



การศึกษารมณเพลงเพื่อพัฒนาธุรกิจ

THE MOODS OF THE LYRICS IN THE FUTURE FOR BUSINESS



กรกัญญา ศิริเกษ

บัณฑิตวิทยาลัย มหาวิทยาลัยศรีนครินทรวิโรฒ

2563

การศึกษารมณเพื่องเพื่อพัฒนาธุรกิจ



สารนิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร
วิทยาศาสตรมหาบัณฑิต สาขาวิชาวิทยาการข้อมูล
คณะวิทยาศาสตร์ มหาวิทยาลัยศรีนครินทรวิโรฒ
ปีการศึกษา 2563
ลิขสิทธิ์ของมหาวิทยาลัยศรีนครินทรวิโรฒ

THE MOODS OF THE LYRICS IN THE FUTURE FOR BUSINESS



A Master's Project Submitted in Partial Fulfillment of the Requirements
for the Degree of MASTER OF SCIENCE
(Data Science)

Faculty of Science, Srinakharinwirot University

2020

Copyright of Srinakharinwirot University

สารนิพนธ์
เรื่อง
การศึกษาอารมณ์เพลงเพื่อพัฒนารัฐกิจ
ของ
กรกัญญา ศิริเกษ

ได้รับอนุมัติจากบัณฑิตวิทยาลัยให้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร
ปริญญาวิทยาศาสตรมหาบัณฑิต สาขาวิชาวิทยาการข้อมูล
ของมหาวิทยาลัยศรีนครินทรวิโรฒ

.....
(รองศาสตราจารย์ นายแพทย์ฉัตรชัย เอกปัญญาสกุล)

คณบดีบัณฑิตวิทยาลัย

.....
คณะกรรมการสอบปากเปล่าสารนิพนธ์

.....
(อาจารย์ ดร.วีระ สอึ้ง)

ที่ปรึกษาหลัก

..... ประธาน
(อาจารย์ ดร.รัตนชัยนันท์ ธรรมสุจริต)

..... กรรมการ
(ผู้ช่วยศาสตราจารย์ ดร.วิรัช เจริญเรืองกิจ)

ชื่อเรื่อง	การศึกษาอารมณ์เพลงเพื่อพัฒนาธุรกิจ
ผู้วิจัย	กรกัญญา ศิริเกษ
ปริญญา	วิทยาศาสตร์มหาบัณฑิต
ปีการศึกษา	2563
อาจารย์ที่ปรึกษา	อาจารย์ ดร. วีระ สอิ่ง

ปัจจุบันผู้คนมักจะฟังเพลงตามอารมณ์ของเพลง การศึกษานี้มีวัตถุประสงค์เพื่อวิเคราะห์การแยกเพลงโดยใช้การประมวลผลภาษาธรรมชาติ (NLP) เพื่อรับข้อมูลอารมณ์จากเนื้อเพลง ชาร์ตเพลงต่างๆ เช่น Billboard สามารถใช้เพื่อค้นหาเพลงและอัลบั้มที่ได้รับความนิยมสูงสุดจากทุกประเภทแนวเพลง ขั้นแรกอักขระพิเศษและใช้ term-frequency / inverse-document frequency (TFIDF) จากนั้นใช้ Latent Dirichlet Allocation (LDA) เพื่อเชื่อมต่อกับคำกับฉลากอารมณ์ เราทำการจำแนกอารมณ์ตามบทเพลงบนตัวจำแนกการเรียนรู้ของเครื่องในท้องถิ่นเช่น Random forest, Decision tree, Naïve Bayes, Logistic Regression, AdaBoost และ XGBoost หลังจากใช้การค้นหาแบบกริดเพื่อปรับอัตราพารามิเตอร์ที่ดีที่สุดของแบบจำลอง ผลการวิจัยพบว่า XGBoost แสดงระดับความแม่นยำสูงสุด สามารถพิสูจน์ได้ว่า XGBoost มีประสิทธิภาพที่ดีกว่าการเรียนรู้ของเครื่องในท้องถิ่นในงานวิจัยนี้ สุดท้ายนี้เรายังแสดงกระจายอารมณ์ของเพลงที่มีอยู่ในชาร์ตบิลบอร์ดระหว่างปี 2000 ถึง 2017

คำสำคัญ : อารมณ์, เนื้อเพลง, การเรียนรู้ของเครื่องจักร, อัลกอริทึม XGBoost, บิลบอร์ด, การประมวลผลภาษาธรรมชาติ

Title	THE MOODS OF THE LYRICS IN THE FUTURE FOR BUSINESS
Author	KORNKANYA SIRIKET
Degree	MASTER OF SCIENCE
Academic Year	2020
Thesis Advisor	Vera Sa-ing , Ph.D.

Nowadays, people tend to listen to music based on the mood of the tracks. The purpose of this study is to analyze song extraction using Natural Language Processing (NLP) to acquire mood information from the lyrics. Various music charts, such as the Billboard, can be used to find the most popular songs and albums across all genres. First, removing special characters and using term-frequency/inverse-document frequency (TFIDF) and then Latent Dirichlet Allocation (LDA) were used to connect words to mood classes. We perform a lyric-based mood classification on local machine learning classifiers such as Random forest, Decision tree, Naïve Bayes, Logistic Regression, AdaBoost and XGBoost. After using a grid search for tuning the best parameter yields, the results revealed that XGBoost showed the highest level of accuracy. It can prove that boosting algorithms had a better performance than local machine learning in this research. Finally, we also provided the mood distribution of songs available on the Billboard chart between 2000 and 2017.

Keyword : Mood, Song Lyric, Machine Learning, Boosting algorithm, Billboard, Natural language processing

กิตติกรรมประกาศ

สารนิพนธ์นี้สำเร็จลุล่วงได้ด้วยความช่วยเหลือจาก ดร.วีระ สะอั้ง อาจารย์ที่ปรึกษาที่ให้คำปรึกษา คำแนะนำในการทำสารนิพนธ์ ตลอดจนสนับสนุนข้อมูลทางวิชาการ และข้อมูลสำหรับทำสารนิพนธ์นี้

ขอกราบขอบพระคุณคณะกรรมการสอบสารนิพนธ์ที่ได้ให้คำแนะนำและข้อเสนอแนะสำหรับการปรับปรุงสารนิพนธ์

ขอกราบขอบพระคุณบัณฑิตวิทยาลัย มหาวิทยาลัยศรีนครินทรวิโรฒ สำหรับทุนสนับสนุนการนำเสนอผลงานวิจัยของนิสิตบัณฑิตศึกษาทำให้ได้รับประสบการณ์นำเสนอสารนิพนธ์รวมถึงได้แลกเปลี่ยนความรู้ทางด้านเทคโนโลยีต่างๆ

กรกัญญา ศิริเกษ



สารบัญ

หน้า

บทคัดย่อภาษาไทย.....	ง
บทคัดย่อภาษาอังกฤษ	จ
กิตติกรรมประกาศ	ฉ
สารบัญ	ช
สารบัญตาราง.....	ญ
สารบัญรูปภาพ	ฎ
บทที่ 1 บทนำ.....	1
1.1 ความเป็นมาและความสำคัญของปัญหาการวิจัย.....	1
1.2 วัตถุประสงค์ของการวิจัย	2
1.3 ประโยชน์ของการวิจัย	2
1.4 ขอบเขตของการวิจัย.....	3
บทที่ 2 เอกสารและงานวิจัยที่เกี่ยวข้อง	6
บทที่ 3 วิธีดำเนินการวิจัย.....	11
3.1 ข้อมูลที่เอามาใช้ในงานวิจัย.....	11
3.1.1 ข้อมูล Billboard 2000-2018	11
3.1.2 ข้อมูล Moody Corpus	12
3.1.3 ข้อมูลรวมเดต้า.....	13
3.2 กระบวนการทำงานของงานวิจัย.....	14
3.3 การสำรวจข้อมูล.....	14
3.3.1 ข้อมูลของ Billboard	14
3.3.1 ข้อมูลของ Moody Corpus	16

3.3.3 ข้อมูลรวม	16
3.3.3.1 การทำความสะอาดข้อมูล	18
3.3.3.2 การทำ Word Cloud กับข้อมูลรวม.....	19
3.4 Feature Selection	20
3.4.1 TF-IDF	20
3.4.2 LDA (Latent Dirichlet Allocation)	21
บทที่ 4 การทดลอง.....	22
4.1 Model Classification.....	22
4.1.1 Traditional Machine Learning	22
4.1.1.1 Decision tree	22
4.1.1.2 Naïve bayes	23
4.1.1.3 Logistic Regression.....	24
4.1.1.4 Random Forest	24
4.1.2 AdaBoost (Adaptive Boosting Algorithm).....	24
4.1.3 XGBoost (Extreme Gradient Boosting).....	25
4.2 Grid Search.....	26
4.3 Model Evaluate.....	26
บทที่ 5 สรุปผลการทดลอง.....	27
5.1 สรุปผลการวิจัย.....	27
5.2 อภิปรายผลการวิจัย.....	28
5.3 ข้อเสนอแนะ.....	34
บรรณานุกรม.....	36
ประวัติผู้เขียน	40



สารบัญตาราง

หน้า

ตาราง 1 Feature Classification	3
ตาราง 2 RUSSEL'S EMOTION ADJECTIVES	4
ตาราง 3 ชื่อคอรัลล์นั้ ของ Billboard 2000-2018	11
ตาราง 4 คอรัลล์นั้ Lyrics ที่เลือกเอามาใช้ทำการทดลอง.....	12
ตาราง 5 ตาราง Moody Corpus	12
ตาราง 6 คอรัลล์นั้ Mood ที่เลือกเอามาใช้ทำการทดลอง.....	13
ตาราง 7 คอรัลล์นั้ที่เอามาใช้ ในการเทรน.....	13
ตาราง 8 จำนวนอารมณ์ในแต่ละอารมณ์ของข้อมูล Moody Corpus(Çano & Morisio, 2017) ..	16
ตาราง 9 จำนวนอารมณ์ในแต่ละอารมณ์ของข้อมูลรวม.....	17
ตาราง 10 แจกแจงเพลงมีความสุขและเพลงไม่มีความสุข	30

สารบัญรูปภาพ

หน้า

ภาพประกอบ 1 Russel's Model(Russell, 1980)	4
ภาพประกอบ 2 ตารางการแบ่งอารมณ์ทางจิตวิทยาทั้ง 23 ประเภท(Yang & Lee, 2009).....	8
ภาพประกอบ 3 การปรับมุมมองในการจำแนกอารมณ์จาก แบบจำลองพื้นฐานของ Plutchick (Oh et al., 2013).....	9
ภาพประกอบ 4 กระบวนการทำงานของงานวิจัย.....	14
ภาพประกอบ 5 จำนวนเพลงในแต่ละปี	15
ภาพประกอบ 6 กราฟแสดงการแจกแจงความถี่ของคำในเนื้อเพลงของข้อมูล Billboard.....	15
ภาพประกอบ 7 อารมณ์ในแต่ละปีข้อมูลรวม.....	18
ภาพประกอบ 8 กราฟแสดงการแจกแจงความถี่ของคำในเนื้อเพลงในข้อมูลรวม	18
ภาพประกอบ 9 ตาราง Word Cloud แสดงคำที่ใช้ถึง 10 อันดับในแต่ละอารมณ์.....	19
ภาพประกอบ 10 การแยกแยะผลลัพธ์ของ Decision Tree	23
ภาพประกอบ 11 หลักการทำ Random Forest.....	24
ภาพประกอบ 12 รูปเปรียบเทียบการทำงานระหว่าง Bagging และ Boosting	25
ภาพประกอบ 13 ผลการทำนายอารมณ์ของเพลงแต่ละเพลงแยกตามจำนวนอารมณ์ในแต่ละปี	29
ภาพประกอบ 14 กราฟแสดงร้อยละของอารมณ์เพลงในแต่ละปี.....	29
ภาพประกอบ 15 จำนวนเพลงที่มีความและไม่มีความสุขในแต่ละปี	33
ภาพประกอบ 16 กราฟแท่งเพลงมีความสุข.....	33
ภาพประกอบ 17 กราฟแท่งเพลงไม่มีความสุข	34

บทที่ 1

บทนำ

การสตรีมคือการถ่ายทอดสดออนไลน์ซึ่งช่วยเพิ่มความสามารถในการเข้าถึงเพลงที่อยู่ทั่วโลกได้หรือเพลงที่ยากฟังได้ และเว็บไซต์เนื้อเพลงเป็นสิ่งที่เข้าถึงได้ง่าย ทั้งสองอย่างนี้เข้าถึงได้ง่ายถ้าเรามีอินเทอร์เน็ต ถ้าในการถ่ายทอดออนไลน์หรือในเว็บไซต์เนื้อเพลงมีการจัดหมวดหมู่อารมณ์ของเพลง มันจะเป็นสิ่งที่ผู้ใช้งานเข้าถึงสิ่งนี้คือความแปลกใหม่ในการจัดหมวดหมู่เพลงและอารมณ์ของเพลง เป็นการรองรับฐานผู้ใช้งานที่หลากหลาย การจัดหมวดหมู่เพลงมีจุดมุ่งหมายเพื่อการจัดกลุ่มอารมณ์เพลงเข้าด้วยกันในแต่ละปี ประเภทศิลปินอัลบั้มอารมณ์ ฯลฯ และดูความนิยมของแต่ละอารมณ์เพลงเป็นหลักว่าอารมณ์แบบไหนจะสามารถเป็นความนิยมในอนาคต โดยการทำนายผ่านการเรียนรู้ของเครื่องจักรซึ่งพวกนี้รวดเร็วและมีประสิทธิภาพในการทำนาย สิ่งที่ทำนายออกมาเป็นอารมณ์ของเพลงสามารถอธิบายรสนิยมของผู้ฟังได้ดีขึ้น และจากการทำนายหมวดหมู่อารมณ์นี้ไม่คำนึงถึงประเภทศิลปินหรืออัลบั้มแต่สนใจที่เนื้อเพลงเป็นหลักนั้นเพื่อนำมาทำการทำนายอารมณ์เพลง นั้นจะสามารถแสดงให้เห็นได้ว่าเป็นการแสดงถึงรสนิยมทางการรับรู้เนื้อเพลงของผู้ฟังเพลงและแนวทางในการดำเนินธุรกิจไปในแนวทางใด

