



วารสารวิทยาศาสตร์ประยุกต์

THE JOURNAL OF APPLIED SCIENCE

ISSN 1513-7805 [Print] ISSN 2586-9663 [Online]

ปีที่ 21 ฉบับที่ 1 [2565] Vol.21, No.1 [2022]

THE JOURNAL IS COVERED BY:



Approved by TCI during
2020-2024



TCI

ศูนย์ดัชนีการอ้างอิงวารสารไทย
Thai-Journal Citation Index Centre



ASEAN
CITATION
INDEX

Available online at <https://www.tci-thaijo.org/index.php/JASCI>

Consulting Editors

from King Mongkut's University of Technology North Bangkok, Thailand

Assoc.Prof.Dr. Surapun Yimman

Prof.Dr. Saowanit Sukparungsee

Assoc.Prof.Dr. Kitti Bodhipadma

Editor

from King Mongkut's University of Technology North Bangkok, Thailand

Assoc.Prof. Narumol Kreua-ngarjnkool

Co-Editor

from King Mongkut's University of Technology North Bangkok, Thailand

Asst.Prof.Dr. Peerapong Pornwongthong

Assoc.Prof.Dr. Nonchanutt Chudpooti

Associate Editor

from King Mongkut's University of Technology North Bangkok, Thailand

Prof.Dr. Yupaporn Areepong

Assoc.Prof.Dr. Nutsuda Sumonsiri

Assoc.Prof.Dr. Sekson Sirisubtawee

Assoc.Prof.Dr. Pianpool Kamoljitprapa

Asst.Prof.Dr. Anusara Srisrual

Asst.Prof.Dr. Walaiporn Prissanaroon Ouajai

Asst.Prof.Dr. Tanapat Anusas-amornkul

Asst.Prof.Dr. Suvimol Phanyaem

Asst.Prof.Dr. Theerawut Phusantisampan

Asst.Prof.Dr. Jintawat Tanamatayarat

Asst.Prof.Dr. Sukanya Thepwatee

Editorial Advisory Board

Prof.Dr. Paul Pigram

La Trobe University, Australia

Prof.Dr. Melvin Pascall

The Ohio State University, USA

Prof.Dr. Sergey Meleshko

Suranaree University of Technology, Thailand

Emeritus Prof.Dr. Yongwimon Lenbury

Mahidol University, Thailand

Prof.Dr. Sutthisak Phongthanapanich

King Mongkut's University of Technology North Bangkok, Thailand

Prof.Dr. Anuvat Sirivat

Chulalongkorn University, Thailand

Assoc.Prof.Dr. Yaowadee Temtanapat

Thammasat University, Thailand

Assoc.Prof.Dr. Khongsak Srikaeo

Pibulsongkram Rajabhat University, Thailand

Assoc.Prof.Dr. Weeradej Meeinkuirt

Mahidol University, Thailand

Editorial Advisory Board (cont.)

Assoc.Prof.Dr. Wararit Panichkitkosolkul
Thammasat University, Thailand

Assoc.Prof.Dr. Montip Tiensuwan
Mahidol University, Thailand

Assoc.Prof.Dr. Piyapatr Busababodhin
Mahasarakham University, Thailand

Assoc.Prof.Dr. Wantida Chaiana
Chiang Mai University, Thailand

Asst.Prof.Dr. Kwan Arayathanitkul
Mahidol University, Thailand

Asst.Prof.Dr. Narumon Emarat
Mahidol University, Thailand

Asst.Prof.Dr. Kitiporn Plaimas
Chulalongkorn University, Thailand

Asst.Prof.Dr. Chatchai Khunboa
Khon Kaen University, Thailand

Asst.Prof.Dr. Julaluk Khemacheewakul
Chiang Mai University, Thailand

Asst.Prof.Dr. Jutarat Iewkittayakorn
Prince of Songkla University, Thailand

Asst.Prof.Dr. Yanisa Laoong-u-thai
Burapha University, Thailand

Asst.Prof.Dr. Sararat Mahasaranon
Naresuan University, Thailand

Dr. Worawut Srisukkhham
Chiang Mai University, Thailand

Coordinator and Manager

from King Mongkut's University of Technology North Bangkok, Thailand

Mr. Phirom Pabu

Mr. Piyawat Sutha

Miss.Wilaiporn Khunkaew

Miss.Wilawan Sae-jeng

Miss.Kitsiya Chuchuaysuwan

Contact info

Faculty of Applied Science

King Mongkut's University of Technology North Bangkok

1518 Pracharat 1 Road, Wongsawang, Bangsue, Bangkok 10800 : Thailand

Phone: 02-555-2000#4211

Website

www.tci-thaijo.org/index.php/JASCI

JOURNAL POLICY

Focus and Scope

The Journal of Applied Science is an academic journal published biannually by the Faculty of Applied Science, King Mongkut's University of Technology North Bangkok. The Journal publishes original research and review papers either in Thai or English covering all areas of applied science and technology. The journal will not accept articles, which have been published or are being considered for publication by another journal, nor should papers published here be submitted to other journals.

"The Journal of Applied Science does not have policy to collect publication fee"

"Each articles must be evaluated by two peers (double blinded) before accepted for publication"

"Article must be revised and sent back to the journal within 4 weeks after the return for revision unless the article will be rejected"

"The Journal of Applied Science published both as hard -copies [ISSN 1513-7805 (Print)] and electronic journal [ISSN 2586-9663 (Online)] available on ThaiJO system"

Scope of the Journal

To publish academic, research and review articles covering all area of both basic and applied sciences and technology including pure and applied mathematics, statistics, chemistry and applied chemistry, physics and industrial physics, environmental science and technology, biotechnology, agro-industrial and food technology, medical science and applications, health and beauty technology, computer and informatic sciences and materials science

Peer Review Process

Each article must be double blind peer reviewed by at least 2 reviewers from the related field.

Human and animal ethics in research

The authors are requested for the appropriate disclosures and declarations if their research involves human participants and animals having an impact on their right, prestige, safety, and health. Hence, they must receive approval from the appropriate ethics committee for research involving humans and/or animals and submit the evidence of approval with the manuscript. Starting from May, 1st 2022)

Language

Both Thai and English

Publication Frequency

Twice a year (Biannual)

First Issue: January to June

Second Issue: July to December

Courtesy by the Faculty of Applied Science,

King Mongkut's University of Technology North Bangkok

Remarks

Authors have to be responsible for any legal effects that may occur due to their opinions expressed in the articles.

TABLE OF CONTENT

Modelling the Household Debts in Thailand using Fay-Herriot models

Annop Angkunsit and Jiraphan Suntornchost.....244482

Significant Attributes from Multiple Regression and Marginal Effects on Multinomial Logit Model for Chemical Servitization

Tanyaporn Kanignat, Pongsa Pornchaiwiseskul and
Kamonchanok Suthiwartnarueput.....244434

Zero Inflated and Zero Truncated Negative Binomial Weighted Garima Distributions for Modeling Count Data

Pornpop Saengthong.....244341

Finite Dimensional Simple Poisson Modules of a Poisson Algebras A

Nagornchat Chansuriya and Supaporn Fongchanta.....244580

An approximation of average run length using numerical integration methods on EWMA control chart for MAX (q, r) process

Piyaphon Paichit and Kanita Petcharat.....247274

Non – Stationary Booking Demand in Two-Class Overbooking Model

Piyachat Leelasilapasart and Murati Somboon.....247376

Comparison of classification model for steam trap valve opening sound

Punyanut Damnong, Phimpaka Taninpong, Anucha Promwungkwa and
Jakramate Bootkrajang.....244672

Wastewater Treatments for Textile Industry by Electrocoagulation Process

Thipsuree Kornboonraksa and Tidarad Chompurach.....244292

Nutrient status in rubber growing soil, rubber leaf and nutrient concentration in latex under different rubber-based intercropping plantation

Pranchalee Saeteng, Jumpen Onthong and Khwunta Khawmee.....244327

Development of cracker products from purple yam flour (*Dioscorea alata* L.)

Ruethaipak Chansri, Rassarin Chatthongpisut and Ratchawan Jantakhat.....244491

Study on properties and cracked heels treatment of insole foam made from urea filled natural rubber latex

Hasan Daupor, Sitisaiyidah Saiwari, Sukree Semmard and Ajaman Adair.....244287

Artificial Intelligent Techniques for Thai Fake News Detection

Phayung Meesad, Phnom Kleechaya, Aongart Aun-a-nan and
Kamonrat Kijrungpaisarn.....244590

TABLE OF CONTENT (cont.)

Rice canopy height model assessment using unmanned aerial vehicle (UAV) images

Parinchat Juntaworn and Kritchayan Intarat.....244575

Factors affecting the picking travel distance in the leaf warehouse

Makusee Masae, Witchuda Boonreung, Phakthiya Hmadhloo and
Panupong Vichitkunakorn.....244190

A comparing on performance of t-test and Wilcoxon rank sum test for two independent populations using counting data

Khemika Urawong, Jularat Chumnaul, Chiranan Phaochamrun and
Asalaya Singhabumrung.....245800

A Comparison of Time Series Forecasting Models between Hybrid Model and Individual Model for Forecasting Daily Incoming Call Volume

Pornpimol Chaiwuttisak.....244662

Forecasting Model for the Export Values for Sugar of Thailand

Warangkhan Riansut.....244424

Average Run Length of CUSUM Control Chart for SARX (P,1) L Model using Numerical Integral Equation Method

Phornpimon Buranapol and Suvimol Phanyaem.....246470

Forecasting model for Export Condom Quantity of Thailand on the COVID-19 situation

Pradthana Minsan.....244500

Forecast of the price of jasmine rice with Data Mining Techniques

Chalermwut Comemuang and Wachirarak Orosram.....244210

Research Article

Modelling the Household Debts in Thailand using Fay-Herriot models

Annop Angkunsit and Jiraphan Suntornchost*

Department of Mathematics and Computer Science, Faculty of Science, Chulalongkorn University, Bangkok 10330, Thailand

*E-mail: Jiraphan.S@chula.ac.th

Received: 22/05/2021; Revised: 17/10/2021; Accepted: 28/03/2022

Abstract

In this study, we investigate the performance of the Fay-Herriot area level model to estimate the Household Debts of Thailand by using the 2017 household socio-economic survey from the National Statistics Office (NSO) of Thailand. Two versions of the model are considered which are the classical Fay-Herriot model and the Bayesian approach to the model, called as the Hierarchical Bayesian Fay-Herriot model. Results from our study show that the two models can reduce mean square errors of the direct estimates, particularly for the areas with large sampling variances. Moreover, we apply results from the two models to discuss behavior of household debts in Thailand by comparing household debts among different regions of Thailand. Moreover, we study the difference in household debts between municipal area and non-municipal area. From our analysis, household debts in municipal areas are higher than those of non-municipal areas. Comparing among different regions, we find that Bangkok has much higher average household debt than other regions. The northern and the northern east regions have higher household debts than the country average while the central, the east, and the southern regions have average household debts lower than the country average.

