

บทคัดย่อ

ประสิทธิภาพของ ระบบรู้จำเสียงพูดอัตโนมัติและ ระบบสังเคราะห์เสียงพูด ขึ้นอยู่กับ ความครอบคลุมของหน่วยเสียงจากคลังข้อความที่เหมาะสม งานวิจัยนี้เสนอการสร้างคลังข้อความ อัตโนมัติ จากการกระจายตัวของหน่วยเสียงตามที่กำหนดการกระจายตัวของหน่วยตามที่กำหนด นั้น สามารถกำหนดได้จากชนิดของหน่วยเสียง ขนาดของคลังข้อความ เกณฑ์ขั้นต่ำของจำนวน หน่วยเสียง และรูปแบบของการกระจายตัวเป้าหมาย ได้คัดเลือกข้อความมาจากข้อมูลจาก อินเทอร์เน็ต โดยข้อความนั้นจะถูก จัดเก็บอย่างต่อเนื่องโดย กระบวนการดึงบทความจากหน้า เว็บบนอินเทอร์เน็ต จนกระทั่งได้คลังข้อความที่เหมาะสม ในวิทยานิพนธ์นี้ยังได้ประยุกต์ใช้วิธีการ เชิงละโมบ เพื่อเลือกประโยคที่เหมาะสมที่จะทำให้เกิดการกระจายตัวของหน่วยเสียงตาม เป้าหมาย ในการทดลองได้ใช้ข้อความจากฐานข้อมูล Large Vocabulary Continuous Speech Recognition (LVCSR) corpus for Thai language ในการสร้างเป้าหมายของการกระจายตัว หน่วยเสียง ผลการทดลองที่ได้คือ จำนวนของข้อมูลข้อความที่ดึงมาจากอินเทอร์เน็ตที่เพิ่มขึ้น สามารถทำให้การกระจายตัวของหน่วยเสียงเป็นไปตามเป้าหมายได้ และเกิดความครอบคลุมทาง หน่วยเสียงคู่ ถึง 99.13% คลังข้อความที่ถูกสร้างขึ้นนี้ จึงสามารถนำไปใช้ในการสร้างคลัง เสียงพูดได้อย่างมีประสิทธิภาพ

ABSTRACT

Performance of Automatic Speech Recognition (ASR) and Text-to-Speech (TTS) systems depend on appropriate text corpus. This article explains about the automated text corpus generating method using custom phonetic distribution. This distribution is defined by phonemes type, corpus size, minimum criterion number of phonemes, and target phonetic distribution. Generally, the system selects text data from the internet by continuously downloading them using web crawler. The greedy algorithm is applied to extract the proper sentences, in order to fit with the target phonetic distribution until the appropriate text corpus is established. The experiment is done by using the text from Large Vocabulary Continuous Speech Recognition (LVCSR) corpus for Thai language to generate target phonetic distribution. The result shown that, the increased number of data drawn from the internet is able to accomplish target phonetic distribution and generate diphone coverage for 99.13%. This text corpus then generate speech corpus efficiently.