1.1 ความเป็นมาและความสำคัญของปัญหาการวิจัย

ดนตรีมีส่วนเกี่ยวข้องกับชีวิตประจำวันของมนุษย์มายาวนาน โดยเริ่มตั้งแต่ยุคก่อนประวัติศาสตร์ สิ่งหนึ่งที่ทำให้เกิดดนตรีขึ้นครั้งแรกคือ มันเกิดจากความหวาดกลัวจากการที่มีการเปลี่ยนแปลงของฤดูกาล เช่น พายุแลบ พายุร้อน ฝนตก แผ่นดินไหว เป็นต้นดังนั้น สิ่งพวกนี้สร้างความวิตกกังวลมนุษย์ยุคนั้นเป็นอันมาก แต่วัตถุประสงค์ที่สร้างดนตรีขึ้นมาก็เพื่อผ่อนคลายความตึงเครียดและผ่อนคลายไปกับอารมณ์ของแต่ละเพลงนั้น ๆ แต่ละคนมองเพลง ในมุมมองที่แตกต่างกันเช่น มุมมองด้านวัฒนธรรม มุมมองด้านความคิดและจินตนาการ มุมมองด้านอารมณ์ มุมมองความบันเทิง มุมมองธุรกิจ จึงทำให้ชีวิตของคนส่วนใหญ่มีดนตรีเข้ามาเกี่ยวข้องอยู่เสมอ เพื่อคลายความเครียด ความกลัว ความเหน็ดเหนื่อย หรือแม้กระทั่งความทุกข์ทั้งหลาย เพลงหรือดนตรีจึงมีความสำคัญและมีคุณค่าต่อมนุษย์เป็นอย่างมาก หรือถือว่ามันมีบทบาทในการดำเนินชีวิตของคนเราตั้งแต่เกิดจนตาย โดยทุกที่ทั่วโลกไม่ว่าจะได้ยินดนตรีจากที่ไหน ที่นั่นก็จะมีความรื่นเริง งานฉลองต่าง ๆ จะเห็นได้ว่าคนทั่วโลกจะมีความสัมพันธ์กับดนตรีหรือเพลงมาโดยตลอด

เมื่อดนตรีประกอบเนื้อเพลงจึงทำให้เกิดความสมบูรณ์มากขึ้นเนื่องจากองค์ประกอบของเพลงๆ หนึ่งจะประกอบไปด้วยเนื้อเพลงและทำนองดนตรี ในเนื้อเพลงนั้นก็ประกอบไปด้วย

อารมณ์ของแต่ละคำที่ผู้เขียนเนื้อเพลงได้ทำการกลั่นกรองออกมาจากอารมณ์ เราจึงทำการวิจัยขึ้นนี้โดยคำนึงของอารมณ์เพลงที่มาจากเนื้อเพลงเป็นหลัก ในข้อมูลนั้นมีการกำหนดฉลากเป็น 4 ฉลาก คือ เศร้า, ความสุข, ผ่อนคลาย และ โกรธ นั้นจึงเป็นประโยชน์สำหรับการแยกประเภทโดยการใช้เนื้อเพลง

ปัจจุบันการฟังเพลงผ่านอุปกรณ์โทรศัพท์ หรือ ผ่านคอมพิวเตอร์ โดยช่องทางออนไลน์เป็นที่นิยมในหมู่คนสมัยใหม่ มันทำให้การฟังเพลงแบบถ่ายทอด (Streaming) เช่น Spotify, Apple Music หรือ Tidal โปรแกรมประยุกต์พวกนี้กลายเป็นสิ่งที่เป็นที่นิยมและเป็นสิ่งแปลกใหม่และเข้าถึงง่ายในปัจจุบัน แต่เพราะการแสดงตัวเล็กแค่นวเพลง หรือนักร้อง หรืออัลบั้ม อาจจะไม่ใช่เพียงพอเนื่องจากมีผู้ใช้งานหลากหลาย จึงมีการจัดอันดับและจัดหมวดหมู่เพลงและอารมณ์เพลง สิ่งเหล่านี้ก็เพื่อจัดกลุ่มเพลงเข้าด้วยกัน ผู้คนสามารถฟังเพลงได้ตั้งใจ อาจจะฟังเพลงตามอารมณ์ เพื่อให้รู้สึกได้เข้าถึงอารมณ์ของตัวเองอย่างแท้จริง

1.2 วัตถุประสงค์ของการวิจัย

เพื่อศึกษาความนิยมของแต่ละอารมณ์เพลงในแต่ละปี ในข้อมูล Billboard ว่าอารมณ์แบบไหนมีความนิยมในปียears โดยเราจะทำการทำนายอารมณ์เพลงจากเนื้อเพลง โดยใช้การเรียนรู้ของเครื่องจักรที่มีประสิทธิภาพมาก และศึกษาการทำงานของแต่ละแบบจำลอง เมื่อเราได้แบบจำลองที่ทำความแม่นยำ เราจะนำไปทำนายอารมณ์เพลงกับข้อมูลเพลง ซึ่งผลลัพธ์ที่ออกมาเป็นสิ่งที่สามารถอธิบายรสนิยมของผู้ฟังได้ดีขึ้น การจัดหมวดหมู่อารมณ์นี้ไม่คำนึงถึงประเภทศิลปินและอัลบั้มแต่สนใจที่เนื้อเพลงเป็นหลักนั้นเป็นการแสดงถึงรสนิยมและอารมณ์ทางการรับรู้เนื้อเพลงของผู้ฟัง สิ่งนี้จะเป็ประโยชน์และสามารถนำมาทำนายอารมณ์เพลงที่มีความนิยมในอนาคตได้ เมื่อใช้ในทางธุรกิจเราสามารถนำสิ่งเหล่านี้จัดหมวดหมู่ในโปรแกรมประยุกต์ (Application) ที่เอาไว้ถ่ายทอด (Stream) เพลง และเว็บไซต์ (Website) ในการวิเคราะห์อารมณ์เพลงวิเคราะห์โดยการใช้ การประมวลภาษาธรรมชาติ (NLP) ซึ่งทำการประมวลภาษาคำศัพท์ในเนื้อเพลงที่มีฉลากอารมณ์กำกับอยู่เพื่อเอาไปใช้ดูความนิยม

1.3 ประโยชน์ของการวิจัย

ประโยชน์ของงานวิจัยชิ้นนี้จัดทำขึ้นมาเพื่อวิจัยในความสามารถของการเรียนรู้ของเครื่องจักรเพื่อค้นหาแบบจำลองที่ดีที่สุดเพื่อใช้ในการวิเคราะห์จัดจำแนกเนื้อเพลง ในเชิงธุรกิจวงการเพลงจะเราจะเห็นได้ว่าเพลงไหนที่อยู่ในกระแสและเพลงไหนอยู่นอกกระแส และเราสามารถนำเนื้อเพลงในแต่ละปีจนถึงปีล่าสุดมาทำการจัดจำแนกอารมณ์ประเภทโดยใช้เนื้อเพลง

ซึ่งจะสามารถคิดได้ว่าต่อไปจะมีความนิยมไปในแนวอาร์มไปแนวทางไหน และประโยชน์ในอีกทางหนึ่งก็คือการจัดจำแนกเพลงของเว็บไซต์(Website) หรือโปรแกรมประยุกต์(Application) โดยแบ่งแยกอารมณ์ของเพลงแต่ละเพลงผ่านเนื้อเพลง การจัดจำแนกประเภทนี้จะเป็นผลดีต่อผู้ที่เป็นเจ้าของโปรแกรมประยุกต์(Application) และ เว็บไซต์(Website) สามารถทำงานได้ง่ายและรวดเร็วขึ้นขึ้น ทำให้ผู้เข้าใช้งานสามารถมีทางเลือกในการอ่านหรือฟังเพลงว่าจะดูเพลงอารมณ์ไหน

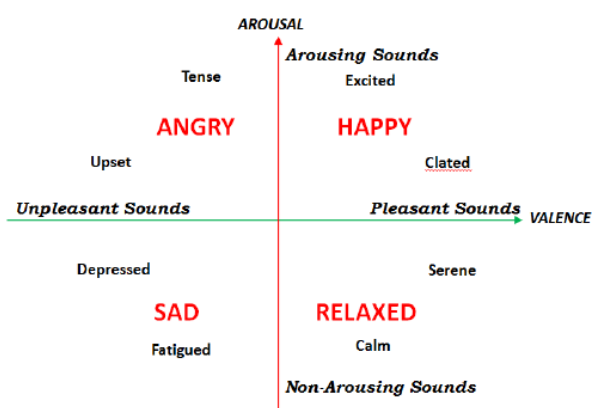
1.4 ขอบเขตของการวิจัย

การศึกษาค้นคว้าครั้งนี้มุ่งศึกษาเฉพาะข้อมูลใน Billboard ปี 2000-2018 จากเว็บไซต์เดต้าเวิร์ล (Tauberg, 2018) คือข้อมูลการจัดอันดับจัดอันดับเพลงและอัลบั้มยอดนิยมตามแนวเพลงในแต่ละสัปดาห์ และ Moody Corpus (Çano & Morisio, 2017) คือ ข้อมูลที่มีชื่อเพลง,ชื่อนักร้อง และอารมณ์ของเพลง และข้อมูลที่ทำกรรวมข้อมูลระหว่าง Billboard ปี 2000-2018 และ Moody Corpus ซึ่งได้กำหนด Column ที่จะเอามาทำการทดลอง ดังตารางที่ 1

ตาราง 1 Feature Classification

ตัวแปร	ความหมาย
Lyrics	เนื้อร้อง
Moods	อารมณ์ของเพลง

ตัวแปร Lyrics และ Moods จะมีกำหนดให้ในข้อมูลที่เรามี ตัวแปรที่เราจะเอามาทำเป็น Label เราจะใช้ตัว Mood ซึ่งมี 4 ฤดูกาล ประกอบด้วย Angry, Happy, Sad และ Relaxed ดังภาพประกอบที่ 1 คือ แบบจำลองรัชเชล



ภาพประกอบ 1 Russel's Model(Russell, 1980)

ตาราง 2 RUSSEL'S EMOTION ADJECTIVES

Mood Label	Emotion
Relaxed	sleepy, calm, relaxed, satisfied, content, at ease, serene, glad, pleased
Happy	aroused, astonished, excited, delighted, happy
Sad	miserable, sad, gloomy, depressed, bored, droopy, tired
Angry	alarmed, tense, angry, annoyed, afraid, distressed, frustrated

เราใช้แบบจำลองรัสเซลล์(Russell, 1980)ในการกำหนดเพราะเป็นแบบจำลองสภาวะทางอารมณ์ที่ส่งผลกระทบอย่างรอบด้าน สภาวะอารมณ์ที่เกิดขึ้นนั้นเกิดจากประสาทสัมผัสวิทยาพื้นฐานสองระบบคือ Arousal(ความตื่นตัวหรือการกระตุ้นอารมณ์) และ Valence(ความพึงพอใจและความไม่พึงพอใจ) การตื่นตัวอย่างในแนวแกนตั้งและความพึงพอใจอยู่ในแนวนอน โดยให้เส้นแบ่งจากตามแกนแนวทั้งสองในการมองแต่ละสภาวะอารมณ์ ถ้ามาทางแกนนอนทางซ้ายบอกได้ว่ามีความรู้สึกที่พอใจ ถ้าทางแกนนอนไปทางซ้ายบอกได้ว่ามีความเป็นความรู้สึกไม่พอใจ และถ้าแกนตั้งไปทางด้านบนบอกได้ว่าเป็นความรู้สึกตื่นเต้น ถ้าแกนตั้งไปทางด้านล่างบอกได้ว่าเป็นความรู้สึกธรรมดาไม่ตื่นเต้น สิ่งนี้เองถูกนำมาใช้แพร่หลายในการทำ Emotion Word

สำหรับแบบจำลองในการจัดจำแนกอันดับของคำในแต่ละอารมณ์ในแต่ละปี ผู้วิจัยจึงใช้เทคนิค Supervised Learning Classification โดยเลือกใช้เทคนิค Random Forest, Logistic Regression, Decision Tree, Naïve Bayes, XGBoost และ AdaBoost โดยแบ่งชุดข้อมูลเรียนรู้เป็น 80%และชุดข้อมูลทดสอบเป็น 20% และใช้ภาษา Python ในการสร้างแบบจำลองการทำนาย