Keywords: Household debt, Fay-Herriot model, empirical best linear unbiased predictor, Bayesian.

Introduction

Due to current economic crisis, the rise of household debt is currently considered as one of the most important economic issues in many countries including Thailand. The problem of household debts affects the whole economy and living standards of people in the country. Therefore, several studies of debts in Thailand have been conducted in many different aspects. For example, Kiatipong et al. (2007) studied the wealth and debt of Thai household. Muthitacharoen et al. (2015) discussed the effect of the rise of household debts to the economic stability of Thailand. Visitwarakorn (2015) studied the debt of Thai's government due to the change of 2013 salary policy. Moreover, the Thailand's National Statistical Office also conducts regular surveys of the household debts of the Thai population and reports an estimate of annual household debt computed by using the information of household which is collected from all provinces in Thailand. Based on our current knowledge, these studies are performed based on the direct estimates from a survey. Therefore, the estimates are heavily influenced by

the sampling design and the sampling errors particularly for areas with small sample sizes. To overcome the issues, alternative model-based methods have been proposed in the context of "small area statistics" (Rao & Molina, 2015) which are the statistical methods used when direct estimates from a survey are not of adequate precision, or in the situation of small sample sizes. The motivation of the method is to use statistical models to link the variable of interest with auxiliary information to define the model-based estimator that "borrow strength" from the related area. One well-known small area model applied in different applications is the Fay-Herriot model. The Fay-Herriot model was first introduced by Fay and Herriot in 1979 to predict the mean per capita income in small places within counties in USA. Since then, the model has been widely applied in different countries. For example, Srivastava et al. (2007) used the Fay-Herriot model to obtain the model-based district level estimates of the amount of loan outstanding per household in India. D'Alò et al. (2008) applied the Fay Herriot model to study socio-economic indicators in Italy. Wawrowski (2016) applied the spatial Fay-Herriot model to study poverty in Poland. For Thai data, Angkunsit & Suntornchost (2021) applied the bivariate Fay-Herriot model to obtain the model-based estimates of the average household income and average household expenditure for Thai data. In this paper, we study the performance of Fay-Herriot area level model to estimate the Household Debts of Thailand. Two versions of the Fay-Herriot model are considered which are the classical Fay-Herriot model and the hierarchical Bayesian Fay-Herriot model. The data used in our study are the Household Socio-Economic Survey 2017 and Population and Housing Census 2010 collected by the Thailand's National Statistical Office.

The organization of this paper is as follows. In Section 2, we introduce the Fay-Herriot model and the hierarchical Bayesian Fay-Herriot model. In this section, we also discuss estimation methods and model diagnostics of the two models. In Section 3, we describe the data source and auxiliary variables used in this study, model comparison, and discuss applications of the two models to the household debt of Thailand. Section 4 gives discussions and conclusions of our study.

Small Area Models

In this section, we first describe the small area models used in this paper which are the Fay-Herriot and the hierarchical Bayesian Fay-Herriot model.

2.1 Fay-Herriot model

The Fay-Herriot model was proposed by Fay & Herriot (1979) to predict the mean per capita income in small places within counties in USA. The structure of the Fay-Herriot model is as follows. Suppose that the population is partitioned into m subpopulations, called as "area". For any area i ($i=1, \dots, m$), let θ_i be the true area mean of i th area and y_i be a direct estimator of θ_i . The model assumes that θ_i is linearly related to the auxiliary variables $\mathbf{x}'_i = (1, x_{i1}, \dots, x_{ip})'$, for $i=1, \dots, m$, through the model

$$\theta_i = \mathbf{x}'_i \boldsymbol{\beta} + v_i, \quad v_i \sim N(0, A), \quad i=1, \dots, m,$$

where $\boldsymbol{\beta} = (\beta_0, \beta_1, \dots, \beta_p)$ is the vector of regression coefficients and A is the regression variance. The true mean is linked to direct estimator y_i through the following sampling model

$$y_i = \theta_i + e_i, \quad e_i \sim N(0, D_i), \quad i = 1, \dots, m,$$

where e_i is the sampling error with the known sampling variance D_i .

The Fay-Herriot model can be rewritten as

$$y_i = \mathbf{x}_i' \boldsymbol{\beta} + v_i + e_i, \quad i = 1, \dots, m, \quad (1)$$

where the area random effects v_i are independent of the sampling errors e_i .

If the area random effect variance A is known, the true area mean is estimated by minimizing the mean squared error (MSE) of θ_i proposed in Henderson (1975). The obtained estimate is called as the best linear unbiased prediction, or BLUP, defined as follows.

$$\tilde{\theta}_i = \frac{A}{A + D_i} y_i + \frac{D_i}{A + D_i} \mathbf{x}_i' \tilde{\boldsymbol{\beta}}, \quad (2)$$

where $\tilde{\boldsymbol{\beta}} = \left(\sum_{i=1}^m \mathbf{x}_i \mathbf{x}_i' (A + D_i)^{-1} \right)^{-1} \left(\sum_{i=1}^m \mathbf{x}_i y_i (A + D_i)^{-1} \right)$. However, A is unknown in general practice. Therefore, it is estimated by an estimate $\hat{A}$. Therefore, by substituting $\hat{A}$ into (2), we obtain the empirical BLUP or EBLUP $\hat{\theta}_i$ of θ_i . That is,

$$\hat{\theta}_i = \frac{\hat{A}}{\hat{A} + D_i} y_i + \frac{D_i}{\hat{A} + D_i} \mathbf{x}_i' \hat{\boldsymbol{\beta}}, \quad (3)$$

where $\hat{\boldsymbol{\beta}} = \tilde{\boldsymbol{\beta}}(\hat{A}) = \left(\sum_{i=1}^m \mathbf{x}_i \mathbf{x}_i' (\hat{A} + D_i)^{-1} \right)^{-1} \left(\sum_{i=1}^m \mathbf{x}_i y_i (\hat{A} + D_i)^{-1} \right)$.

There are many estimation methods to obtain the estimate $\hat{A}$ available in literatures, for example, the Prasad-Rao method-of-moments estimator (Prasad & Rao, 1990), the Fay-Herriot method-of-moment estimator (Fay & Herriot, 1979), the profile maximum likelihood estimator (Hartley & Rao, 1967), and the residual maximum likelihood estimator (Patterson & Thompson, 1971). Among these methods, the residual maximum likelihood estimation method, called as REML method, has been shown to outperform other methods because REML estimator is a second-order unbiased estimator of A . Therefore, in this study, we use the REML method to obtain the estimator of variance model variance $\hat{A}$.

The residual maximum likelihood estimator $\hat{A}$ can be obtained by maximizing the residual log-likelihood function

$$\ell(A) = -\frac{n}{2} \log(2\pi) - \frac{1}{2} \log \left| \sum_{i=1}^m \frac{\mathbf{x}_i \mathbf{x}_i'}{A + D_i} \right| - \frac{1}{2} \log \left| \text{diag } A + D_i \right|_{1 \leq i \leq m} - \frac{1}{2} \sum_{i=1}^m \frac{(y_i - \mathbf{x}_i' \tilde{\boldsymbol{\beta}})' (y_i - \mathbf{x}_i' \tilde{\boldsymbol{\beta}})}{A + D_i}.$$

Using the REML estimator of A , a second-order unbiased (or nearly unbiased) estimator of MSE of the EBLUP is given by (Rao & Molina, 2015)

$$\hat{M}[\hat{\theta}_i] = g_{1i}(\hat{B}_i) + g_{2i}(\hat{B}_i) + 2g_{3i}(\hat{B}_i), \quad (4)$$

where $\hat{B}_i = \frac{D_i}{\hat{A} + D_i}$, $g_{1i}(\hat{B}_i) = D_i(1 - \hat{B}_i)$ is the MSE estimator of θ_i ,

$g_{2i}(\hat{B}_i) = \hat{B}_i^2 x_i' \left(\sum_{j=1}^m x_j x_j' \left(\frac{\hat{B}_j}{D_j} \right) \right)^{-1} x_i$ is the excess in MSE estimator due to estimation of β ,

$g_{3i}(\hat{B}_i) = 2 \left(\frac{\hat{B}_i^3}{D_i} \right) \left(\sum_{j=1}^m \left(\frac{\hat{B}_j}{D_j} \right)^2 \right)^{-1}$ is the excess in MSE estimator due to estimation of A .

2.2 Hierarchical Bayesian Fay-Herriot model

In this section, we discuss the hierarchical Bayesian approach (Gelman et al., 1995; Albert, 2009) to the Fay-Herriot model, called as the hierarchical Bayesian Fay-Herriot model (Rao & Molina, 2015, Chapter 10). The model is then written as follows.

- (i) $y_i | \theta_i, D_i \sim N(\theta_i, D_i)$, $i = 1, \dots, m$, independent
- (ii) $\theta_i | \beta, A \sim N(x_i' \beta, A)$, $i = 1, \dots, m$, independent
- (iii) $A \sim f(A)$, $\beta \sim g(\beta)$,

where f and g are prior distributions of the model variance A and regression coefficient β , respectively. The prior distributions may be informative priors or noninformative priors. The informative priors can be determined based on substantial prior information such as previous study related to the data. However, informative priors are infrequently available. Therefore, noninformation priors are commonly used in many applications.

