าอย่างชัดเจน

โครงสร้าง VFNet ที่ใช้ kernel หลายขนาดพร้อมกันมีแนวโน้มใช้พารามิเตอร์สูงมาก (เช่น 116,256 ในแบบ 3 ชั้น) แต่กลับไม่ได้ให้ประสิทธิภาพที่เหนือกว่า ขณะที่โมเดลแบบอนุกรมเช่น 7→5→3 ใช้พารามิเตอร์เพียง 26,560 แต่ให้ F1-score สูงถึง 82.23%

ผลการวิเคราะห์ Receptive Field ยังช่วยสนับสนุนว่า RF ที่เหมาะสมควรอยู่ในช่วง 30–38 (แกนความถี่) เพื่อให้จับบริบทได้ดีและไม่สูญเสียความละเอียด

การออกแบบโครงสร้างโมเดล CNN ที่พิจารณาทั้งขนาด kernel, ลำดับการเรียง, RF และจำนวนพารามิเตอร์สามารถช่วยเพิ่มประสิทธิภาพในการจำแนกสำเนียง และลดต้นทุนการประมวลผลในระบบจริงได้

เอกสารอ้างอิง

- [1] A. Ahmed, P. Tangri, A. Panda, D. Ramani, and S. Karmakar, "VFNet: A Convolutional Architecture for Accent Classification," in *Proc. IEEE 16th India Council Int. Conf. (INDICON)*, Nov. 1-4 2019, pp. 1-4.
- [2] A. Suwanbandit, B. Naowarat, O. Sangpetch, and E. Chuangsuwanich, "Thai Dialect Corpus and Transfer-based Curriculum Learning for Dialect ASR," in *Proc. Interspeech 2023*, Dublin, Ireland, Aug. 20-24 2023, pp. 4069-4073, doi: 10.21437/Interspeech.2023-1828.
- [3] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, MA: MIT Press, 2016.
- [4] Stanford CS231n, "Convolutional Neural Networks for Visual Recognition," [Online]. Available: <http://cs231n.stanford.edu> [Accessed: 19-Jun-2025].



นายธีรภัทร เล็บกลิน นิสิต ปริญญาโท สาขา
วิศวกรรมคอมพิวเตอร์ (วศ.ม.) ม.นเรศวร สนใจด้าน
การประมวลผลภาษาธรรมชาติ (NLP) และการ
เรียนรู้เชิงลึก (Deep Learning)



ดร.สุรเดช จิตประไพกุลศาล อาจารย์ประจำ
ภาควิชาวิศวกรรมไฟฟ้าและคอมพิวเตอร์ ม.นเรศวร
ความเชี่ยวชาญ Mathematical Programming,
วิศวกรรมซอฟต์แวร์, Cybersecurity, Machine
Learning