บทที่ 2

เอกสารและงานวิจัยที่เกี่ยวข้อง

การจำแนกอารมณ์ตามเนื้อเพลงได้รับการศึกษามากมายกว้างขวางในไม่กี่ปีมานี้ (An, Sun, & Wang, 2017) นำเสนอวิธีการจำแนกเพลงจีนตามอารมณ์ของเพลงโดยมีการรวบรวมจากเว็บไซต์ Baidu และได้ใช้การกำหนดแท็กอารมณ์(Tag)จากผู้ใช้งานบนเว็บไซต์ที่หลากหลายเข้ามาช่วยกันติดฉลากอารมณ์ของแต่ละเพลงผู้วิจัยใช้แบบจำลองอารมณ์ของรัสเซล(Russell, 1980) แบ่งอารมณ์ของเพลงแบ่งออกเป็น 3 ประเภท ได้แก่ depression(ความหดหู่), contentment(ความพึงพอใจ) และ exuberance(ความอุดมสมบูรณ์ของชีวิต) โดยใช้อัลกอริทึม Naïve Bayes ในการจัดประเภทข้อมูลที่เป็นข้อความและให้ความแม่นยำเท่ากับ 68% ในการศึกษาครั้งนี้ด้วยในงานวิจัยนี้ยังมีการใช้รูปแบบอารมณ์ของรัสเซลด้วยการจำแนกอารมณ์ออกเป็น 4 ประเภท ได้แก่ Happy, Sad, Angry and Relaxed

(Akella & Moh, 2019)เสนอแบบจำลองการจำแนกอารมณ์สำหรับเพลงภาษาอังกฤษที่รวบรวมเพลงจากเว็บไซต์ Million Song และรวมกับการแท็กหรือการติดฉลากอารมณ์จากเว็บไซต์ Lastfm ในการศึกษาครั้งนี้ใช้เทคนิคต่างๆ ในการแปลงจากข้อความเป็นเวกเตอร์ตัวเลขนอกจากนี้ การศึกษาการเปรียบเทียบผลการจำแนกประเภทจากเทคนิคการเรียนรู้ของ Local Machine learning เช่น Support vector machine (SVM), Random forest, Linear regression และ Naive Bayes ด้วยเทคนิคการเรียนรู้เชิงลึกเช่น Convolutional Neural Network (CNN) แสดงให้เห็นว่า CNN ให้ความแม่นยำสูงสุด 71%

การจัดหมวดหมู่อารมณ์ที่ได้เข้าถึงเทคนิคการประมวลภาษาธรรมชาติ ที่มี Label ของอารมณ์ที่มาจากทฤษฎีอารมณ์ที่เกี่ยวข้องกับ Valance และ Arousal ของ Russel Model (Russell, 1980) และบางวิธีได้ใช้การแนะนำตามอารมณ์ (F. Rachman, Samo, & Fatichah, 2018) ในการใช้แบบจำลองใหม่ล่าสุดสำหรับการกระจายอารมณ์ เรียกว่า Emotion Music Scale อารมณ์ที่เกิดจากดนตรีเพียงอย่างเดียวจะถูกแยกและจัดกลุ่มเป็นป้ายกำกับและป้ายเหล่านี้ถูกนำมาสร้างระบบที่เรียกว่า MoodPlay (Andjelkovic, Parra, & O'Donovan, 2019) MoodPlay ให้ความสำคัญคล้อยถึงกันโดยการใช้ฉลากกำกับจาก GEMS(The Geneva Emotional Music Scale)

มีการเก็บ Mood Label จาก เว็บไซต์ last.fm.com จะทำการจับคู่อารมณ์ในแต่ละเพลง และ เลยทำการเก็บไว้ในรูปของโมเดล Russel (Akella & Moh, 2019) และนำเอาข้อมูลมาเทรน โดย ใช้ Traditional Machine learning(SVM, Random Forest, KNN, MLP, Logistic Regression, N.Bays) ซึ่ง Random forest ได้ค่า Accuracy เบื้องต้นได้ 60.11%

แบบจำลองที่ศึกษาการตรวจจับอารมณ์ของเพลงโดยใช้เสียงและเนื้อเพลง (F. H. Rachman, Sarno, & Fatichah, 2020) ในงานวิจัยนี้ใช้แบบจำลองเทเยอร์ (Thayer model) ในการแบ่งเขตอารมณ์ออกเป็น 4 ส่วนของอารมณ์ Q1, Q2, Q3 และ Q4 ใน 4 พื้นที่นี้ งานวิจัยชิ้นนี้ใช้ฉลากอารมณ์อ้างอิงถึง พื้นที่ที่กระตุ้นอารมณ์ (arousal areas) ประกอบไปด้วย ความเร้าอารมณ์ต่ำ (low arousal) และความเร้าอารมณ์ (high arousal) และนำเนื้อเพลงและเสียงเพลงช่วงกลางเพลง (chorus segments) มาใช้เป็นลักษณะเฉพาะในการเอามาเข้าแบบจำลอง โดยใช้ คุณสมบัติถ่วงน้ำหนักในแต่ละแบบจำลอง คือ Naive Bayes , Support Vector Machine (SVM), Random Forest

การสร้างกรอบการทำงานเพื่อระบุการติดฉลากแบบใหม่เพื่อการติดฉลากในสื่อสังคมออนไลน์ได้ถูกต้อง (Hu & Ogihara, 2012) ในกรอบงานของงานวิจัยอันดับแรกจะทำการลบและกรองฉลากเพลง จากนั้นจะใช้อัลกอริทึมในการจัดกลุ่มแบบลำดับขั้นที่ปรับปรุงแล้วจัดกลุ่มฉลากเพื่อสร้างหมวดหมู่ฉลาก ในงานวิจัยชิ้นนี้ใช้แบ่งประเภทเพลงโดยใช้เนื้อเพลงตามหมวดหมู่ในการขยายข้อมูลความหมายของเนื้อเพลงงานวิจัยนี้ใช้ CLOPE เพื่อจัดกลุ่มเนื้อเพลง และใช้ เซนทรอยด์ของคลัสเตอร์ที่เกี่ยวข้องแสดงเนื้อเพลง ตามวิธีการของ Naive Bayes ผลจากการจัดประเภทเพื่อพิจารณาว่าฉลากที่ใช้กำหนดเพลงถูกต้องหรือไม่ จากผลการทดลองสรุปได้ว่ามีประสิทธิภาพมาก

งานวิจัยที่มุ่งเน้นไปที่การระบุอารมณ์ โดยมาจากประสบการณ์ของผู้คนในขณะที่ฟังเพลง (Chauhan & Chauhan, 2016) มุ่งหมายที่จะวิเคราะห์เนื้อเพลง ของเพลงที่ใช้ภาษาฮินดี ตรวจจับอารมณ์ของผู้ฟัง ในงานวิจัยนี้ใช้ Unigram และ ความถี่ของคำ (TF) เป็นคุณสมบัติหลัก และเราใช้ การ ลดมิติ ใช้เฉพาะคำที่เกี่ยวข้องในการตรวจจับอารมณ์ ซึ่งใช้ แบบจำลอง unsupervised คือ topic-modeling (Latent Dirichlet Allocation model) เพื่อทำการค้นหาอารมณ์จากทุกเพลงไหน คลังที่เรามี และจะได้ทำการสร้างชุดข้อมูลจำนวน 1900 เพลง ที่เป็นภาษาฮินดี และจากการทดลองนี้ ที่ ใช้ การทำ LDA (Latent Dirichlet Allocation model) ปรากฏว่าผลลัพธ์ที่ออกมาค่อนข้างที่จะดี

งานวิจัยที่จะใช้การเรียนรู้ของเครื่องแทนความเชี่ยวชาญของมนุษย์ในการที่จะจัดจำแนกอารมณ์ของเพลงออกมา (Yang & Lee, 2009) การจัดประเภทขึ้นอยู่กับการจัดจำแนกทางจิตวิทยาของอารมณ์ ซึ่งแบ่งอารมณ์เป็น 23 ประเภทอย่างเจาะจงผลลัพธ์ของงานวิจัยชิ้นนี้จากการที่วิเคราะห์ข้อมูลของเพลงอย่างเฉพาะเจาะจงนั้นมีแนวโน้มที่ดี ที่แบบจำลองจำแนกผลลัพธ์ออกมา ในประเภทที่มนุษย์สามารถเข้าใจได้และสร้างผลลัพธ์ที่สอดคล้องกับสัญชาตญาณทั่วไปเกี่ยวกับ

อารมณ์โดยงานวิจัยชิ้นนี้ใช้แบบจำลอง Decision tree มีในการจำแนกประเภทเพื่อทำความเข้าใจกับแบบจำลองการจำแนกประเภทในแต่ละประเภท ผลลัพธ์ของเราในการหาคำอธิบายความโคลงสั้นๆ สำหรับอารมณ์ที่เฉพาะเจาะจงนั้นมีแนวโน้มที่จะสร้างแบบจำลองการจำแนกประเภทที่เข้าใจได้โดยมนุษย์และสร้างผลลัพธ์ที่สอดคล้องกับสัญชาตญาณทั่วไปเกี่ยวกับอารมณ์ที่เฉพาะเจาะจงได้

Positive emotions		Negative emotions	
Allmusic	EMO	Allmusic	EMO
Brash, Bravado, Swaggering	Proud	Acerbic, Aggressive, Bitter, Fiery, Nasty, Outraged, Rebellious, Snide	Anger
Bright, Effervescent, Gleeful, Humorous, Irreverent, Party/Celebratory, Rambunctious, Raucous, Rollicking, Silly, Sweet, Whimsical, Witty	Joy	Aggressive	Aggressive
Yearning, Sexy, Provocative, Sleazy, Sensual	Arousal	Bleak, Distraught, Sombre, Poignant, Melancholy, Plaintive	Sad
Exuberant	Exuberant	Gloomy	Gloomy
Passionate	Passionate	Cynical, Ironic	Alienation
Happy	Happy	Manic, Paranoid, Spooky, Unsettling	Fear
Wistful	Reflective	Ominous	Ominous
Romantic, Precious, Reverent, Sentimental, Warm	Love	Eerie	Eerie
Soothing	Soothing	Tense	Tense
Laidback/Mellow, Gentle, Refined/Mannered, Reserved, Restrained	Calm	Greasy	Disgust
Confident	Confident	Eccentric, Fractured, Trippy, Freakish, Spacey	Psychotic
Detached	Reflective	Aggressive, Boisterous, Brittle, Brooding, Cold, Confrontational, Harsh, Wry, Dramatic, Theatrical, Volatile, Thuggish, Trashy, Visceral, Earthy, Rustic	
Complex, Enigmatic, Eccentric, Street-smart, Stylish, Earnest, Ramshackle, Smooth, Slick, Sophisticated, Gritty, Spiritual, Searching, Playful, Rousing, Free-wheeling			

ภาพประกอบ 2 ตารางการแบ่งอารมณ์ทางจิตวิทยาทั้ง 23 ประเภท (Yang & Lee, 2009)

(Oh, Hahn, & Kim, 2013) นำเสนอถึงวิธีการจำแนกอารมณ์ของเพลงใหม่ โดยเลือกใช้เฉพาะเนื้อเพลงท่อนแรกและบางส่วนของเนื้อเพลงเพื่อจัดประเภทเพลงเนื่องจากเนื้อเพลงท่อนแรกจะสร้างบรรยากาศของเพลงแต่ละเพลงและในส่วนของเนื้อเพลงจะมีคำหลักที่สำคัญที่สุดของเพลง ดังนั้น แนวทางและสิ่งที่จะนำมาเสนอนี้ เพื่อที่จะเพิ่มความแม่นยำในการจัดจำแนกประเภทโดยที่ไม่คำนึงถึงคำที่ไม่มีความหมาย งานวิจัยนี้แบ่งอารมณ์ ได้ 8 แบบ คือ Joy, Acceptance, Anticipation, Anger, Disgust, Sadness, Surprise, and Fear มาจากแบบจำลองอารมณ์พื้นฐานของ Plutchick และเพื่อที่จะวิเคราะห์ความคล้ายคลึงกันระหว่างเนื้อเพลงและอารมณ์วิธีนี้จะพิจารณาความสัมพันธ์ระหว่างคำสั่งและกำหนดอารมณ์ของแต่ละอารมณ์ เป็นหมวดหมู่ตามภาพประกอบ 3

<i>User's mood</i>	<i>Happy</i>	<i>Angry</i>	<i>Sad</i>
Recommended mood of songs	joy + acceptance + anticipation	anger + disgust	case1 Sadness + fear + surprise case2 joy + acceptance + anticipation

ภาพประกอบ 3 การปรับมุมมองในการจำแนกอารมณ์จาก แบบจำลองพื้นฐานของ Plutchick (Oh et al., 2013)

เทคนิคการจำแนกอารมณ์ที่รวมคุณสมบัติของบทกวีและเสียงสำหรับ Music Information Retrieval (MIR) ได้รับการศึกษาอย่างกว้างขวางในช่วงไม่กี่ปีที่ผ่านมา (Ying, Doraisamy, & Abdullah, 2012) บทความนี้จะตรวจสอบการแสดงของประเภทดนตรีและการจำแนกอารมณ์โดยใช้คุณสมบัติของเนื้อเพลงเท่านั้น ในการศึกษาเบื้องต้นนี้คุณลักษณะ Part-of-Speech (POS) ใช้สำหรับการจำแนกคอลเลกชันเพลง 600 เพลง ประเภทดนตรีและอารมณ์สิบประเภทได้รับการคัดเลือกตามลำดับโดยพิจารณาจากบทสรุปจากวรรณกรรม อารมณ์ของเนื้อเพลงแต่ละเพลงนั้นๆ จัดจำแนกมาจากแบบจำลอง Russell งานวิจัยนี้ข้อมูล K-NN, Naïve Bayes และ SVM ปรากฏว่า ค่า Accuracy ที่ดีที่สุดสำหรับ งานวิจัยนี้คือ SVM