The hierarchical Bayes (HB) estimate of θ_i is the posterior mean given by

$$E[\theta_i | y] = \int \theta_i f(\theta_i | y) d\theta_i,$$

where $f(\theta_i | y)$ is the posterior density.

The corresponding posterior variance of θ_i is given by

$$\text{Var}[\theta_i | y] = \int (\theta_i - E[\theta_i | y])^2 f(\theta_i | y) d\theta_i,$$

Numerical estimates of posterior mean and variance can be obtained via the Markov chain Monte Carlo (MCMC) method (Rao & Molina, 2015). Having drawn an MCMC sample $\{(\beta^{(k,l)}, A^{(k,l)}), k = d + 1, \dots, d + K, l = 1, \dots, L\}$ from the Markov Chain Monte Carlo method using L chains with the sample size where K of each chain and the burn in size of d , the posterior mean of θ_i can be obtained as

$$\hat{\theta}_i^{\text{HB}} = \frac{1}{LK} \sum_{l=1}^L \sum_{k=d+1}^{d+K} \hat{\theta}_i^{\text{B}}(\beta^{(k,l)}, A^{(k,l)}) := \theta_i^{(\cdot, \cdot)}. \quad (6)$$

The corresponding posterior variance is

$$\text{Var}[\theta_i | y] = \frac{d-1}{d} W_i + \frac{1}{d} B_i, \quad (7)$$

where $B_i = \frac{d}{L-1} \sum_{l=1}^L (\theta_i^{(\cdot, l)} - \theta_i^{(\cdot, \cdot)})^2$ is the between-run variance and $W_i = \frac{1}{L(d-1)} \sum_{l=1}^L \sum_{k=d+1}^{d+K} (\theta_i^{(k, l)} - \theta_i^{(\cdot, l)})^2$ is the within-run variance with $\theta_i^{(k, l)}$ is the k th retained value in the l th run of the length $2d$ with the first d burn-in iterations deleted such that $\theta_i^{(\cdot, l)} = \frac{1}{K} \sum_{k=d+1}^{d+K} \theta_i^{(k, l)}$ and $\theta_i^{(\cdot, \cdot)} = \frac{1}{L} \sum_{l=1}^L \theta_i^{(\cdot, l)}$.

The convergence of Markov chain Monte Carlo output can be examined by the “potential scale reduction factor (PSRF)” originally proposed by Gelman & Rubin’s (1992). The PSRF is an estimated factor by which the scale of the current distribution for the target distribution. Each PSRF reduces to 1 as the number of iterations approaches infinity. The multivariate extension, multivariate potential scale reduction factor (MPSRF), was introduced by Brooks & Gelman (1998). The values of PSRF and MPSRF close to 1 indicate convergence of the chain.

The model validation for Bayesian models can be done via many different approaches such as the posterior predictive assessment introduced by Datta et al. (1999) and the divergence measure introduced by Laud & Ibrahim (1995) described below.

A posterior Predictive Assessment approach. (Datta et al., 1999)

We define y_{obs} and y_{new} to be the observed and the generated data, respectively. Then define $f(\theta | y_{\text{obs}})$, $f(d(y_{\text{obs}}, \theta) | y_{\text{obs}})$, and $f(d(y_{\text{new}}, \theta) | y_{\text{obs}})$ to be the posterior (predictive) predictions of θ , $d(y_{\text{obs}}, \theta)$, and $d(y_{\text{new}}, \theta)$, respectively. The discrepancy measure is

$$d(y, \theta) = \sum_{i=1}^m \sum_{j=1}^n \frac{(y_{ij} - \theta_{ij})^2}{\sigma_{ij}},$$

where m is the length of y , n is the total number of iterations of the θ values, and σ_{ij} is the variance of y_{ij} . The value $P(d(y_{\text{new}}, \theta) \geq d(y_{\text{obs}}, \theta) | y_{\text{obs}})$ of a model close to 0.5 indicates that the model is adequately fit to the data. The extreme value (close to 0 or 1) indicates a lack of fit.

Divergence measure. (Laud & Ibrahim, 1995)

The divergence measure between the observed data y_{obs} and the generated data y_{new} is defined as

$$d(y_{new}, y_{obs}) = E[m^{-1} \|y_{new} - y_{obs}\|^2 | y_{obs}],$$

where m is the length of y_{obs} . An adequate model should produce a small value of the estimated divergence measure.

Data Analysis

In this section, we apply the Classical Fay-Herriot model and the Hierarchical Bayesian Fay-Herriot model to the average household debt data in Thailand demonstrated in previous section. The structure of this section is as follows. We first explain data description and variable selection of explanatory variables used in the study. We then discuss model diagnostics of the two models, compare numerical results of the two models with the classical direct estimators. Finally, we discuss the results on the average household debt in Thailand using estimates from the two models.

3.1 Data Description

The data used in this paper is the average household debt data in Thailand from the Household Socio-Economic Survey (SES) 2017 (National Statistical Office, 2017). The SES is conducted yearly by the National Statistical Office Thailand. The total sample of SES 2017 is 43,210 households which are distributed in 77 provinces. Each province (except Bangkok) is divided into two parts according to the type of local administration area, namely, municipal area and non-municipal area. The design sampling of SES is a stratified two-stage sampling. The basic concepts of stratified two-stage sampling is as follows. (Sukhatme, 1984; Särndal et al., 1992; and Cochran, 1997). The population U is partitioned into H separate strata U_h ($h=1, \dots, H$). For each stratum U_h , consists of N_h separate primary sampling units, called as PSU, U_{hi} for $i=1, \dots, N_h$. For each PSU U_{hi} , consists of M_{hi} secondary sampling units, called as SSU. Let y_{hij} be the values of the quantity of interest in the j th SSU of the i th PSU from the h th stratum. The population mean per second-stage unit in the h th stratum is $\bar{Y}_h = \frac{1}{M_{h0}} \sum_{i=1}^{N_h} \sum_{j=1}^{M_{hi}} y_{hij}$, where $M_{h0} = \sum_{i=1}^{N_h} M_{hi} = N_h \bar{M}_h$. The estimate of the population mean per second-stage unit in the h th stratum is

$$\bar{y}_h = \frac{1}{N_h \bar{M}_h} \frac{N_h}{n_h} \sum_{i=1}^{n_h} \frac{M_{hi}}{m_{hi}} \sum_{j=1}^{m_{hi}} y_{hij} = \frac{1}{n_h \bar{M}_h} \sum_{i=1}^{n_h} \frac{M_{hi}}{m_{hi}} \sum_{j=1}^{m_{hi}} y_{hij}, \quad (8)$$

where n_h is a sample size of the PSUs from the h th stratum, and m_{hi} is the sample size of SSUs from the i th PSU from the h th stratum. The variance of estimate of the population mean is given by

$$\text{Var}[\bar{y}_h] = \left(\frac{1}{n_h} - \frac{1}{N_h} \right) S_h^2 + \frac{1}{n_h N_h} \sum_{i=1}^{N_h} \frac{M_{hi}^2}{\bar{M}_h^2} \left(\frac{1}{m_{hi}} - \frac{1}{M_{hi}} \right) S_{hi}^2, \quad (9)$$

where $S_h^2 = \frac{1}{N_h-1} \sum_{i=1}^{N_h} \left(\frac{M_{hi}}{\bar{M}_h} \bar{y}_{hi} - \bar{Y}_h \right)^2$ is the variance among primary unit means and $S_{hi}^2 = \frac{1}{M_{hi}-1} \sum_{j=1}^{M_{hi}} \left(y_{hij} - \bar{y}_{hi} \right)^2$ is the variance among subunits within primary units. The estimator of the variance is obtained by replacing S_h^2 and S_{hi}^2 with their sample estimators and summing up by sampled PSUs in each h th stratum.

The area-specific explanatory variables used in this paper are chosen by an Akaike's Information Criteria (AIC) backward selection from 12 variables available in the Population and Housing Census 2010 file (National Statistical Office, 2010). The 12 variables in the Housing Census 2010 file, before backward variable selection, are as follows: 1) the proportion of households with detached house; 2) the proportion of households that own living quarters; 3) the proportion of households that rent living quarters; 4) the proportion of households that cement or brick dwellings; 5) the proportion of households that own land; 6) the proportion of households using gas for cooking; 7) the proportion of households using sitting toilet; 8) the proportion of male populations; 9) the average population per private household; 10) the proportion of households that working household head; 11) the proportion of populations that graduated from upper secondary level; 12) the proportion of populations that graduated from higher level.

Having performed the backward variable selection, there are six significantly explanatory variables to be used in our models as shown in Table 1.

Table 1 Explanatory variables

Variable	Notation
Proportion of households that rent living quarters	X_1
Proportion of households that cement or brick dwellings	X_2
Proportion of households that own land	X_3
Proportion of households using gas for cooking	X_4
Proportion of households using sitting toilet	X_5
Proportion of male populations	X_6

3.2 Model diagnostics.

3.2.1 Model diagnostics for the Fay-Herriot model

In this section, we apply the Fay-Herriot model (1) to the average household debt for 153 areas: Bangkok and two areas from each of other 76 provinces which are municipal area and non-municipal area. The direct estimates and sampling variances are obtained from (8) and (9), respectively. The EBLUP estimates for the household debts and a second-order unbiased estimators of MSE of the EBLUP are obtained from (3) and (4), respectively. To diagnose the model, we apply the residual analysis discussed in Erciulescu et al. (2021). We consider the standardized residuals

$$r_i = \frac{y_i - x_i' \hat{\beta}}{\sqrt{\text{Var}[y_i - x_i' \hat{\beta}]}}$$

where $\text{Var}[y_i - x_i' \hat{\beta}] = D_i + A$. Since the parameter A is unknown, the variance is estimated by $D_i + \hat{A}$, where $\hat{A}$ is the estimator of A .

The normality analysis used are 1) Q-Q plot; 2) histogram; 3) Kolmogorov-Smirnov test; 4) Shapiro-Wilk normality test, of standardized residuals of direct estimate from synthetic part.

