

บทที่ 5

สรุปผลการวิจัยและข้อเสนอแนะ

สรุปผลการวิจัย

จากผลการทดลองในขั้นตอนของระบบ LVCSR จะเห็นได้ว่าผลลัพธ์ที่ดีที่สุดจากการทดลองในเบื้องต้นในบทที่ 3 ข้อมูลที่ใช้ในการฝึกฝนแบบจำลองเสียง คือข้อมูล LOTUS Corpus 50 ชั่วโมง ข้อมูลที่ใช้ในการฝึกฝนแบบจำลองภาษา คือข้อมูล BN Corpus 9,115 ประโยค 170,219 คำ 7,123 คำศัพท์ และคำศัพท์ในพจนานุกรมเสียงอ่านมาจาก คำศัพท์ใน LOTUS Corpus รวมกับคำศัพท์ของชุดข้อมูลทดสอบ ซึ่งจะได้เปอร์เซ็นต์ความถูกต้อง 68.64% และเปอร์เซ็นต์ความแม่นยำ 63.56%

ในเบื้องต้นผู้วิจัยตั้งสมมติฐานว่า ถ้าเพิ่มข้อมูลของข่าวในการฝึกฝนแบบจำลองภาษาให้มีคำศัพท์มากขึ้นอาจจะทำให้เปอร์เซ็นต์ความถูกต้องและเปอร์เซ็นต์ความแม่นยำเพิ่มมากขึ้น ดังนั้นจึงเพิ่มข้อมูลข่าวของหนังสือพิมพ์ไทยรัฐเข้าไปฝึกฝนแบบจำลองภาษา แต่ผลการทดลองที่ได้นั้นให้เปอร์เซ็นต์ความถูกต้องและเปอร์เซ็นต์ความแม่นยำลดลง เพราะชุดข้อมูลการทดสอบมีขนาดเล็ก คำศัพท์ที่มากเกินไปจนความจำเป็น ทำให้เกิดความผิดพลาดในการถอดรหัสได้มากขึ้นด้วย

ดังนั้นจึงเลือกแบบจำลองในเบื้องต้นมาทำการปรับปรุง โดยใช้แบบจำลองเสียงและแบบจำลองภาษาเหมือนเดิม แต่เปลี่ยนพจนานุกรมเสียงอ่าน โดยนำคำหยุด (Stop word list) 125 คำ ออกจากพจนานุกรมเสียงอ่าน เพราะคำเหล่านี้ไม่น่าจะเกิดเป็นคำสำคัญได้ พจนานุกรมเสียงอ่านที่ใช้ในการถอดรหัสนั้นจะมีคำศัพท์ 9,717 คำ ซึ่งแบบจำลองที่ปรับปรุงนี้ให้เปอร์เซ็นต์ความถูกต้อง 78.32 % และเปอร์เซ็นต์ความแม่นยำ 74.34 % ซึ่งเป็นเกณฑ์ที่น่าพอใจ จึงนำแบบจำลองนี้ไปใช้ในการค้นหาคำสำคัญต่อไป

ขั้นตอนในการค้นหาคำสำคัญ จากการใช้คำสำคัญทั้งคำไปค้นหา จะได้ผลลัพธ์ในการค้นหาไม่ตึก เพราะผลที่ได้จากระบบ LVCSR นั้นไม่ได้ให้มีความถูกต้อง 100% คำสำคัญที่นำไปค้นหานั้นอาจจะไม่ตรงกับที่ต้องการ จึงใช้วิธีการ n-best ซึ่งในที่นี้ $n = 10$ จะได้คำสำคัญ 10 คำ มาช่วยในการค้นหา วิธีการค้นหานี้จะให้ผลที่ดีขึ้น แต่ยังไม่เป็นที่น่าพอใจ ดังนั้นจึงใช้คำสำคัญย่อย (Sub keyword) มาช่วยในการค้นหาเพิ่มเติม ผลลัพธ์ที่ได้ในการค้นหานี้เป็นที่ยอมรับได้

ข้อเสนอแนะ

เนื่องจากฐานข้อมูลเสียงภาษาไทยในขณะนี้ยังมีไม่มากนัก ข้อมูลที่ใช้ในงานวิทยานิพนธ์นี้เป็นฐานข้อมูลข่าวเพียงอย่างเดียว การค้นหาข้อมูลจึงได้เฉพาะในเรื่องข่าว ดังนั้นถ้ามีฐานข้อมูลเสียงอื่นๆ มาฝึกฝนแบบจำลองนั้น จะทำให้หมวดหมู่ของข้อมูลที่ต้องการค้นหาวางขึ้น สามารถพัฒนาต่อจากงานวิทยานิพนธ์ต่อไป

ในงานวิทยานิพนธ์นี้ได้ใช้ชุดข้อมูลทดสอบขนาดเล็ก ซึ่งเป็นเพียงกลุ่มคำสั้นๆ จึงอาจจะทำให้ประสิทธิภาพของระบบ BN - LVCSR นั้นไม่ดีนัก ถ้าเพิ่มชุดข้อมูลทดสอบให้มีขนาดใหญ่ขึ้น อาจจะทำให้เปอร์เซ็นต์ความถูกต้องและเปอร์เซ็นต์ความแม่นยำเพิ่มขึ้นได้ และแบบจำลองหน่วยเสียงของวิทยานิพนธ์นี้ไม่มีการใช้ข้อมูลของเสียงวรรณยุกต์ ดังนั้นถ้าเพิ่มข้อมูลของเสียงวรรณยุกต์ในแบบจำลองหน่วยเสียงแล้ว จะทำให้เปอร์เซ็นต์ความถูกต้องและเปอร์เซ็นต์ความแม่นยำเพิ่มขึ้นได้อีกด้วย

และในงานวิทยานิพนธ์สามารถนำไปใช้กับภาษาอื่นได้ โดยเฉพาะภาษาที่มีลักษณะโครงสร้างคล้ายคลึงกับภาษาไทย คือเป็นเป็นภาษาที่มีระดับเสียงของคำแน่นอนหรือมีวรรณยุกต์และออกเสียงแยกคำต่อคำ ตัวอย่างเช่น ภาษาจีน แต่อย่างไรก็ตามเนื่องด้วยข้อจำกัดและเอกลักษณ์เฉพาะของแต่ละภาษา ซึ่งอาจจะมีข้อกำหนดหรือข้อยกเว้นต่างๆ ไม่เหมือนกัน จึงทำให้ผลลัพธ์ที่ได้อาจจะไม่สมบูรณ์เทียบเท่ากับการใช้กับภาษาไทย ทั้งนี้อาจจะต้องมีการประยุกต์ใช้กับภาษาอื่นๆ เพื่อเพิ่มประสิทธิภาพในการค้นหามากขึ้น

ผู้วิจัยหวังเป็นอย่างยิ่งว่างานวิจัยนี้จะเป็นรากฐานให้แก่งานด้านการสืบค้นด้วยเสียง ซึ่งจะเป็นประโยชน์แก่ผู้ที่ต้องการใช้งานเฉพาะทางต่างๆ เช่น การสืบค้นข้อมูลในโทรศัพท์มือถือ หรือการสืบค้นด้วยเสียงเพื่อหาข้อมูลในห้องสมุด หรือในสถานที่ต่างๆ เช่น พิพิธภัณฑ์ แหล่งท่องเที่ยวต่างๆ