แยกคุณลักษณะจากเนื้อร้อง คลังเพลงเพื่อจัดการการผลที่จะมาเฉพาะเนื้อเพลงและเสนอวิธีการประเมินอารมณ์จากวลีแต่ละบท (Matsumoto & Sasayama, 2018) ทำมาเพื่อแก้ปัญหาเพลงบางเพลงที่ไม่มีอารมณ์เพลงกำกับอยู่ และประเมินประสิทธิภาพของแต่ละชุดของคลังข้อมูลที่ใช้สำหรับการเรียนรู้คำเวกเตอร์ (word vector learning) และขั้นตอนวิธีการฝึกอบรมเวกเตอร์คำ (the corpus word vector training algorithm) หรือ คำสั่งอัลกอริทึม การประเมินอารมณ์ (emotion estimation algorithm) โดยการทดลองการประมาณอารมณ์จำนวนหนึ่ง ผลการทดลองประเมินของเรายืนยันประสิทธิภาพของวิธีการที่นำเสนอ (lyric corpus word vector + 3-NN) เนื่องจากมันสร้างค่า F-measure สูงกว่าวิธีอื่นๆ เราพบว่าวิธีการใช้คลังข้อมูลเนื้อเพลง (lyric corpus) สามารถประเมินอารมณ์ได้อย่างแม่นยำในระดับสูงโดยใช้วิธี เพื่อนบ้านที่ใกล้ที่สุด

(K- Nearest Neighbor method)

ในการศึกษานี้แทนที่จะใช้การเรียนรู้เชิงลึกเราได้เพิ่มความแม่นยำของผลการจำแนกโดยใช้เทคนิค Ensemble เช่น วิธี Boosting ซึ่ง อัลกอริทึม Boosting กำลังได้รับความนิยมและ

แสดงผลความแม่นยำได้ดีในการจัดประเภทต่างๆในปัจจุบัน (Al-Salemi, Mohd Noah, & Ab Aziz, 2016; Elmitwally, 2020; Salminen et al., 2020)



บทที่ 3

วิธีดำเนินการวิจัย

3.1 ข้อมูลที่เอามาใช้ในงานวิจัย

3.1.1 ข้อมูล Billboard 2000-2018

ข้อมูล Billboard ปี 2000-2018 จากเว็บไซต์เดต้าเวิร์ล (Tauberg, 2018) Billboard คือ ข้อมูลการจัดอันดับเพลงและอัลบั้มยอดนิยมตามแนวเพลงในแต่ละสัปดาห์ ซึ่งประกอบด้วย เพลงจำนวน 7,573 เพลง ในข้อมูล Billboard จะมีข้อมูลที่ซ้ำกันเนื่องจากการจัดอันดับเป็นรายสัปดาห์ หลังจากจัดเนื้อเพลงที่ซ้ำ จึงเหลือเพลง 5,743 เพลง จาก นักร้อง 1,386 คน จาก ปี 2000 ถึงปี 2018 ซึ่งจะมีรายละเอียด ดังตารางที่ 3

ตาราง 3 ชื่อคอลัมน์ ของ Billboard 2000-2018

Field name			
year	analysis_url	rank	danceability
title	energy	change	key
Simple title	liveness	spotify_link	duration_ms
artist	tempo	spotify_id	loudness
main artist	speechiness	video_link	valence
peak_pos	acousticness	Genre	mode
last_pos	instrumentalness	bard_genre	lyrics
weeks	time_signature		

จากตารางที่ 3 คอลัมน์ Billboard มีจำนวนทั้งหมด 30 คอลัมน์ งานวิจัยนี้เราเลือกคอลัมน์ title ของข้อมูล Billboard ไปจับกับ ข้อมูล title ของข้อมูล Moody corpus เพื่อทำการจับเพลงเหมือนกัน และจะนำคอลัมน์ Lyrics ทำไปใช้ต่อไปในแบบจำลอง

ตาราง 4 คอลัมน์ Lyrics ที่เลือกเอามาใช้ทำการทดลอง

Field name	ความหมาย
Lyrics	เนื้อเพลง

3.1.2 ข้อมูล Moody Corpus

ซึ่งเป็นข้อมูล และอารมณ์ของเพลงโดย Erion Çano (Çano & Morisio, 2017) ประกอบไปด้วยอารมณ์ 4 แบบคือ Angry, Sad, Relaxed และ Happy จากเพลงทั้งหมด 2,595 นักร้องทั้งหมด 1,665 คน ตามชุดข้อมูลนี้มีขนาด 3 คอลัมน์ ตามรายละเอียดตารางที่ 5

ตาราง 5 ตาราง Moody Corpus

Field name	ความหมาย
Artist	นักร้อง
Title	ชื่อเพลง
Mood	อารมณ์ของเพลง

งานวิจัยนี้เราเลือกคอลัมน์ title ของข้อมูล Moody Corpus ไปจับกับ ข้อมูล title ของข้อมูล Billboard เพื่อทำการจับเพลงเหมือนกัน และจะนำคอลัมน์ Mood ทำไปใช้ต่อไปในแบบจำลอง

ตาราง 6 คอลัมน์ Mood ที่เลือกเอามาใช้ทำการทดลอง

Field name	ความหมาย
Mood	อารมณ์ของเพลง

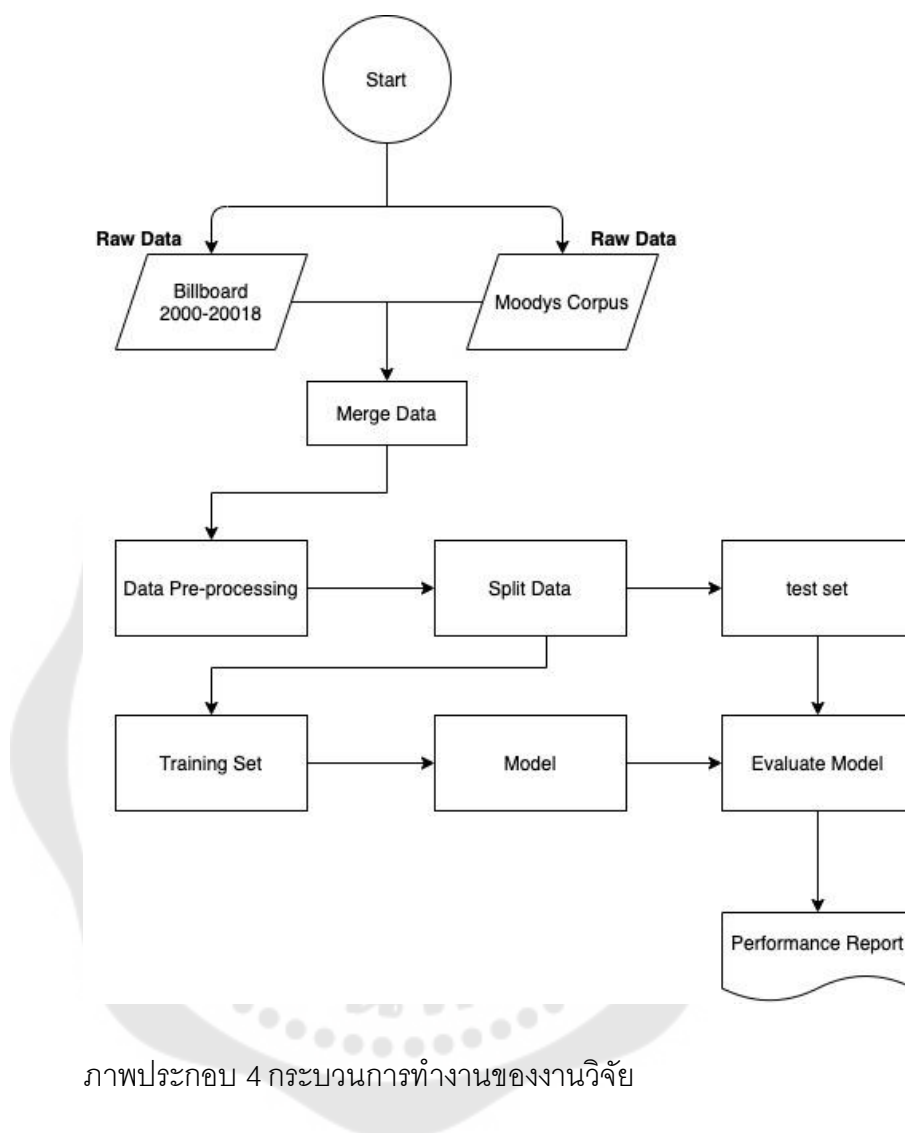
3.1.3 ข้อมูลรวมเดต้า

ซึ่งเป็นข้อมูลที่น่าข้อมูล Billboard และ Moody Corpus (Çano & Morisio, 2017) มาทำการจับชื่อเพลงที่เหมือนและเราเลือกที่จะนำมาใช้เพียง 2 คอลัมน์ จึงได้ข้อมูลทั้งหมด 209 แถว ในข้อมูลรวมเดตตานี้จะมี 2 คอลัมน์ดังตาราง 6 คอลัมน์ Lyrics

ตาราง 7 คอลัมน์ที่เอามาใช้ ในการเทรน

Field name	ความหมาย
lyrics	เนื้อเพลง
Mood	อารมณ์ของเพลง

3.2 กระบวนการทำงานของงานวิจัย



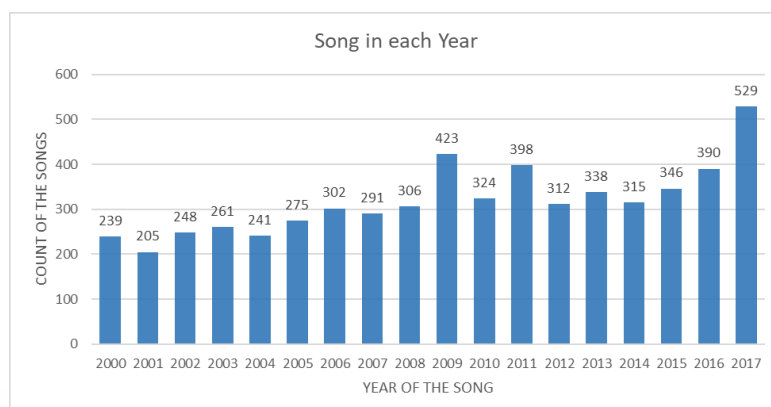
จากภาพกระบวนการทำงานของงานวิจัย เมื่อเรามีข้อมูลดิบ 2 ข้อมูล เราย้นำข้อมูลทั้ง 2 ข้อมูลมาทำการรวมกันได้ 209 แถว และทำการเตรียมข้อมูล ทำการแบ่งข้อมูลออกเป็น 2 ส่วน ข้อมูลฝึก(Training data) และข้อมูลทดสอบ(test data) และนำข้อมูลมาทำการประเมินจนได้ค่าแสดงประสิทธิภาพของแบบจำลอง

3.3 การสำรวจข้อมูล

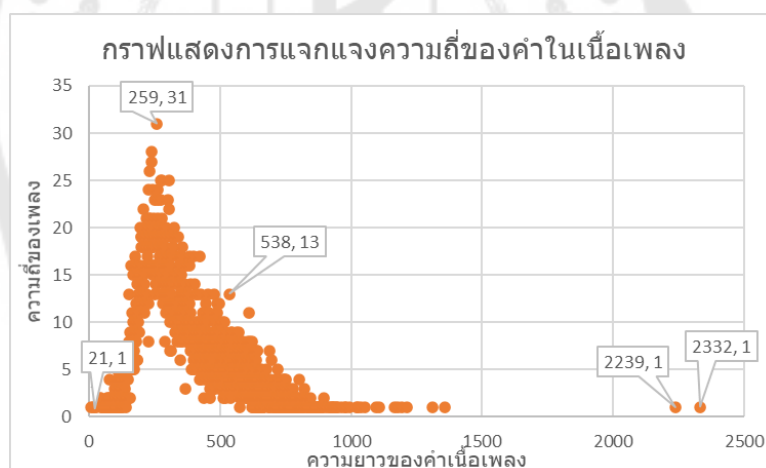
3.3.1 ข้อมูลของ Billboard

จากข้อมูลเพลงในแต่ละปีของ Billboard หลังจากได้ทำการทำความสะอาดข้อมูลที่มีเพลงซ้ำกันทำให้เหลือข้อมูลเหลือแค่ปี 2000-2017 เราพบว่าในปี 2017(Tauberg, 2018) มี

เพลงเยอะมากที่สุดคือ 529 เพลง รองลงมาคือ ปี 2009 มีเพลง 423 เพลง และปี 2011 มีเพลง 398 เพลง ตามลำดับ และในปีที่มีเพลงน้อยที่สุดคือปี 2001 ซึ่งมีเพลง 205 เพลง ดังภาพประกอบที่ 3



ภาพประกอบ 5 จำนวนเพลงในแต่ละปี



ภาพประกอบ 6 กราฟแสดงการแจกแจงความถี่ของคำในเนื้อเพลงของข้อมูล Billboard

จากภาพประกอบที่ 4 เป็นการแจกแจงความถี่ของคำในเนื้อเพลงแต่ละเพลงของข้อมูล Billboard พบว่าช่วงความถี่ค่าน้อยที่สุดคือ 21 คำ พบ 1 เพลง ความถี่ค่าที่พบเยอะที่สุดคือ 259 พบ 31 เพลง และเมื่อไล่ระดับความเยอะขึ้นไปเรื่อย ๆ ความถี่ของคำที่ 538 พบ 13 เพลง เพลงความถี่ของคำเนื้อเพลงที่เยอะที่สุด 2 อันดับ 2,239 และ 2,332 พบ อย่างละ 1 เพลง จากกราฟสรุปได้ว่าส่วนมากเพลงความถี่ของคำในเนื้อเพลงถ้ายิ่งเยอะมากเพลงจะยิ่งมีน้อย