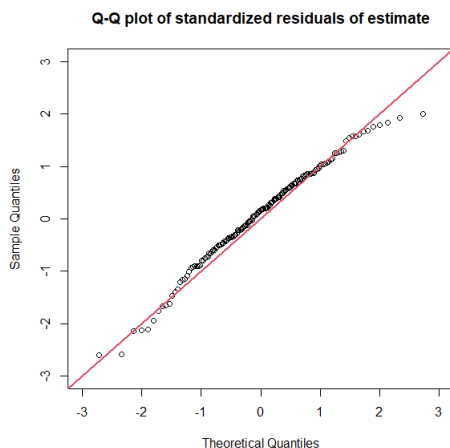


Figure 1 Q-Q plot of standardized residuals of estimate

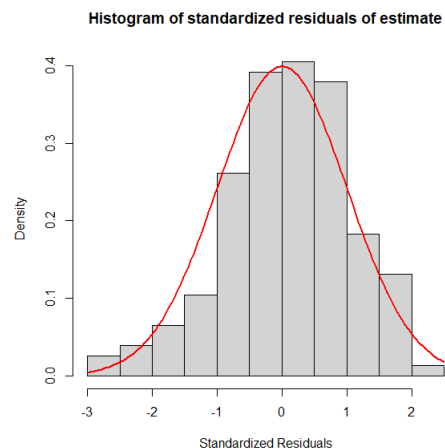


Figure 2 Histogram of standardized residuals of estimate

From Figures 1 - 2, we can see that the standardized residuals from the synthetic part follow normal distribution. Moreover, we also perform the Kolmogorov-Smirnov and Shapiro-Wilk normality tests for the standardized residuals. The corresponding p-values of the tests are 0.3940 and 0.0784, respectively. From the three analyses, we can conclude that the standardized residuals follow normal distribution at the significant level $\alpha = 0.05$.

3.2.2 Model diagnostic for the Hierarchical Bayesian Fay-Herriot model.

In this section, we apply the hierarchical Bayesian Fay-Herriot model (5) to the average household debts. Since there is no further information available, we use noninformative prior distributions for the unknown parameters. Several choices of noninformation priors can be applied such as uniform distribution, normal distribution, and inverse-gamma distribution. In this study, we use common choices which are the uniform prior for the variance and normal prior for the regression coefficients. To be specific, the prior distributions of the model variance A and regression coefficients β are $A \sim U(0, 10^8)$ and $\beta \sim N_7(0, 10^6 I_7)$, respectively. The MCMC samples are obtained by using **R2OpenBUGS** (Thomas, 2020), **R2WinBUGS** (Ligges, 2015) and **coda** (Plummer, 2020) packages in R program (R Core Team, 2020). The setting of the Monte Carlo simulation is as follows: the number of Markov chains is 3, the sample size per chain is 500000, the length of burn in is 250000, and the thinning rate is 2. Having obtained MCMC sample, we apply the PSRF and MPSRF convergence tests and the model diagnostics mentioned in Section 2. The values of PSRFs are in the range of (0.9999,

1.0001) and the MPSRF is 1.0007 which are close to 1. Therefore, the PSRFs and MPSRF results show the convergence of the MCMC chains. The posterior Predictive Assessment value is 0.4792 which is close to 0.5, and the divergence measure value is 0.3922 which is relatively small. The two model validations indicate an adequate fit of the model to the data.

3.3 Model Comparisons.

In this section, we compare performances of the two models with direct estimates. The values of direct, EBLUP and HB estimators of average household debts ordered by the magnitudes of sampling errors are demonstrated in Figure 3. The corresponding MSEs are shown in Figure 4.

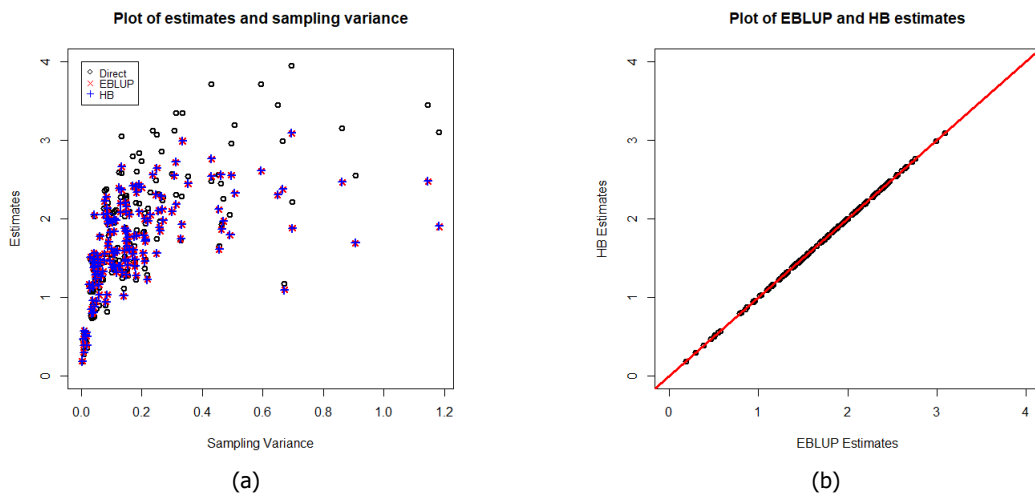


Figure 3 Plot of direct, EBLUP and HB estimates of average household debt

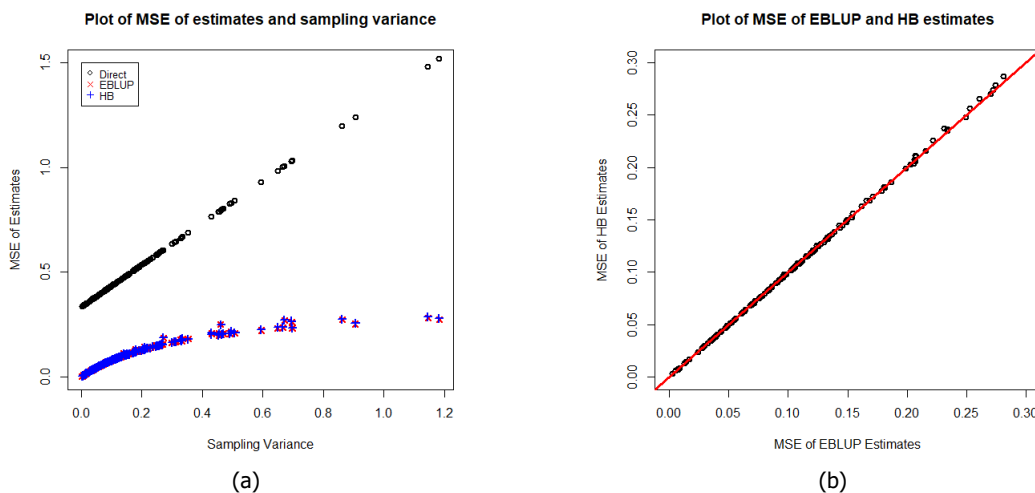


Figure 4 Plot of MSE of direct, EBLUP and HB estimates of household debt

From Figure 3 (a), we can see that EBLUP and HB estimates are close to direct estimates for the areas with small sampling variances but the two estimates are different from the direct

estimates for the areas with large sampling variances. From Figure 3 (b), we can see that the EBLUP and HB estimates are almost identical as the plot between the two estimates goes along with the 45-degree line. From Figure 4 (a), we can see that the MSEs of the direct estimates are larger than MSEs of the EBLUP and HB estimates particularly for areas with large sampling variances. From Figure 4 (b), the MSEs of the direct estimates steady increase when sampling variances increase, while the MSEs of the EBLUP and HB estimates seem to be stable for a certain value of sampling variance, at 0.3 in this dataset. Results from Figure 3 and Figure 4 indicate the lack of fit of the direct estimates for data for areas with large sampling variances while the Fay-Herriot model and the hierarchical Bayesian Fay-Herriot model can improve estimates for such cases by reducing the MSEs. Comparing the the Fay-Herriot model and the hierarchical Bayesian Fay-Herriot model, Figure 3 (b) and Figure 4 (b) show that two estimates are approximately the same.

3.4 Discussions on the Household Debts in Thailand

In this section, we discuss the household debts in Thailand by using estimates from the Fay-Herriot model and the hierarchical Bayesian Fay-Herriot. Table 2 shows average household debts for Bangkok and the five regions of Thailand. The average household debts are presented for municipal area, non-municipal area, and the combined estimates. Figure 5 shows the boxplots of average household debts for each region. From Table 2 and Figure 5, we can see that the average household debts in municipal areas are higher than those of non-municipal areas for all regions. Moreover, the average household debts in Bangkok (BKK) are higher than average debts in other regions of Thailand. Comparing with the country average (1.7023 for EBLUP, and 1.7030 for HB), we can see that Bangkok, the northern, and the northeast regions have higher debts than the country average while the central, the east and the southern regions have lower debts than the country average.