3.3.1 ข้อมูลของ Moody Corpus

มีเพลงทั้งหมด 2,595 นักร้องทั้งหมด 1,665 คน เมื่อทำการนับจากอารมณ์เพลงแต่ละเพลง จากตาราง 7 เห็นได้ว่าอารมณ์ happy มีเยอะที่สุดในจำนวน 819 เพลง ตามมาด้วยอารมณ์ sad และอารมณ์ relaxed อารมณ์ที่มีจำนวนน้อยสุดคืออารมณ์ angry มีจำนวน 575 เพลง

ตาราง 8 จำนวนอารมณ์ในแต่ละอารมณ์ของข้อมูล Moody Corpus (Çano & Morisio, 2017)

อารมณ์ของเพลง	จำนวนของอารมณ์เพลง
angry	575
happy	819
relaxed	597
sad	604
ทั้งหมด	2,595

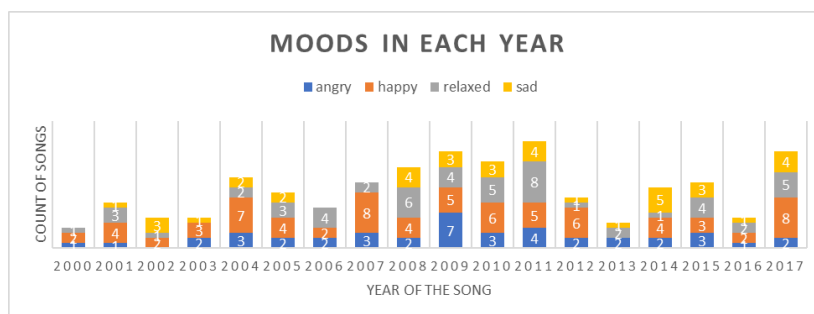
3.3.3 ข้อมูลรวม

หลังจากทำการรวมข้อมูลแล้ว ทำการแจกแจงอารมณ์ของเพลงได้ดังตารางที่ 8 เราพบว่าอารมณ์ happy มีเพลงมากที่สุดมีจำนวน 75 เพลง อารมณ์ relaxed อยู่ในอันดับที่ 2 มีจำนวน 54 เพลง ส่วนอันดับที่ 3 คืออารมณ์ angry มีจำนวน 42 เพลง และเพลงที่น้อยที่สุดคืออารมณ์ sad มีจำนวน 38 เพลง เมื่อมาเทียบกับข้อมูล Moody Corpus ความมาน้อยของแต่ละอารมณ์เพลงจะไม่เท่ากันเนื่องจากการที่ได้ทำการจับเพลงกับตัวข้อมูล Billboard

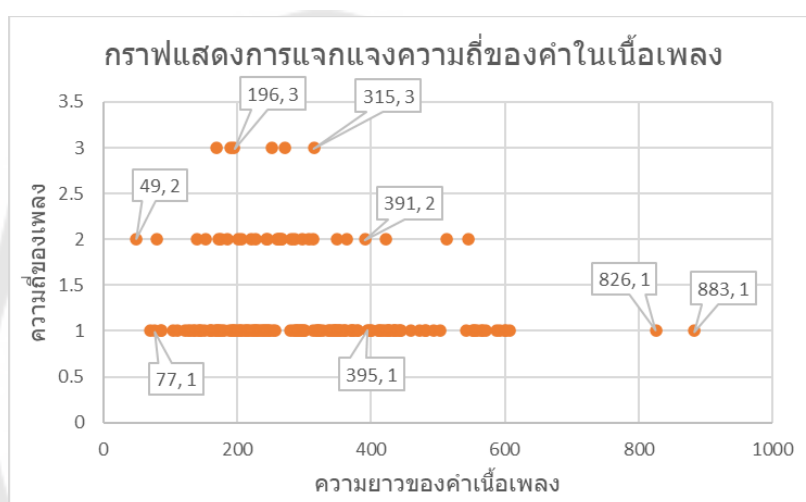
ตาราง 9 จำนวนอารมณ์ในแต่ละอารมณ์ของข้อมูลรวม

อารมณ์ของเพลง	จำนวนของเพลง
angry	42
happy	75
relaxed	54
sad	38
ทั้งหมด	209

ในปี 2011 ในข้อมูลรวมมีเพลงเยอะมากที่สุดและมีเพลงอารมณ์ relaxed มากที่สุด อารมณ์ happy มีมากที่สุด ในปี 2017 และ 2009 มีจำนวนเพลงเท่ากัน แต่ในปี 2017 อารมณ์ happy เยอะที่สุด และในปี 2009 อารมณ์ angry เยอะมากที่สุด ปีที่มีเพลงน้อยที่สุดคือปี 2000 มีอารมณ์เพลง happy มากที่สุดและไม่มีอารมณ์ sad ในปีนี้ ในปี 2014 มีอารมณ์เพลง sad มากที่สุด สรุปแล้วดูในแต่ละปีจะมีทั้ง 4 อารมณ์ประกอบกันแล้วโดยเฉลี่ยทุกปีจะเท่าๆกัน ดังภาพประกอบที่ 5



ภาพประกอบ 7 อารมณ์ในแต่ละปีข้อมูลรวม



ภาพประกอบ 8 กราฟแสดงการแจกแจงความถี่ของคำในเนื้อเพลงในข้อมูลรวม

จากภาพประกอบที่ 6 เป็นการแจกแจงความถี่ของคำในเนื้อเพลงแต่ละเพลงข้อมูลรวม พบว่าช่วงความถี่คำที่น้อยที่สุดคือ 49 คำ พบ 2 เพลง ความถี่คำที่ 77 คำ พบ 1 เพลง ความถี่ 196 คำที่ พบ 3 เพลง ความถี่ 315 คำที่ พบ 3 เพลง ความถี่ 391 คำที่ พบ 2 เพลง ความถี่ 395 คำที่ พบ 1 เพลง และเพลงความถี่ของคำเนื้อเพลงที่เยอะที่สุด 2 อันดับ 826 และ 883 พบ อย่างละ 1 เพลง จากกราฟสรุปได้ว่าส่วนมากเพลงความถี่ของคำในเนื้อเพลงถ้ายิ่งเยอะมากเพลงจะยิ่งมีน้อย

3.3.3.1 การทำความสะอาดข้อมูล

ในเนื้อเพลงนั้นเราสังเกตเห็นได้ว่ามีคำเยอะมากในเนื้อเพลง ซึ่งแต่ละเพลงนั้นก็มีความเป็นอารมณ์นั้นแฝงอยู่ และเพื่อที่จะนำมันเข้าไปทำการจัดจำแนกอารมณ์ของเพลง ดังนั้นเราจึงจัดการลบคำบางคนออกจากเนื้อเพลง ซึ่งคำพวกนั้นไม่ได้มีความหมายหรือว่าไม่ได้ส่งผล

กระทบในการนำคำมาเข้าโมเดล ในเนื้อเพลงที่นำเข้ามานั้นจะมีการเว้นวรรคซึ่งคำนั้นจะมีคำ “/n” ซึ่งแปลว่า เว้นบรรทัด พบได้ทั่วไปในการทำ Sentiment ซึ่งคำนี้ไม่มีความหมายหากจะเพียงคำสั่งให้เว้นบรรทัดเท่านั้นและไม่นำเข้ามาคำนวณ และ คำจำพวกที่เราสามารถเห็นกันได้บ่อย เช่น “Gonna”, “Wanna” และ “Gotta” เป็นต้น และคำจำพวก Stop word “a”, “and”, “the” และ “to” เป็นต้น คำทั้งหลายเหล่านี้พบได้บ่อยในเนื้อเพลงสามารถที่จะทำการลบหรือตัดออกไปได้เพราะจะไม่ทำให้เนื้อความของเพลงแต่ละเพลงไม่ผิดความหมายเนื่องจากจะคำแต่ละคำไม่ได้มีการให้ความสื่อความหมายเท่าไร

3.3.3.2 การทำ Word Cloud กับข้อมูลรวม

การทำ word cloud คือ การที่จับเอากลุ่มคำต่างๆมาเรียงจากมากที่สุดจนน้อยที่สุดมันจะสามารถแสดงให้เห็นถึงคำทั้งหมดในแต่ละเพลงหรือโดยรวมแล้วแต่จะกำหนดเพื่อให้มองเห็นความถี่ในการใช้คำของคำทั้งหมดที่เรา มี หลังจากทำ Word cloud ทั้ง 4 อารมณ์ประกอบไปด้วยอารมณ์ happy, relaxed, sad และ angry ในอารมณ์เพลง happy พบคำว่า love มากที่สุด อารมณ์เพลง angry พบคำว่า burn มากที่สุด อารมณ์เพลง relaxed พบคำว่า im มากที่สุด อารมณ์เพลง sad พบคำว่า home มากที่สุด ในการทำ Word Cloud นี้เป็นการแสดงให้เห็นว่า คำๆ ใดที่เราพบเจอบ่อยในเนื้อเพลงแต่ละอารมณ์

Word of Happy	Number of words	Word of angry	Number of words
love	1767	burn	563
im	931	im	461
know	593	na	443
dont	489	got	400
like	390	dont	348
got	370	know	342
na	367	fire	330
oh	356	let	302
life	337	like	293
yeah	317	oh	257

Word of sad	Number of words	Word of relaxed	Number of words
im	400	home	1060
lonely	217	im	860
like	177	got	447
don't	160	youre	411
got	139	take	392
girl	135	know	381
na	127	go	361
baby	114	baby	352
know	106	like	346
hold	104	dont	306

ภาพประกอบ 9 ตาราง Word Cloud แสดงคำที่ใช้ถึง 10 อันดับในแต่ละอารมณ์

3.4 Feature Selection

การทำ Feature Selection คือการลดจำนวน Feature ลงเพื่อเมื่อเวลาเอาไปสร้างแบบจำลองแล้วมันจะไม่ทำให้เครื่องเกิดการค้างหรือ Error ไม่ก็เกิดการ Overfitting เพราะจำนวนมิติข้อมูลが多เกินไป เป็นกระบวนการที่สำคัญอย่างหนึ่งในการทำการจัดจำแนกประเภทในงานวิจัยนี้เราใช้เทคนิค 2 ด้วยกัน ดังนี้

3.4.1 TF-IDF

TF-IDF หรือ Term Frequency – Inverse Document Frequency เป็นกระบวนการทางสถิติรูปแบบหนึ่ง นำมาใช้เพื่อหาความสำคัญของประโยค หรือ ข้อความ การหาค่า TF-IDF เป็นวิธีที่นิยมใช้ที่สุดในการอธิบายเนื้อเพลงให้อยู่ในรูปแบบเวกเตอร์ TF หรือ Term Frequency , มันคือการหาความถี่ของการใช้คำ นั้นหมายความว่า TF และ IDF เป็นการคำนวณความถี่ของคำทั่วไปหรือคำที่ไม่ค่อยปรากฏในเนื้อเพลงซึ่งคำนวณจากคำทั้งหมดในเนื้อเพลงหารด้วยจำนวนคำที่ปรากฏในเนื้อเพลง จากนั้นจะคำนวณค่าต่างโดยการทำลอการิทึมและการคูณด้วยความถี่ของคำที่ปรากฏในเนื้อเพลง สุดท้ายก็จะมีการระบุอัตราส่วนของจำนวนครั้งที่คำปรากฏในเนื้อเพลงนี้คือความถี่ของคำที่ปรากฏในเนื้อเพลง

สรุปคือ TF-IDF คือการหาน้ำหนักเพื่อพิจารณาว่าคำใดมีความหมายสำคัญสำหรับเนื้อเพลง ถ้ามีน้ำหนักมากคำนั้นก็มีความสำคัญในเนื้อเพลง ซึ่งเป็นกระบวนการนำความถี่ของคำที่อยู่ในเนื้อเพลง ซึ่งในเนื้อเพลงนั้นอาจมีความยาวที่แตกต่างกันไป ดังนั้นการคำนวณความถี่ของคำ จึงมักหารด้วยความยาวของเอกสาร (จำนวนคำทั้งหมดในเอกสารนั้น)

การหาค่าน้ำหนักสามารถคำนวณได้ตามสมการที่ 1 ดังนี้

$$TF - IDF = TF \times \log \frac{N}{DF} \quad (1)$$

โดย ความหมายของตัวแปรแต่ละตัว

N = จำนวนคำทั้งหมดในเนื้อเพลง

DF = จำนวนคำที่ปรากฏในเนื้อเพลงแต่ละเพลง

TF = จำนวนความถี่ของคำที่ปรากฏในเนื้อเพลงแต่ละเพลง

3.4.2 LDA (Latent Dirichlet Allocation)

LDA เป็น Latent Probabilistic Model ซึ่งเป็นงานวิจัยในปี 2003 แต่งโดย David Blei, Andrew Ng และ Michael Jordan (Blei, Ng, & Jordan, 2001) LDA ถูกใช้ในการทำ Topic Modelling โดย Topic นี้ อาจจะเป็น หัวข้อ หรือ กลุ่มคำที่ถูก ซ่อนอยู่ในเอกสารและเป็นสิ่งที่ต้องการสกัดออกมาจากเอกสาร โดย Topic เหล่านี้เป็น Latent Topic ซึ่งไม่สามารถสังเกตได้อย่างชัดเจน มันจึงจัดกลุ่มของเนื้อเพลงที่มีหัวข้อที่เกี่ยวข้องกันให้มาอยู่ในกลุ่มเดียวกัน โดยการคำนวณหาค่าความน่าจะเป็น จากคำที่ปรากฏในเนื้อเพลง ซึ่งข้อจำกัดของวิธีนี้คือ จะแสดงเฉพาะข้อมูลที่มีความเกี่ยวข้องกันมาก ๆ เท่านั้น