Table 2 The model-based estimates of average household debt (Unit: 100,000 Baht)

Regions	Municipality	Size	Mean		Standard Deviation	
			EBLUP	HB	EBLUP	HB
Bangkok		1	2.0470	2.0467	-	-
Central	Total	36	1.6759	1.6771	0.6792	0.6807
	Municipal area	18	1.7899	1.7910	0.6527	0.6542
	Non- Municipal area	18	1.5619	1.5631	0.7044	0.7061
East	Total	14	1.4838	1.4839	0.3974	0.3980
	Municipal area	7	1.5983	1.5980	0.3673	0.3677
	Non- Municipal area	7	1.3694	1.3697	0.4205	0.4215
North	Total	34	1.7448	1.7455	0.5813	0.5819
	Municipal area	17	1.9571	1.9577	0.5673	0.5680
	Non- Municipal area	17	1.5325	1.5333	0.5286	0.5293
Northeast	Total	40	1.8136	1.8144	0.4798	0.4813
	Municipal area	20	2.0366	2.0381	0.4537	0.4555
	Non- Municipal area	20	1.5907	1.5907	0.4025	0.4033
South	Total	28	1.6223	1.6231	0.6074	0.6082
	Municipal area	14	1.8638	1.8642	0.5789	0.5793
	Non- Municipal area	14	1.3807	1.3819	0.5529	0.5544
Total	Total	153	1.7023	1.7030	0.5726	0.5737
	Municipal area	76	1.8882	1.8890	0.5509	0.5520
	Non- Municipal area	76	1.5118	1.5125	0.5364	0.5375

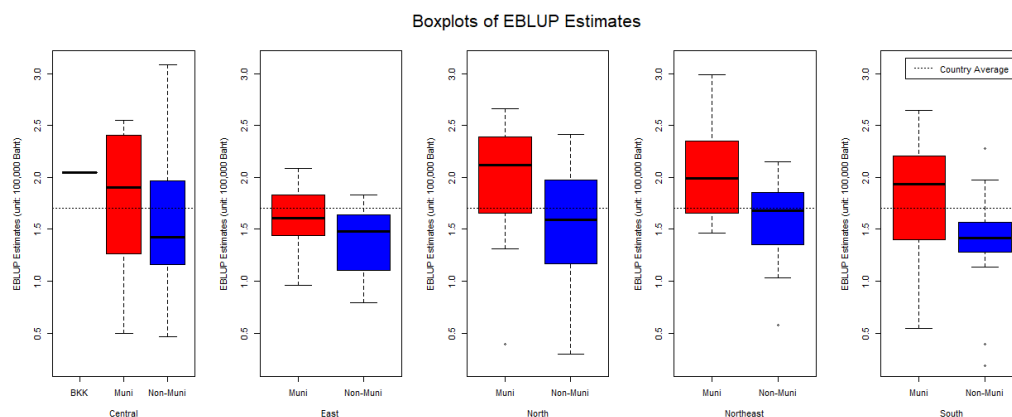


Figure 5 Boxplots of EBLUP estimates for average household debts in different regions

Discussion and Conclusions

In this paper, we applied the two area level models: the Fay-Herriot model and the hierarchical Bayesian Fay-Herriot model to obtain are level estimates of household debts where the area considered in this paper is the municipality by province. The model diagnostics of the two models indicate adequate fit to the data. Model comparisons show that the two small area models outperform the direct estimates especially for areas with large sampling variances. The study indicates that model-based estimates can improve the direct estimates for the household debt data. From our study, we believe that the Fay-Herriot model can be applied to other applications with small sample sizes to overcome the bias of the direct survey estimates caused by sampling errors. Some improvements of our study can be performed such as incorporating time effect into the model to study trends of household debts over years.

Acknowledgement

The authors would like to thank the National Statistical Office Thailand for providing the Household Socio-Economic Survey 2017 data and Population and Housing Census 2010 data used in our study.

References

- Albert, J. (2009). *Bayesian Computation with R* (2nd ed.). Springer, New York.
- Angkunsit, A., & Suntornchost, J. (2021). Bivariate Fay-Herriot models with application to Thai socio-economic data. *Naresuan University Journal: Science and Technology*, 29(1), 36-49. <https://doi.org/10.14456/nujst.2021.3>
- Brooks, S. P., & Gelman, A. (1998). General methods for monitoring convergence of iterative simulations. *Journal of Computational and Graphical Statistics*, 7(4), 434-455. <https://doi.org/10.1080/10618600.1998.10474787>
- Cochran, W. G. (1997). *Sampling Techniques* (3rd ed.). John Wiley and Sons, Inc., New York.
- Datta, G., Lahiri, P., Maiti, T., & Lu, K. (1999). Hierarchical Bayes estimation of unemployment rates for the states of the U.S. *Journal of the American Statistical Association*, 94(448), 1074-1082. <https://doi.org/10.2307/2669921>
- D'Alò, M., Consiglio, L. D., Falorsi, S., & Solari, F. (2008). Small area estimation methods for socio-economic indicators in household surveys. *Rivista Internazionale Di Scienze Sociali*, 116(4), 419-441.

- Erciulescu, A. L., Franco, C., & Lahiri, P. (2021). Use of Administrative Records in Small Area Estimation. In A. Y. Chun & M. Larsen (Eds.), *Administrative Records for Survey Methodology*. John Wiley and Sons, Inc., New York.
- Fay, R. E., & Herriot, R. A. (1979). Estimates of income for small places: an application of James-Stein procedure to census data. *Journal of the American Statistical Association*, 74(366), 269-277. <https://doi.org/10.2307/2286322>
- Gelman, A., Carlin, J., Stern, H., & Rubin, D. (1995). *Bayesian data analysis*. Chapman & Hall, London.
- Gelman, A., & Rubin, D. B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, 7(4), 457-472. <https://doi.org/10.1214/ss/1177011136>
- Hartley, H. O., & Rao, J. N. K. (1967). Maximum-likelihood estimation for the mixed analysis of variance model. *Biometrika*, 54(1-2), 93-108.
- Henderson, C. R. (1975). Best linear unbiased estimation and prediction under selection model. *Biometrics*, 31(2), 423-447. <https://doi.org/10.2307/2529430>
- Kiatipong, A., Wilatluk, S., & Nalin, C. (2007). *The Wealth and Debt of Thai Household: Risk Management and Financial Access*. Bangkok, Thailand: Bank of Thailand Symposium 2007.
- Laud, P. W., & Ibrahim, J. G. (1995). Predictive model selection. *Journal of the Royal Statistical Society. Series B (Methodological)*, 57(1), 247-262. <https://doi.org/10.1111/j.2517-6161.1995.tb02028.x>
- Liggers, U. (2015) *Running 'WinBUGS' and 'OpenBUGS' from 'R' / 'S-PLUS'*. <https://cran.r-project.org/web/packages/R2WinBUGS/index.html>
- Muthitacharoen, A., Nuntramas P., & Chotewattanakul P. (2015). Rising household debt: Implications for economic stability. *Thailand and The World Economy*, 33(3), 43-65.
- National Statistical Office. (2010). Population and Housing Census 2010. Thailand.
- National Statistical Office. (2017). The Household Socio-Economic Survey 2017. Thailand.
- Patterson, H. D., & Thompson, R. (1971). Recovery of inter-block information when block sizes are unequal. *Biometrika*, 58(3), 545-554.
- Prasad, N. G. N., & Rao, J. N. K. (1990). The estimation of the mean squared error of small-area estimators. *Journal of the American Statistical Association*, 85(409), 163-171.
- Plummer, M. (2020). *Output Analysis and Diagnostics for MCMC*. <https://cran.r-project.org/web/packages/coda/index.html>
- R Core Team, R. (2020). *R: A language and environment for statistical computing*. Vienna, Austria. Retrieved from <https://www.R-project.org/>
- Rao, J. N. K., & Molina, I. (2015). *Small Area Estimation* (2nd ed.). John Wiley and Sons Inc., New York.
- Särndal, C. E., Swensson, B., & Wretman, J. (1992). *Model Assisted Survey Sampling*. Springer-Verlag, New York.
- Srivastava, A. K., Sud, U. C., & Chandra, H. (2007). Small area estimation – an application to national sample survey data. *Journal of the Indian Society of Agricultural Statistics*, 61(2), 249-254.
- Sukhatme, P. V. (1984). *Sampling Theory of Surveys with Applications*. Ioxa State University Press, Ames, IA.
- Thomas, N. (2020). *Running OpenBUGS from R*. <https://cran.r-project.org/web/packages/R2OpenBUGS/index.html>

- Visitwarakorn, W. (2015). Debt of government officers after the 2013 salary policy: A case study of the office of the permanent secretary, ministry of natural resources and environment, *Varidian E Journal*, 1, 515-529.
- Wawrowski, L. (2016). The spatial Fay-Herriot model in poverty estimation. *Folia Oeconomica Stetinensia*, 16(2), 191-202. <https://doi.org/10.1515/foli-2016-0034>

Research Article

Significant Attributes from Multiple Regression and Marginal Effects on Multinomial Logit Model for Chemical Servitization

Tanyaporn Kanignant^{1*}, Pongsa Pornchaiwiseskul², Kamonchanok Suthiwartnarueput³

¹Logistics and Supply Chain Management Interdisciplinary Program, Chulalongkorn University, Bangkok, Thailand

²Faculty of Economics, Chulalongkorn University, Bangkok, Thailand

³Chulalongkorn Business School, Chulalongkorn University, Bangkok, Thailand

*E-mail: tanya.kanignant@gmail.com; ppongasa@pop.co.th; kamonchanok.s@chula.ac.th

Received: 11/05/2021; Revised: 20/07/2021; Accepted: 19/08/2021

Abstract

Chemical industry traditionally produces and sells tangible goods. Recently, firms in chemical industry provide additional services to their customers. Several manufacturers change from tangible product suppliers to both product and service providers. This movement is called servitization (Vandermerwe & Rada, 1988). Chemical servitization levels can be classified into 4 categories which are product only, service added to the product, service differential the product and service is the product (Thoben, Eschenbacher, & Jagdev, 2001). The objectives of this paper are to construct servitization framework for chemical suppliers to shift to product service integration and to examine factors affecting chemical service levels to provide guidance to chemical suppliers to implement product service system (Kortman, Theodori, Ewijk, Verspeek, & Uitinger, 2006). The first part of the framework is to develop servitization levels for chemical industry in Thailand. The second part is to define servitization levels for suppliers to offer to their customers. Questionnaire surveys were distributed to chemical dealers, sub-dealers, and end-users, and the sample size was 200. To accomplish the research objective, descriptive statistics, Analytical Hierarchy Process (AHP), Multiple Regression Analysis, and Multinomial Logit Model (MNL) used in this research. The finding includes seven significant factors which were identified in order to analyze the service level of customer needs. Implications and suggestions for suppliers who want to change their business model to providing chemical solution should offer chemical blending, chemical storage, chemical documentation, and environmental and safety program as bundle services with chemical products.

Keywords: Chemical Servitization, Servitization Levels, Product Service System, Extended Product, Multinomial Logit Model, Multiple Regression Analysis

Introduction

Servitization concepts have been introduced to explain the idea that manufacturers or producers turn out to be service providers (Buschak & Lay, 2014; Goedkoop, 1999; Tukker, 2004; Vandermerwe & Rada, 1988). This concept has been applied in many industries including chemical industry. Chemical servitization is a new trend for companies in chemical industry to change their focus to gain competitive advantages and leave out cost competition to win against competitors (Kortman, Theodori, Ewijk, Verspeek, & Uitinger, 2006; Robinson, Clarke-Hill, & Clarkson, 2002). Chemical is one of the most important industry that its products are widely used in our daily lives. Consumers are influenced by chemicals in many ways such that we consume food, housekeeping, painting pharmaceuticals agriculture, construction, adhesive, and textile products. As the range of chemical product chain is too wide to concentrate, this study will focus only on commodity chemicals products in B2B business type in a perception that chemicals are used as raw materials for manufactures to produce finished goods.

The organizational changes in traditional manufacturers to new trend of servitization have been developed since the last two decades. "Power by hour" is a shift in offering from aero engines selling to providing a total care package developed by Rolls-Royce Aerospace to its customers such as Boeing

and Airbus. The new business model combines tangible products with intangible services of maintenances. Revenue generates from charging customers based on hours of engine used. Another organizational change is the example of SEFCHEM, Dow Chemical's subsidiary company which provides solvent service solutions. SEFCHEM cooperated with Pero AG established a new company to offer cleaning services. The new cleaning full-service company provides cleaning services, uses cleaning machines from Pero AG, and attains chemical supplies from SEFCHEM (Buschak & Lay, 2014).

Chemical products are defined as commodity products which are used as raw materials for manufactures and can be transformed to intermediate and specialty chemicals. They are also sold by volume with standardized quality with few variants. The commodities are in high market competition because price is the key buying criteria for buyers. Thus, any suppliers who offer lower prices will be more attractive to customers than the suppliers who charge higher prices. When a chemical firm cannot charge customers in high prices, the firm is in a struggle situation namely "commodity trap". Robinson et al. (2002) studied servitization model which is a strategy that helps companies to drip out the commodity trap, achieve competitive advantages and seek for differentiation instead. The servitization strategy is an approach for companies changing from traditionally cost oriented to service and relationship management. Servitization is also one element of logistics 4.0 trends for sustainable business model to transform enterprises from tangible product to service-oriented that can increase the value proposition by integrating services and manufacturing processes in their offers (Strandhagen et al., 2017).

Major problems of Thai chemical providers are high competitiveness markets, price sensitivity, volume based selling with low margin, limited services with low value, and tangible goods business model. Under the uncertainty economic condition, Thai chemical providers are also facing the same problems as others in other parts of the world. These companies need to change their focus of their business as well.

The objectives of this research are to construct servitization framework for chemical suppliers and to investigate determinant affecting chemical service levels to provide a guidance to companies in chemical industry to implement product service system.

Materials and methods

Chemical Servitization

Servitization is popularly adopted for innovative business model development in chemical industry to help customers avoid chemical waste. It is used as a link between physical offers and additional services provided to customers (Buschak & Lay, 2014). The innovative business models for chemicals industry can be described as chemical product services (CPS) where business models shift from selling chemical products by volume to combining with some basic services to fulfil customers and suppliers' requirements (Kortman et al., 2006); chemical management services (CMS) where business models create a long-term collaboration between customers and chemical service providers to supply and manage chemical related services (Stoughton & Votta, 2003); and chemical leasing where chemical companies supply special services and substances but hold the ownership of chemicals. The traditional business models for chemical industry was focusing on selling chemical products by volume. This leads to conflicts between customers and suppliers because the customers want to decrease chemical volume and cost while the suppliers want to maximize sales volume (Kortman et al., 2006; Reiskin, White, Johnson, & Votta, 2000; Toffel, 2008).

Chemical product service (CPS) business model aligns the interests of customers and chemical suppliers that both of them receive benefit from reducing chemical sales volume and cost. The suppliers are no longer focusing on selling chemical product by volume (Kortman et al., 2006), see Figure 1.

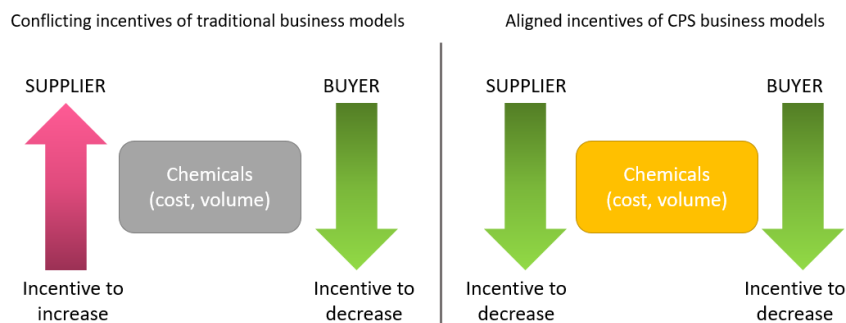


Figure 1 Traditional and CPS Business Models (Kortman et al., 2006)

Business models for CPS are variety by adding some more extra services, and they can be related to the several stages in the chemical life cycle. Kortman et al. (2006) recommended some extra chemical services; for example, chemical packaging, chemical management, chemical inventory and storage, chemical advice on process tuning, chemical transportation, chemical recycling and waste treatment, chemical health concern, environmental and safety programs, and worker's training. Kortman et al. (2006) classified CPS into two different types as CPS-I and CPS-II as the transition from traditional business models to CPS models (See Figure 2). CPS-I is a business model that chemical producers or suppliers are still selling chemical products by volume. To increase value of the chemical product, some related services are added to the products. CPS-II is a business model chemical suppliers provide product service integrated solutions regarding to customer requirements instead of offering products by volume. The ownership of chemical product is fully transferred from suppliers to customers.

Chemical management service (CMS) is a business model for chemical products that both suppliers and customers collaborate each other to improve and develop chemical product services (Stoughton & Votta, 2003) in term of partnership (Reiskin et al., 2000). Example of CMS are chemical supply, chemical quality monitoring, chemical adjustment, removal of applied chemical, chemical recycling, and chemical solution network. The benefits of CMS mentioned by Kortman et al. (2006) are liability reduction, storage space decreasing, chemical labor reduction, and heal and environmental saving.

Chemical leasing is a business model that the ownership of chemical product is still on the suppliers, not customers. This means the main concentration is not on selling chemical product by volume but on the integrated services offered with the products. Thus, profit is not directly from selling chemical product in large volume, but it comes from bundled services instead (Kortman et al., 2006).

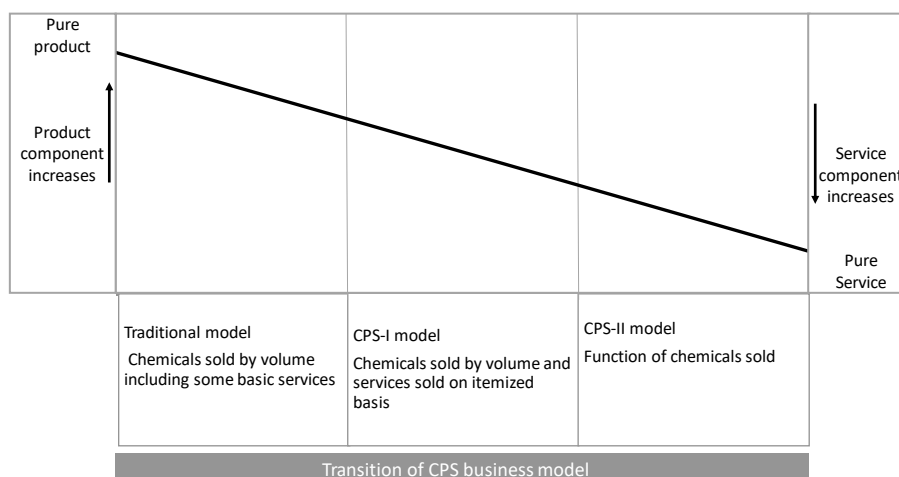


Figure 2 Transition from the traditional business models to the CPS models (Kortman et al., 2006)

Many literatures also talk about the same concepts but in different terms. Product Service System (PSS) (Tukker, 2004), Product Transition (Oliva & Kallenberg, 2003), and Extended Product (Thoben, Eschenbacher, & Jagdev, 2001) are phrases commonly used when the manufacturing firms apply the servitization concepts. The below section is related theory applied for the servitization used in this study.

Extended Product Theory

To have advantages in competitive market, manufacturers and suppliers have to integrate their core products with additional services to make their products more valuable and attractive. This concept is defined as Extended Product (Thoben, Eschenbacher, & Jagdev, 2001), which consists of three layers, the kernel as an illustration of the core and functionalities of product (tangible), the middle layer describing the product shell including packaging of the core product (packaging), and the outer layer representing all the intangible assets of the offer (services), see Figure 3.

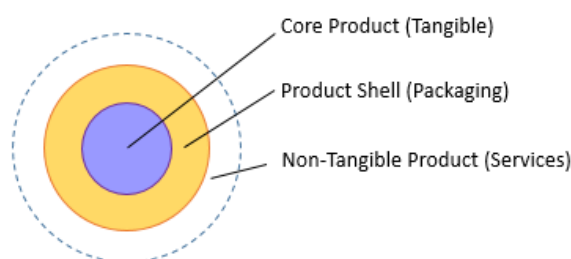


Figure 3 Extended Product Theory (Thoben et al., 2001).

A combination of core product and the product shell is called narrow sense which tangible products are offered to the market, whereas a blending between product shell and non-tangible product is named a broader sense as a product solution that both tangible and intangible products are integrated together (Thoben et al., 2001). Figure 4 illustrates dimension of migration process based on the expanded product concept transforming from tangible product to intangible services and finally service as product (Chen & Gusmeroli, 2015).