หลักการทำงานของ LDA คือทำการเลือกจำนวนหัวข้อขึ้นมา คือ k จำนวน หัวข้อ ที่เราต้องการ เหมือนกับการหาคลัสเตอร์ จากนั้น LDA จะทำการอ่านทุกๆ คำในเอกสารและจำกัดหัวข้อให้แต่ละคำแบบสุ่ม ซึ่งผลลัพธ์ที่ได้คือ กลุ่มของหัวข้อคำ (Topic Representation) ซึ่งในหัวข้อนั้นๆประกอบด้วยคำอะไรบ้าง โดยที่คำแต่ละคำที่อยู่ในแต่ละหัวข้อจะไปอยู่รวมกันในรูปแบบของ ถุงคำ (Bag of Words) และจะได้หัวข้อของแต่ละเอกสาร (Document represent in term of topics) แต่เนื่องจากว่ากลุ่มคำในหัวข้อที่เลือกไว้ได้มาจากการสุ่ม ดังนั้นกลุ่มคำที่อาจจะไม่มีความเกี่ยวข้องกัน นั้นทำให้ LDA จำเป็นต้องวิเคราะห์แต่ละเอกสารอีกครั้งเพื่อหาข้อมูลตามต่อไป

1. $P(\text{หัวข้อ/เอกสาร}) =$ ร้อยละของจำนวนคำศัพท์ทั้งหมดในเอกสาร ที่ถูกจัดให้อยู่ในหัวข้อ

2. $P(\text{คำศัพท์/หัวข้อ}) =$ ร้อยละของจำนวนครั้งที่คำศัพท์ ถูกจัดให้อยู่ในหัวข้อ สำหรับทุกเอกสาร

หมายเหตุ คือ $P(\text{คำศัพท์/หัวข้อ})$ คือความน่าจะเป็นแบบมีเงื่อนไข (conditionality probability)

LDA จะย้ายคำศัพท์จากหัวข้อ A ไป B เมื่อ $P(\text{หัวข้อ A|เอกสาร}) \times P(\text{คำศัพท์|หัวข้อ A})$ น้อยกว่า $P(\text{หัวข้อ B|เอกสาร}) \times P(\text{คำศัพท์|หัวข้อ B})$ หลังจากทำแบบนี้ไปสักพักหนึ่ง LDA จะได้ถุงคำศัพท์ (Bag of Words) และหัวข้อเอกสารที่ดีขึ้น เมื่อถึงจุดๆหนึ่ง ระบบก็จะหยุดการค้นหา

บทที่ 4

การทดลอง

4.1 Model Classification

งานวิจัยชิ้นนี้เราสามารถมองได้ว่าเป็นการจำแนกแบบหลายแบบ (Multi-class) การจำแนกแบบหลายแบบอธิบายถึงงานในการจำแนกข้อมูลออกเป็นแต่ละคลาส โดยมีคลาสนับจำนวนมากและแต่ละจุดข้อมูลสามารถกำหนดคลาสดูได้เพียงคลาสดูเดียว ซึ่งในการจัดจำแนกประเภทนี้เราใช้แบบจำลองจำนวน 6 แบบ ในการจัดจำแนกประเภทอารมณ์ของเพลง ในแต่ละแบบเราจะใช้วิธี Grid Search ที่ทำการหาค่าพารามิเตอร์ที่ดีที่สุดของ

4.1.1 Traditional Machine Learning

เป็นการศึกษาทางวิทยาศาสตร์เกี่ยวกับอัลกอริทึมและแบบจำลองทางสถิติที่ใช้ระบบคอมพิวเตอร์ในการทำงาน การเรียนรู้ด้วยเครื่องจักรเป็นส่วนหนึ่งของปัญญาประดิษฐ์ (Artificial Intelligence) โดยอัลกอริทึมการเรียนรู้ของเครื่องจักรจะสร้างแบบจำลองทางคณิตศาสตร์จากข้อมูลเรียนรู้ (Training Data) จากนั้นแบบจำลองที่ได้จะสามารถนำไปใช้เพื่อทำนายหรือตัดสินใจในอนาคตได้ โดยไม่ต้องเขียนโปรแกรมคำสั่งใหม่ เพื่อการทำงานเพื่อการทำงานที่รวดเร็ว

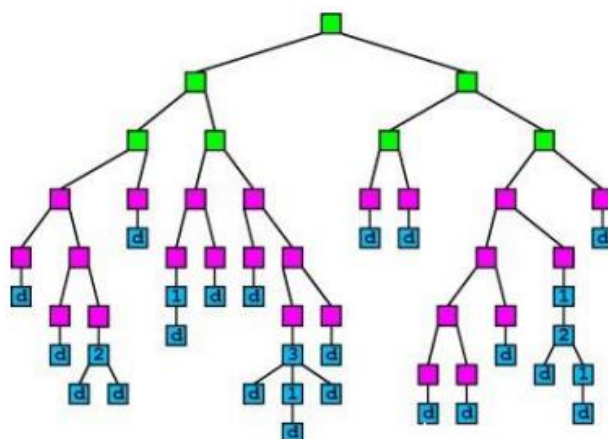
อัลกอริทึมที่รู้จักกันดีเช่น Decision tree, Random forest, Linear regression และ Naive Bayes ถูกนำมาใช้กันอย่างแพร่หลายในการจัดประเภทข้อความ ในการศึกษาครั้งนี้เราได้เปรียบเทียบประสิทธิภาพของอัลกอริทึมข้างต้น ทำการค้นหาแบบ Grid Search เพื่อค้นหาไฮเปอร์พารามิเตอร์ที่เหมาะสมที่สุด

เพื่อเพิ่มความแม่นยำในการจำแนกแบบจำลอง เราใช้ Ensemble Method โดยทั่วไปโมเดลที่เป็น Ensemble Method จะถูกมองว่า มีประสิทธิภาพสูงสุดใน Machine Learning โดยทั่วไปมีประสิทธิภาพดีกว่าแบบจำลองเดี่ยว (Subasi, 2020) วิธีการ Ensemble 2 ประเภทคือ Bagging และ Boosting ซึ่ง Random forest ตัวอย่างของแบบจำลอง Bagging นั้นเอง

4.1.1.1 Decision tree

Decision tree หรือ ต้นไม้ตัดสินใจ คือแบบจำลองต้นไม้ที่สนับสนุนการตัดสินใจ ซึ่งมีลักษณะเป็นต้นไม้กลับหัวรากจะอยู่ด้านบนสุดและใบอยู่ด้านล่างสุด ในต้นไม้จะประกอบไป

ด้วยโหนด (Node) ซึ่งคือคุณสมบัติต่างๆ ที่เป็นจุดแยกกว่าจะให้ข้อมูลไปในทิศทางใด ซึ่งโหนดที่อยู่สูงกว่าเรียกว่าโหนดราก (Root node) กิ่งของต้น (Branch) คือ คุณสมบัติของโหนดที่แตกออกมา โดยจำนวนกิ่งจะเท่ากับคุณสมบัติของโหนด ใบของต้นไม้ (Leaf) คือ กลุ่มของผลลัพธ์ในการแยกแยะจำแนกข้อมูลโดย จะแสดงส่วนประกอบของต้นไม้ ดังรูปด้านล่าง



ภาพประกอบ 10 การแยกแยะผลลัพธ์ของ Decision Tree

4.1.1.2 Naïve bayes

การจัดจำแนกแบบ Naive Bayes (เบย์เซียน) เป็นแบบจำลองหนึ่งในการจัดจำแนกที่ใช้ความน่าจะเป็น Naive bayes เช่น ถ้าต้องการแบ่งกลุ่มว่าคนไข้ที่เข้ามานั้นเป็นไข้หวัดใหญ่หรือไม่ ซึ่งเราจะต้องถามอาการคนไข้มาให้ได้มากที่สุดว่าอาการเป็นอย่างไร แล้วเราถึงจะคาดคะเนจากข้อมูลอาการที่ได้ว่ามีความน่าจะเป็น ไข้หวัดใหญ่เท่าไร เช่นความน่าจะเป็น 95% เป็นต้น ซึ่งความน่าจะเป็นดังกล่าว เราจะใช้ Naive Bayes ในการหาค่าความน่าจะเป็นนั่นเอง ซึ่งในงานวิจัยนี้เราทำการจัดหมวดหมู่เนื้อเพลงว่าเป็นแนวอารมณ์ไหนเราก็ต้องดูคำในเนื้อเพลงให้ได้มากที่สุดว่ามีคำอะไรบ้างเราถึงจะสามารถคาดการณ์จากข้อมูลที่เราได้ว่าเป็นไปได้เท่าไรที่จะเป็นอารมณ์เพลงอะไร งานวิจัยนี้ใช้การทดลองโดยคำนวณความน่าจะเป็น ดังนี้

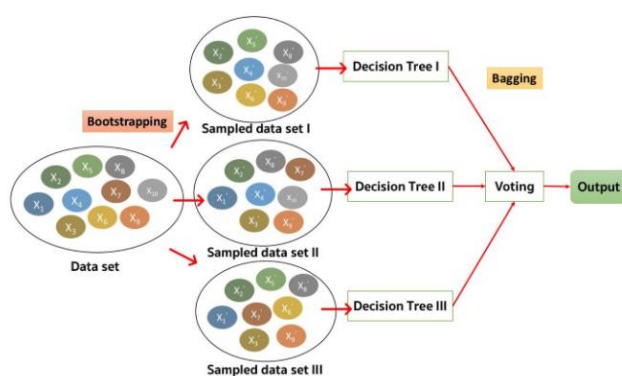
$$P(h|E1, E2, \dots, En) = \frac{P(E1|h) \times P(E2|h) \times \dots \times P(En|h) \times P(h)}{P(h|E1, E2, \dots, En)} \quad (2)$$

4.1.1.3 Logistic Regression

การถดถอยเชิงโลจิสติก หรือ Logistic regression เป็นแบบจำลองพื้นฐานที่นิยมใช้สำหรับปัญหาการจัดจำแนกประเภท นั่นคือ ปัญหาที่มีตัวแปรตาม เป็นตัวแปรประเภทไม่ต่อเนื่อง เช่น ปัญหาการจัดจำแนกประเภทลูกค้าออกเป็นลูกค้า ดี/ไม่ดี หรือ การทำนายสภาพอากาศในวันพรุ่งนี้ โดยคำทำนายที่เป็นไปได้คือ ฝนตก, แดดแรง, มีเมฆมาก เป็นต้น ในงานวิจัยของเราเราทำการจัดจำแนกอารมณ์ของเพลง นั่นคือการจัดจำแนกแบบหลายหมู่ หรือเรียกว่า Multiclass classification ซึ่งใช้ SoftMax function โดย SoftMax จะทำการทำนายอารมณ์ของเพลงที่ได้คะแนนความเป็นไปได้สูงสุด

4.1.1.4 Random Forest

Random Forest คือ แบบจำลองที่เอา Decision Tree หลายต้นมารวมกันและทำการฝึกแบบจำลองด้วยกัน สามารถทำได้มากกว่า 10 ต้น หรือ มากกว่า 1000 ต้น โดยที่แต่ละต้นจะได้รับ feature และ ข้อมูลที่เป็น subset ของ feature และ ข้อมูลทั้งหมด แบบสุ่มตอนทำการจัดจำแนกมันจะทำให้แต่ละ Decision Tree ทำการจัดจำแนกของใครของมัน และเลือกผลการจัดจำแนกครั้งสุดท้ายจากค่าที่ได้รับการโหวตมากที่สุด ซึ่งเทคนิคนี้เรียกว่า bagging หรือ bootstrapping



ภาพประกอบ 11 หลักการทำ Random Forest

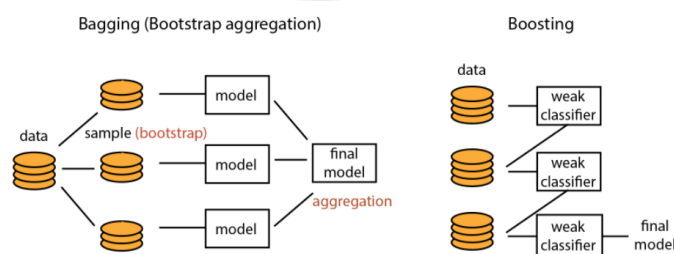
4.1.2 AdaBoost (Adaptive Boosting Algorithm)

แนวคิดเบื้องหลังอัลกอริทึมการส่งเสริมการปรับตัวหรือ AdaBoost สามารถพบได้ในวรรณกรรมหลายเล่ม (Freund & Schapire, 1999; Schapire, 2002) มันเป็น Sequential

Ensemble Method ที่มีการรวมตัวเรียนรู้ที่อ่อนแอหลาย ตัวเข้าด้วยกัน แล้วสร้างเป็นตัวเรียนรู้ที่แข็งแกร่งทำให้ผลลัพธ์ของการจัดจำแนกชั้นสุดท้ายเพิ่มขึ้น แนวคิดหลักคือการปรับน้ำหนักใหม่ซึ่งจะให้ความสำคัญกับข้อมูลที่จัดจำแนกผิดเพื่อให้มีโอกาสที่จะถูกจัดจำแนกในรอบถัดไป โดยจะมีการปรับเพิ่มน้ำหนักของข้อมูลที่ จัดจำแนกผิดและปรับลดน้ำหนักของข้อมูลที่จัดจำแนกถูกต้อง ข้อมูลเหล่านี้ไปเรียนรู้และสร้างตัวเรียนรู้ที่อ่อนแอตัวใหม่ ขั้นตอนเหล่านี้จะทำแบบ Sequential และจะมีการปรับค่าน้ำหนักในทุกๆ รอบ สรุปได้ว่า AdaBoost นั้นรวดเร็วและนำมาใช้ง่าย ใน การศึกษานี้เราใช้ AdaBoost ที่นำมาใช้ใน Scikit-learn (Pedregosa et al., 2011) โดยใช้ Decision Tree เป็น Base Learner