Servitization levels are also mentioned in chemical industry in similar ways as in other manufacturing industries. The starting point is the pure manufacturer traditionally provide chemical product in large volume. The next level is chemical supplier offers some product related services such as transportation and worker training. Chemical supplier may also provide other different services not directly related to chemical product such as product monitoring system. Lastly, in the highest level, chemical suppliers focus on providing intangible services with the add on tangible products (Buschak & Lay, 2014; Chen & Gusmeroli, 2015; Kortman et al., 2006). Example of product as a service is chemical trend that SAFECOM cooperates with Pero AG company to provide cleaning services to their customers rather than selling tangible chemical products (Buschak & Lay, 2014).

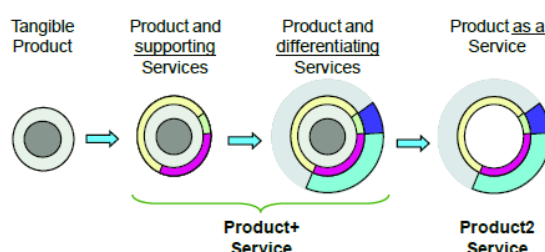


Figure 4 Extended Product Elements (Chen & Gusmeroli, 2015)

Servitization Frameworks

Ryu, Rhim, Park, and Kim (2012) proposed servitization framework adapted from Meyer and Arthur (1999). This framework composed of three components which are markets, product-service-knowledge system (PSK) (see Table 1), and competencies in the supply chain. Chen and Gusmeroli (2015), Oliva

and Kallenberg (2003) proposed a framework for manufacturing servitization which combined three dimensions; 1) the x-axis represents servitization process; 2) the y-axis represents stages of product extension; 3) the z-axis is illustrated service innovation. However, none of these papers proposed guidance or solutions on servitization process. Thus, this research closes the gap by constructing the servitization framework for chemical suppliers observing factors affecting chemical service levels to provide guidance to chemical suppliers to implement product service system.

This study proposed new servitization framework adapted from the previous studies and illustrated as Figure 5. The first part of the framework begins with Chen and Gusmeroli (2015) framework, and ends with Ryu et al. (2012) in the second part. The three dimensions are changed to; 1) the x-axis represents customer segment which classified by 3 different company sizes and 5 types of industry; 2) y-axis represents 4 different types of servitization levels, namely product only, service added to the product, service differential the product, and service is the product; 3) z-axis represents PSK system. PSK is classified to three service types dealing with product, service, and knowledge. The second part was adopted from (Kanignat et al., 2018).

Research Methodology

Scope of the Study

The chemical products mentioned in this research are chemical products which are considered as commodity products. Size of chemical companies are defined as number of employees based on OSMEP (2000) which can be classified into three groups of small (less than 50 employees), medium (50-200 employees), and large (more than 200 employees). Respondents in this research are separated by types of industry which can be divided into five groups of: Industrials including adhesive, ink, packaging, paint, petrochemicals, resin, thinner, tire (wheel); Consumer Product including cosmetics, food, pharmaceutical; Resources for example mining; Technology for example electronics; and Others.

The study focuses in chemical industry only in Thailand and approaches one B2B business company of tier-3 who is a chemical importer or distributor traditionally provides tangible chemical products for their customers in large volume and have high competitive market. Chemical product in this study is defined as commodity product that has similar property. It is also price sensitive and is often sold in bulky amount. The company's customers are: tier-2 firms who provide chemical products as wholesalers, tier-1 companies who perform as sub dealers supplying chemical products to manufacturers, and the end-users who are manufacturers using chemical products as raw materials in production to make products. The study studies servitization strategies for this distributor company to generate product transition for customers.

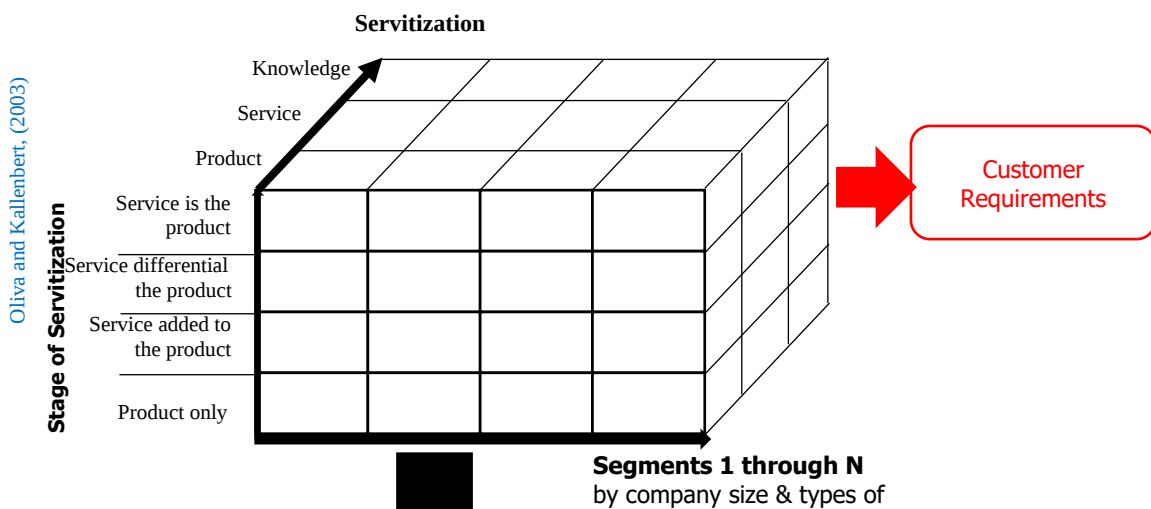
Respondents in this research are separated by position of companies in chemical supply chain, see Figure 6, which can be divided into three groups as 1) end-users or manufacturers, 2) tier-1: sub dealers or suppliers, and 3) tier-2: dealers or wholesalers. This research does not include respondents who are upstream producers or oversea and local makers, tier-3 companies who are importers or distributors, and consumers. Figure 6 illustrates chemical product supply chain for an easier point of view of targeted respondents.

Table 1 Product-Service-Knowledge System

Product	Service	Knowledge
Chemical product only	Chemical document and license	Chemical health risk assessment
Chemical blending	Chemical inventory	Environmental and safety programs
Chemical packaging	Chemical waste treatment	Worker's training
Chemical storage		
Chemical recycling		
Transportation		