4.1.3 XGBoost (Extreme Gradient Boosting)

เป็นอัลกอริทึมที่ใช้โครงสร้างพื้นฐาน (Chen & Guestrin, 2016) XGBoost เป็นส่วนหนึ่งของตระกูลต้นไม้ (Decision tree, Random Forest, bagging, boosting, gradient boosting) Boosting เป็นวิธีการทั้งหมดโดยมีวัตถุประสงค์หลักในการลด Over Fit และ ความแปรปรวน เป้าหมายคือการสร้างตัวเรียนรู้ขึ้นมา และให้ตัวเรียนรู้แต่ละตัวมุ่งเน้นไปที่ความผิดพลาด (ข้อมูลที่จัดประเภทไม่ถูกต้อง) ของตัวเรียนรู้อ่อนแอ หลังจากจัดการกับความผิดพลาดแล้วน้ำหนักข้อมูลจะถูกปรับใหม่หรือที่เรียกว่า "การถ่วงน้ำหนักใหม่" ทั้งหมดสร้างแบบจำลองที่ดีขึ้นเมื่อมีการเรียนรู้ของตัวเรียนรู้อ่อนแอต่อเนื่องกันจนมีความลึกมากพอ และ แบบจำลองจะหยุดเรียนรู้เมื่อไม่เหลือผลของความผิดพลาดจากตัวเรียนรู้อ่อนแอให้เรียนรู้แล้ว ซึ่ง XGBoost ได้รับการพิสูจน์แล้วว่าประสิทธิภาพสูงสุด



ภาพประกอบ 12 รูปเปรียบเทียบการทำงานระหว่าง Bagging และ Boosting

4.2 Grid Search

Grid Search คือ การหาค่าพารามิเตอร์ที่ดีที่สุดให้แบบจำลองของเราที่เรานำมาทำการทดลอง โดยทำการกำหนดค่าพารามิเตอร์ของแต่ละแบบจำลองขึ้นมาจากนั้นทำการรันแบบจำลองในทุกชุดของค่าพารามิเตอร์ที่ให้ค่าดีที่สุด ที่เลือกหยิบแบบจำลองนี้มาใช้เนื่องจากเรามีข้อมูลน้อยและแบบจำลองของเรานั้นไม่ค่อยมีความซับซ้อน ซึ่งใช้เวลาในการรันแบบจำลองน้อย

4.3 Model Evaluate

จากการที่เราทำการใช้แบบจำลองทั้ง 6 แบบจำลองเราใช้การวัดค่าของแต่ละแบบจำลองคือค่าความถูกต้อง(accuracy) ซึ่งเป็นค่าที่บ่งบอกถึงความสามารถของเครื่องมือวัดค่าความถูกต้องที่เข้าใจค่าความเป็นจริง โดยวิธีที่เราทำการหาค่าความถูกต้อง(accuracy) ออกมาคือเราทำการฝึกแบบจำลองของเรากับค่าที่มีหลากหลายกับเมื่อเราฝึกแบบจำลองทุกแบบจำลองเสร็จแล้วเราจะนำข้อมูลตัวเดิมซึ่งมีหลากหลายกับให้แบบจำลองแต่ละแบบนั้นทำการทำนายออกมาว่าค่าจะเหมือนกันมากน้อยขนาดไหน จากนั้นเราก็จะได้ค่าความแม่นยำมาจากการที่เราเอาข้อมูลเดิมของเราที่เราใช้ฝึกก่อนหน้านี้มาทำการทดสอบกับแบบจำลองนั้นๆ

บทที่ 5

สรุปผลการทดลอง

ในบทความงานวิจัยเรื่องนี้เราได้วัดประสิทธิภาพของแบบจำลองแต่ละเทคนิค เพื่อนำมาเปรียบเทียบและสรุปผล โดยสามารถแบ่งหัวข้อในการสรุปผลดังต่อไปนี้

5.1 สรุปผลการวิจัย

5.2 อภิปรายผลการวิจัย

5.3 ข้อเสนอแนะ

5.1 สรุปผลการวิจัย

อารมณ์เพลงเป็นสิ่งที่ช่วยให้ผู้ฟังเพลงเข้าถึงอารมณ์ที่แท้จริงของตัวเองที่รู้สึกอยู่ในขณะนั้นๆ ซึ่งในปัจจุบันการหาเพลงในรายการถ่ายทอดสดหรือตามเว็บไซต์อาจจะยังไม่สามารถเข้าถึงอารมณ์เพราะยังไม่มี การแบ่งประเภทอารมณ์ของเพลงอย่างทั่วถึง นั้นจึงทำให้ผู้ฟังมีทางเลือกในการฟังเพลงได้ไม่มาก มากไปกว่านั้นในอนาคตความต้องการในการฟังเพลงตามอารมณ์นั้นอาจจะมีเพิ่มขึ้น ในงานวิจัยนี้จึงตระหนักถึงการทำนายอารมณ์เพลง และการใช้เครื่องมือในการทำนายอารมณ์เพลงจากเนื้อเพลงอย่างเหมาะสม

ผลจากงานวิจัยนี้หลังจากการประมวลผลข้อความเราพบว่ามีคำโดยเฉลี่ย 70 คำในแต่ละเพลง จาก จำนวนเพลง 209 เพลง แสดงคำที่พบบ่อยที่สุดในแต่ละอารมณ์ ในอารมณ์ Relaxed และอารมณ์ Happy คำว่า "home" และ "love" จะปรากฏมากที่สุดตามลำดับ คำที่พบบ่อยในอารมณ์ Sad เช่น "lonely" หรือ "hold" สุดท้ายคำว่า "burn" หรือ "fire" เป็นคำทั่วไปที่พบในอารมณ์ angry ต่อมาเราทำการแบ่งข้อมูลเนื่องจากเรามีข้อมูลจำนวนน้อยมากเราจึงแบ่ง 80% เป็นชุดฝึก และ 20% เป็นชุดทดสอบ ซึ่งเราได้ทำการฝึกฝนแบบจำลองตามอัลกอริทึม คือ XGBoost, Random Forest, AdaBoost, Decision tree, Naïve Bayes และ Logistic Regression ที่กล่าวมาด้านบน เราประเมินผลการคำนวณโดยใช้ค่าร้อยละของความแม่นยำ โดยหารเพลงที่จัดประเภทอย่างถูกต้องกับเพลงทั้งหมด เราพบว่า XGBoost มีประสิทธิภาพเหนือกว่าโมเดลอื่น ๆ ทั้งหมดด้วยความแม่นยำ 89% โปรดทราบว่าเราทำการค้นหา Grids Search เพื่อค้นหาพารามิเตอร์ที่ดีที่สุดสำหรับ XGBoost ซึ่งมีดังต่อไปนี้ คือ อัตราการเรียนรู้ที่ดีที่สุดคือ 0.1 , ตัวประมาณค่า คือ 80 และ ความลึกสูงสุด คือ 8

Classifier	Accuracy
XGBoost	89.19%
Random forest	81.01%
AdaBoost	63.68%
Decision tree	60.13%
Naïve Bayes	58.02%
Logistic regression	48.41%

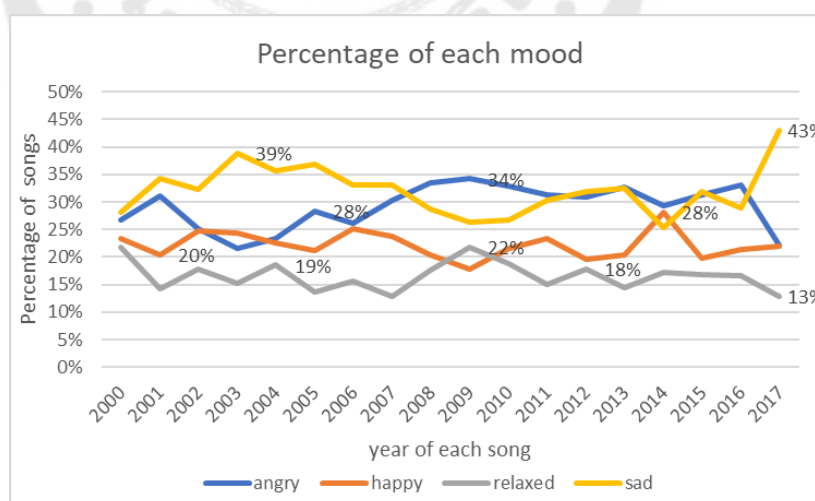
ภาพประกอบ 6 ผลลัพธ์จากการทดสอบแบบจำลอง

5.2 อภิปรายผลการวิจัย

ในงานวิจัยนี้ ได้ตั้งแบบจำลองการจัดจำแนกอารมณ์ของเพลง เพื่อทำนายอารมณ์ของเพลงจากเนื้อเพลง โดยอาศัยการเรียนรู้ของเครื่อง โดยทำการจำแนกประเภท หรือเรียกว่าการทำ Classification ซึ่งเราได้นำข้อมูล Billboard โดยนำแบบจำลอง XGBoost มาใช้ในการจัดจำแนกอารมณ์ของเพลงจากเนื้อเพลง เพราะเนื่องจากมีความแม่นยำมากที่สุดแล้วนำมาทำนายชุดข้อมูล Billboard 2000-2017 โดยก่อนที่เราจะเอาข้อมูล Billboard มาทำการทำนาย เราจะทำการลบเพลงจากข้อมูลรวมที่มีข้อมูลอยู่ในข้อมูล Billboard ออกก่อน มีข้อมูลเท่ากับ 5,533 เพลง และจำแนกประเภทอารมณ์ของเนื้อเพลงดังภาพประกอบที่ 11 พบว่าปี 2009 อารมณ์ angry พบมากที่สุดจำนวน 138 เพลง และ น้อยที่สุดคือ 53 ในปี 2004 พบว่าปี 2017 อารมณ์ happy พบมากที่สุดจำนวน 112 เพลง และ น้อยที่สุดคือ 40 ในปี 2001 พบว่าปี 2009 อารมณ์ relaxed พบมากที่สุดจำนวน 88 เพลง และ น้อยที่สุดคือ 28 ในปี 2001 และ พบว่าปี 2017 อารมณ์ sad พบมากที่สุดจำนวน 2017 เพลง และ น้อยที่สุดคือ 66 ในปี 2000 จากการทดลองกับข้อมูล Billboard นั้น เนื่องจากเรามีข้อมูลจากการที่เรานำมาฝึกฝนแค่เพียง 206 แถว จากการทดลองนี้ เนื่องจาก เรามีข้อมูลในการฝึกฝนกับแบบจำลอง ในเนื้อเพลง อาจมีความหลากหลายของคำในเนื้อเพลงไม่เพียงพอในการฝึกของแบบจำลอง เมื่อเรานำแบบจำลอง ที่เราเคยฝึกมาทำการจัดจำแนกอารมณ์ของเนื้อเพลงโดยใช้ ข้อมูลเนื้อเพลง จาก Billboard มาทำการจัดจำแนกอารมณ์ของเพลงโดยใช้เนื้อเพลง ซึ่งแบบจำลองของเรา อาจจะไม่เคยเจอคำบางคำในเนื้อเพลงจากข้อมูลเนื้อเพลงใน Billboard ที่มีจำนวน 5,533 เพลง การจัดจำแนกออกมาอาจจะยังได้ผลที่ยังไม่ถูกต้องมากนัก เนื่องจากมีข้อมูลที่ไม่เคยเห็นการฝึกแบบจำลองมาก่อน จะได้ผลตามภาพประกอบ 13

Year	angry	happy	relaxed	sad
2000	63	55	51	66
2001	61	40	28	67
2002	61	60	43	78
2003	55	62	39	99
2004	53	51	42	81
2005	75	56	36	97
2006	77	74	46	97
2007	84	66	36	92
2008	97	59	51	83
2009	138	72	88	106
2010	101	66	58	82
2011	118	88	57	114
2012	93	59	54	96
2013	109	68	48	108
2014	89	85	52	77
2015	104	66	56	106
2016	127	82	64	111
2017	113	112	66	219
Total	1618	1221	915	1779

ภาพประกอบ 13 ผลการทำนายอารมณ์ของเพลงแต่ละเพลงแยกตามจำนวน
อารมณ์ในแต่ละปี



ภาพประกอบ 14 กราฟแสดงร้อยละของอารมณ์เพลงในแต่ละปี

จากกราฟเป็นการคิดร้อยละ (percent) ในแต่ละอารมณ์จากรูปประกอบที่ 14 จากการใช้แบบจำลอง XGBOOST ที่เราฝึกมานั้นพบว่าจากปี 2000-2017 แนวโน้มของแต่ละอารมณ์มีลักษณะดังนี้

1.อารมณ์ happy และ อารมณ์ sad มีการเปลี่ยนแปลงไปในทิศทางเดียวกัน ยกเว้น ในปี 2014 – 2017 ซึ่งในปี 2017 นั้นพบว่าเพลงอารมณ์ sad มีสัดส่วนเพิ่มขึ้นอย่างเห็นได้ชัด

2.อารมณ์ angry และ อารมณ์ relaxed มีการเปลี่ยนแปลงไปในทิศทางเดียวกัน ยกเว้นในปี 2000 – 2017 อารมณ์ sad จะมีสัดส่วนสูงที่สุด และในปี 2008 – 2017 อารมณ์ angry จะมีสัดส่วนสูงที่สุด

เนื่องจากเรามีอารมณ์เพลง 4 ประเภท คือ happy, sad, relaxed และ angry ต่อมา เนื่องจากเราอยากทำการปรับมุมมองการมองภาพอารมณ์เพลงเป็น 2 แบบ คือ อารมณ์ที่มีความสุขและอารมณ์ที่ไม่มีความสุข โดยอารมณ์ที่มีความสุขเกิดจากการเอาอารมณ์ happy รวมกับ relaxed และ อารมณ์ที่ไม่มีความสุขคิดจาก การเอาอารมณ์ sad รวมกับ angry ซึ่งเราได้ผลตามตารางด้านล่าง

ตาราง 10 แจกแจงเพลงมีความสุขและเพลงไม่มีความสุข

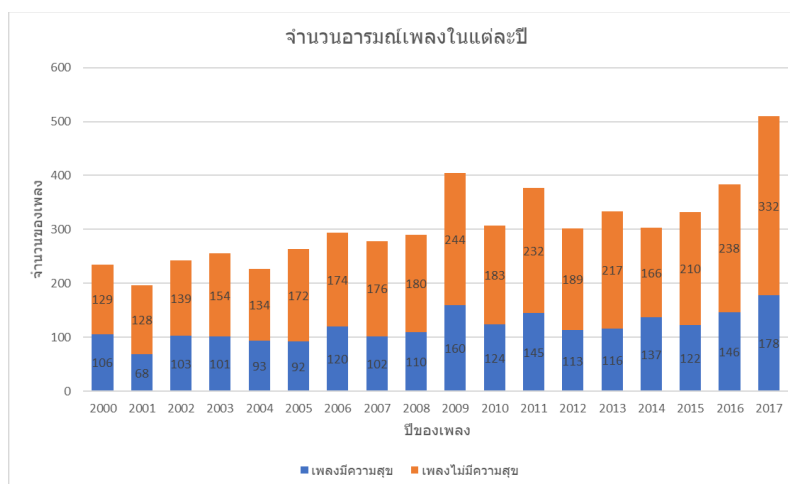
ปี	รวม	เพลงมีความสุข	เพลงไม่มีความสุข	เพลงมีความสุข (%)	เพลงไม่มีความสุข (%)
2000	235	106	129	45.11%	54.89%
2001	196	68	128	34.69%	65.31%
2002	242	103	139	42.56%	57.44%

ตาราง 10 (ต่อ)

ปี	รวม	เพลงมีความสุข	เพลงไม่มี ความสุข	เพลงมีความสุข (%)	เพลงไม่มีความสุข (%)
2003	255	101	154	39.61%	60.39%
2004	227	93	134	40.97%	59.03%
2005	264	92	172	34.85%	65.15%
2006	294	120	174	40.82%	59.18%
2007	278	102	176	36.69%	63.31%
2008	290	110	180	37.93%	62.07%
2009	404	160	244	39.60%	60.40%
2010	307	124	183	40.39%	59.61%
2011	377	145	232	38.46%	61.54%

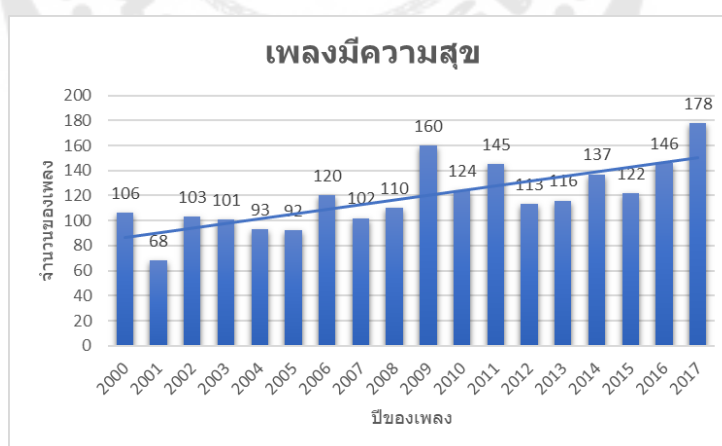
ตาราง 10 (ต่อ)

ปี	รวม	เพลงมีความสุข	เพลงไม่มี ความสุข	เพลงมีความสุข (%)	เพลงไม่มีความสุข (%)
2012	302	113	189	37.42%	62.58%
2013	333	116	217	34.83%	65.17%
2014	303	137	166	45.21%	54.79%
2015	332	122	210	36.75%	63.25%
2016	384	146	238	38.02%	61.98%
2017	510	178	332	34.90%	65.10%
ทั้งหมด	5533	2136	3397	38.82%	61.18%



ภาพประกอบ 15 จำนวนเพลงมีความสุขและไม่มีความสุขในแต่ละปี

ในปี 2017 ปรากฏเพลงอารมณ์ที่ไม่มีมีความสุขขึ้นมาเยอะมากที่สุด 332 เพลงคิดเป็น 65.01% ของอารมณ์เพลงทั้งหมดในปี 2017 และเพลงที่มีอารมณ์ความสุขเยอะมากที่สุด 178 เพลง ในปี 2017 คิดเป็น 34.90% สังเกตได้ว่าเพลงไม่มีความสุขจะมีเยอะกว่าเพลงที่มีมีความสุขในทุกๆปีตามภาพประกอบที่ 15



ภาพประกอบ 16 กราฟแท่งเพลงมีความสุข



ภาพประกอบ 17 กราฟแท่งเพลงไม่มีความสุข

5.3 ข้อเสนอแนะ

5.3.1 สรุปในการทำงานวิจัยครั้งนี้เราได้ทำการเรียนรู้แบบจำลองการจัดจำแนกอารมณ์ของเพลง ซึ่งข้อมูลที่เรานำมาจัดจำแนกนั้น มีจำนวนน้อย ในอนาคตจะต้องมีข้อมูล สำหรับคนการฝึกแบบจำลองที่มากขึ้นกว่าเดิม เพื่อที่จะได้ค่าความถูกต้องของแบบจำลองที่ดีขึ้น สามารถที่จะนำไปพัฒนาต่อในเชิงธุรกิจ เช่น รายการถ่ายทอดสดเพลง หรือ ตามเว็บไซต์เนื้อเพลงต่างๆ จัดจำแนกอารมณ์ของเพลงจากเนื้อเพลงให้มีความแม่นยำมากขึ้น

5.3.2 เนื่องจากข้อมูลที่เราได้นำมาทำการทดสอบการจัดจำแนกอารมณ์เพลงจากเนื้อเพลง โดยใช้ข้อมูลเนื้อเพลงจาก Billboard ในการฝึกฝนแบบจำลองที่เราใช้ข้อมูลจำนวนน้อยมาก ซึ่งค่าบางค่าข้อมูลจาก Billboard อาจมีค่าแปลกๆ ที่แบบจำลองของเรา ไม่เคยเห็นมาก่อนอาจจะทำให้ ความถูกต้อง ในการจัดจำแนก ไม่ดีพอในอนาคตเราจึงมุ่งมั่นที่จะศึกษาปรับปรุงพัฒนาแก้ไขแบบจำลองต่อไป โดย เราจะใช้ Unknown Token เข้ามาร่วมในการฝึกแบบจำลองเพื่อเอามาจับกับคำที่ไม่เคยฝึกมาก่อนจากแบบจำลองก่อนหน้านี้ ซึ่งถ้ามีเนื้อเพลงใหม่เข้ามา แบบจำลองจะทำงานได้ไวขึ้นความถูกต้องของแบบจำลองก็จะดีขึ้น

5.3.3 สร้างแบบจำลองด้วยอัลกอริทึม อื่นๆ ทดสอบดูความถูกต้องเพื่อที่จะได้แบบจำลอง ที่มีค่าความถูกต้องสูงที่สุด

5.3.4 เปลี่ยนการปรับปรับค่าพารามิเตอร์จากการใช้วิธี Grid Search ที่ทำการหาค่าพารามิเตอร์ที่ดีที่สุด ปรับโดยใช้การปรับค่าเองด้วยมือ

5.3.5 ในอนาคต อาจจะมีการ นำข้อมูล เพื่อต่อยอดเรื่องธุรกิจ คือเรื่องการเพิ่มยอดขาย หรือทำกำไรให้กับธุรกิจวงการเพลงโดยใช้ข้อมูลความสุขของคนทั่วโลกในการทำนาย GDP ในอนาคตร่วมกับแนวโน้มในการทำเพลงออกมาสู่ตลาดเพลงในอนาคต



บรรณานุกรม

- Akella, R., & Moh, T. (2019). Mood Classification with Lyrics and ConvNets. *2019 18th IEEE International Conference On Machine Learning And Applications (ICMLA)*, 511-514.
- Al-Salemi, B., Mohd Noah, S. A., & Ab Aziz, M. J. (2016). RFBoost: An improved multi-label boosting algorithm and its application to text categorisation. *Knowledge-Based Systems*, 103, 104-117.
- An, Y., Sun, S., & Wang, S. (2017, 24-26 May 2017). *Naive Bayes classifiers for music emotion classification based on lyrics*. Paper presented at the 2017 IEEE/ACIS 16th International Conference on Computer and Information Science (ICIS).
- Andjelkovic, I., Parra, D., & O'Donovan, J. (2019). Moodplay: Interactive music recommendation based on Artists' mood similarity. *International Journal of Human-Computer Studies*, 121, 142-159. Retrived from <http://www.sciencedirect.com/science/article/pii/S1071581918301654>
- Blei, D., Ng, A., & Jordan, M. (2001). *Latent Dirichlet Allocation* (Vol. 3).
- Çano, E., & Morisio, M. (2017). *Moodylyrics: A sentiment annotated lyrics dataset*. Paper presented at the Proceedings of the 2017 International Conference on Intelligent Systems, Metaheuristics & Swarm Intelligence.
- Chauhan, S., & Chauhan, P. (2016). *Music mood classification based on lyrical analysis of Hindi songs using Latent Dirichlet Allocation*. Paper presented at the 2016 International Conference on Information Technology (InCITe)-The Next Generation IT Summit on the Theme-Internet of Things: Connect your Worlds.
- Chen, T., & Guestrin, C. (2016). *XGBoost: A Scalable Tree Boosting System*. Paper presented at the Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, San Francisco, California, USA. <https://doi.org/10.1145/2939672.2939785>
- Elmitwally, N. S. (2020, 13-15 Oct. 2020). *Building A Multi-class XGBoost Model for Arabic Figurative Language*. Paper presented at the 2020 2nd International Conference

- on Computer and Information Sciences (ICCIS).
- Freund, Y., & Schapire, R. (1999). *A Short Introduction to Boosting*.
- Hu, Y., & Ogiwara, M. (2012). *Identifying accuracy of social tags by using clustering representations of song lyrics*. Paper presented at the 2012 11th International Conference on Machine Learning and Applications.
- Matsumoto, K., & Sasayama, M. (2018). *Lyric Emotion Estimation Using Word Embedding Learned from Lyric Corpus*. Paper presented at the 2018 IEEE 4th International Conference on Computer and Communications (ICCC).
- Oh, S., Hahn, M., & Kim, J. (2013). *Music mood classification using intro and refrain parts of lyrics*. Paper presented at the 2013 International Conference on Information Science and Applications (ICISA).
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., . . . Duchesnay, É. (2011, avec comité de lecture). Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research*.
- Rachman, F., Sarno, R., & Fatichah, C. (2018). Music Emotion Classification based on Lyrics-Audio using Corpus based Emotion. *International Journal of Electrical and Computer Engineering (IJECE)*, 8, 1720.
- Rachman, F. H., Sarno, R., & Fatichah, C. (2020). *Music Emotion Detection using Weighted of Audio and Lyric Features*. Paper presented at the 2020 6th Information Technology International Seminar (ITIS).
- Russell, J. A. (1980). A circumplex model of affect. *Journal of Personality and Social Psychology*, 39(6), 1161-1178.
- Salminen, J., Hopf, M., Chowdhury, S. A., Jung, S.-g., Almerikhi, H., & Jansen, B. J. (2020). Developing an online hate classifier for multiple social media platforms. *Human-centric Computing and Information Sciences*, 10(1), 1.
- Schapire, R. (2002). The Boosting Approach to Machine Learning: An Overview. *Nonlin. Estim. Classif. Lect. Notes Stat*, 171.
- Subasi, A. (2020). *Practical Machine Learning for Data Analysis Using Python*: Elsevier Science.

Tauberg, M. (2018). Billboard Hot-100 Songs 2000-2018 w/ Spotify Data + Lyrics.

[typhon/billboard-hot-100-songs-2000-2018-w-spotify-data-lyrics](#)

Yang, D., & Lee, W.-S. (2009). *Music emotion identification from lyrics*. Paper presented at the 2009 11th IEEE International Symposium on Multimedia.

Ying, T. C., Doraisamy, S., & Abdullah, L. N. (2012). *Genre and mood classification using lyric features*. Paper presented at the 2012 International Conference on Information Retrieval & Knowledge Management.

ปักเข็ม, ผ. ด. น. เอกสารประกอบการสอนวิชาการทำเหมืองข้อมูล เวอร์ชัน 10.



ประวัติผู้เขียน

ชื่อ-สกุล นางสาวกรกัญญา ศิริเกษ
วัน เดือน ปี เกิด 9 ธันวาคม 2539
สถานที่เกิด กรุงเทพมหานคร
วุฒิการศึกษา พ.ศ.2557
วิทยาศาสตรบัณฑิต สาขาวิทยาการประกันภัย
จาก มหาวิทยาลัยธรรมศาสตร์
พ.ศ.2562
วิทยาศาสตรมหาบัณฑิต สาขาวิทยาการข้อมูล
จาก มหาวิทยาลัยศรีนครินทรวิโรฒ