Servitization Model for Chemical Products

Markets Customer Segmentation



Servitization process

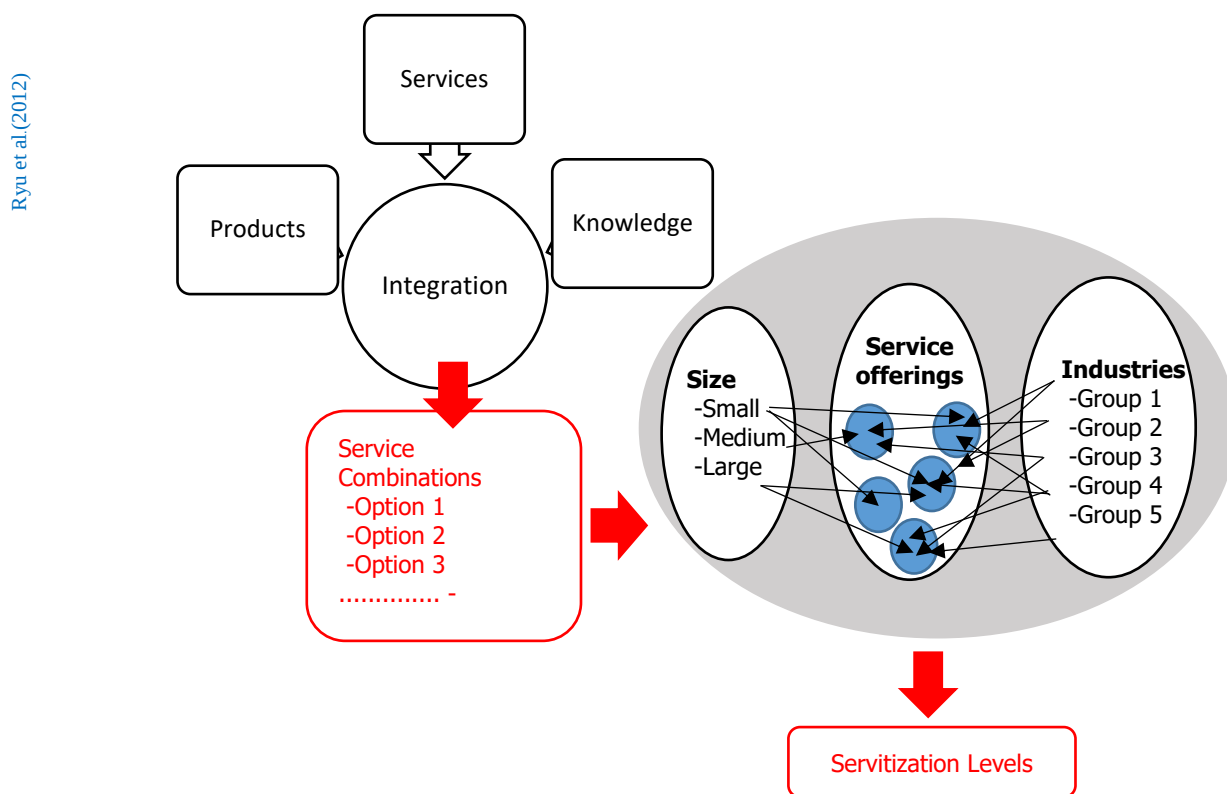


Figure 5 Servitization Model for Chemical Product

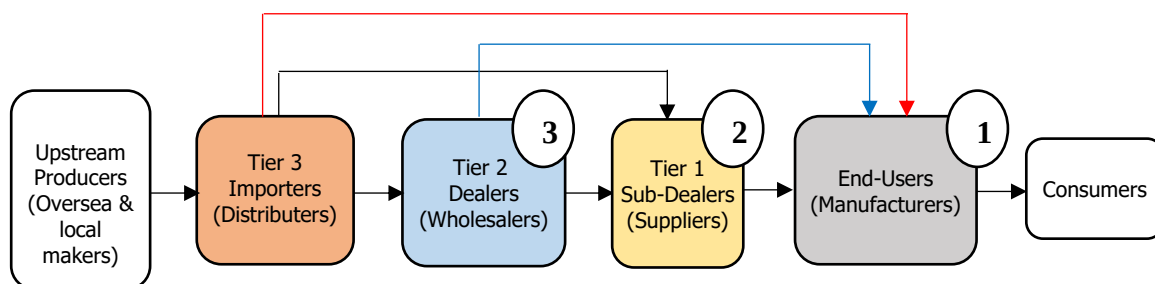


Figure 6 Chemical Supply Chain

Data Collection Research Tools and Design

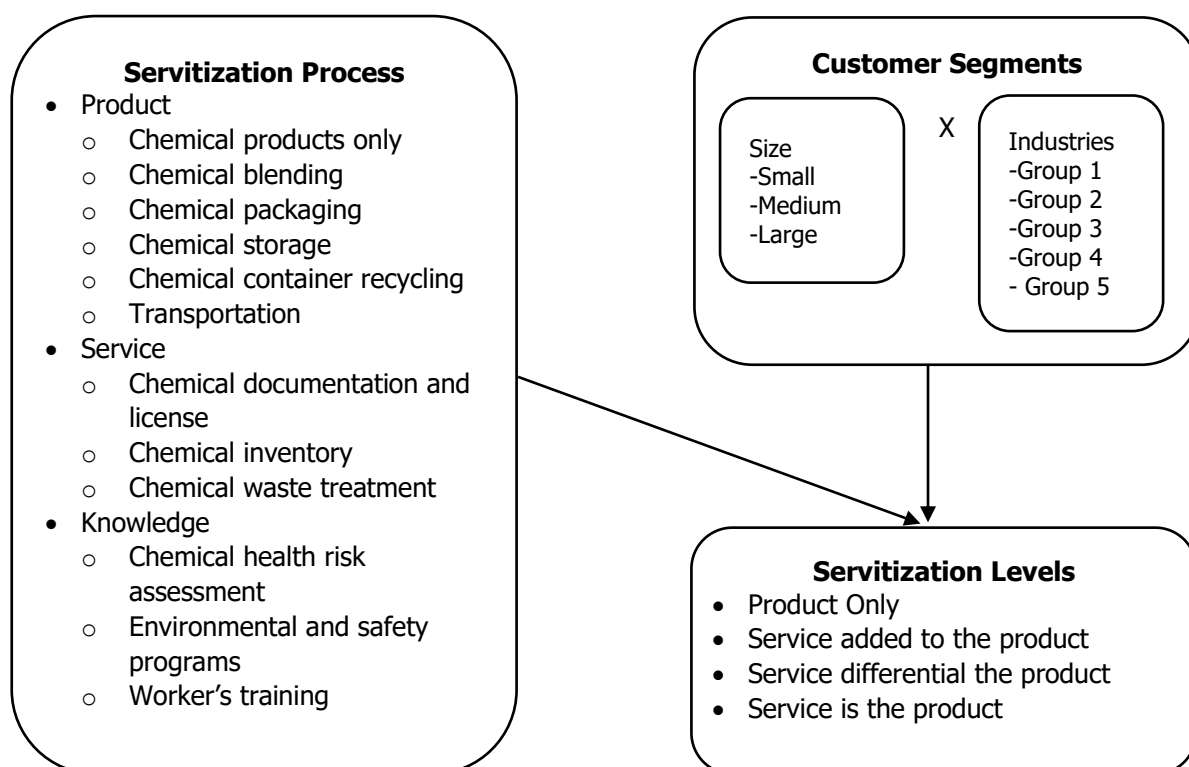


Figure 7 Relationship between Independent and Dependent Variables

The research tools for this study is questionnaire survey distributed to respondents via face to face or an interview. The questionnaire composed of 3 sections; 1) company background, 2) attitude towards product or service needed driven by 10-point Likert Scale ranging from 0 to 10 to employ the questions and scale responses in the survey, and 3) comparison attitude towards servitization levels constructed by Analytical Hierarchy Process (AHP) using the pairwise comparison between 4 service levels. In the questionnaire design process, required data for composing questionnaire is assemble from literature reviews and discussion with the staff of chemical distributor.

Index of Item-Objective Congruence (IOC) was used to analyze the content validity. The questionnaire was reviewed by five experts including two chemical management officers, and three academic experts. The reliability of the questionnaire was examined in order to confirm that the collected responses were reliable and consistent. The researcher distributed 30 pilot questionnaires to staff of the chemical distributor company to ask their customers excluded from the sample group. For the pilot data reliability test, Cronbach's Alpha score of each question was greater than 0.9. This can be assumed that the questionnaire was highly reliable. The chemical distributor company has almost 250 customers in

Thailand in various locations, and the sample size for this study is 200. The survey was started from May to September 2020 through phone interview only according to COVID-19 situation until the sample size was achieved. The collected data is sufficient enough to do the data analysis and estimate parameters of this study.

As mentioned, parameters in the model are defined by the 3-axis as follows:

- X-axis is the independent parameter represents customer segments.
- Y-axis is the dependent parameter contains 4 different types of servitization levels.
- Z-axis is the independent parameter of PSK.

From these three dimensions, the relationship between independent and dependent variables can be illustrated as Figure 7.

Multinomial Logit Model (MNL)

MNL was utilized in this research rather than Multinomial Logistic Regression (MLR) as the sample size was limited. Peduzzi, Concato, Kemper, Holford, and Feinstein (1996) suggested that MLR was a good technique for large sample sets. The discrete choice model or multinomial logit model was developed by McFadden (1973) and applied in the study of travel mode choices, for example; the choice between bus, car, train, or airplane. The objective is to estimate probability of choosing each of the four modes and to calculate the odds ratios for choice of different modes. The simple MNL can be written as:

$$U_{nj} = \beta x_{nj} + \varepsilon_{nj}. \quad (1)$$

Where

$$\begin{array}{ll} U_{nj} & = \text{the utility of alternate } j \text{ to individual } n, \\ x_{nj} & = \text{J-vector of observed attributes of alternative } j \\ \beta & = \text{a vector of utility weights} \end{array} \quad \begin{array}{ll} \varepsilon_{nj} & = \text{an error} \\ n & = 1, \dots, N \\ j & = 1, \dots, J \end{array}$$

The probability that person n chooses alternative j is given by:

$$\Pr(j | x_n) = \frac{e^{\beta x_{nj}}}{\sum_{k=1}^J e^{\beta x_{nk}}} = \frac{e^{g_j(x)}}{\sum_{k=1}^J e^{g_k(x)}}. \quad (2)$$

In this research study, the dependent variables are categories of servitization level: 1 = product only, 2 = services added to the product, 3 = service differential the product, and 4 = service is the product. For each choice of dependent variable, assume that p covariates and has a constant term, denoted by the vector x , of length $p + 1$, where $x_0 = 1$, the multinomial logit model with the value of dependent variable $Y = 1$ as a reference outcome can be expressed as:

$$g_1(x) = \ln \left[\frac{\Pr(Y = 2|x)}{\Pr(Y = 1|x)} \right] = \beta_{10} + \beta_{11}\chi_1 + \beta_{12}\chi_2 + \dots + \beta_{1p}\chi_p = x'\beta_1. \quad (3)$$

$$g_2(x) = \ln \left[\frac{\Pr(Y = 3|x)}{\Pr(Y = 1|x)} \right] = \beta_{20} + \beta_{21}\chi_1 + \beta_{22}\chi_2 + \dots + \beta_{2p}\chi_p = x'\beta_2. \quad (4)$$

$$g_3(x) = \ln \left[\frac{\Pr(Y = 4|x)}{\Pr(Y = 1|x)} \right] = \beta_{30} + \beta_{31}\chi_1 + \beta_{32}\chi_2 + \dots + \beta_{3p}\chi_p = x'\beta_3. \quad (5)$$

Then the conditional probabilities of each outcome category are:

$$\Pr(Y = 1|x) = \frac{1}{1 + e^{g_1(x)} + e^{g_2(x)} + e^{g_3(x)}}. \quad (6)$$

$$\Pr(Y = 2|x) = \frac{e^{g_1(x)}}{1 + e^{g_1(x)} + e^{g_2(x)} + e^{g_3(x)}}. \quad (7)$$

$$\Pr(Y = 3|x) = \frac{e^{g_2(x)}}{1 + e^{g_1(x)} + e^{g_2(x)} + e^{g_3(x)}}. \quad (8)$$

$$\Pr(Y = 4|x) = \frac{e^{g_3(x)}}{1 + e^{g_1(x)} + e^{g_2(x)} + e^{g_3(x)}}. \quad (9)$$

By taking the log and applying the fact that $\sum \Pr(j|x_n) = 1$, all these four equations are associated by consuming the same denominator and by:

$$\Pr(Y = 1|x) + \Pr(Y = 2|x) + \Pr(Y = 3|x) + \Pr(Y = 4|x) = 1. \quad (10)$$

Thus

$$\frac{\partial \Pr(Y = 1|x)}{\partial x} + \frac{\partial \Pr(Y = 2|x)}{\partial x} + \frac{\partial \Pr(Y = 3|x)}{\partial x} + \frac{\partial \Pr(Y=4|x)}{\partial x} = 0. \quad (11)$$

In this study, the outcome of $Y = 1$, product only, is the reference outcome. Marginal effect describes the average effect of changes in independent variables on the changes in the probability of dependent variables in multinomial logit model.

$$\frac{\partial \Pr(Y = 2|x)}{\partial x} = \Pr(Y = 2|x) (1 - \Pr(Y = 2|x))\beta_1 - \Pr(Y = 2|x) \Pr(Y = 3|x) \beta_2 - \Pr(Y = 2|x) \Pr(Y = 4|x) \beta_3. \quad (12)$$

$$\frac{\partial \Pr(Y = 3|x)}{\partial x} = \Pr(Y = 2|x) (Y = 3|x)\beta_1 - \Pr(1 - \Pr(Y = 3|x)) \Pr(Y = 3|x) \beta_2 - \Pr(Y = 3|x) \Pr(Y = 4|x) \beta_3. \quad (13)$$

$$\frac{\partial \Pr(Y = 4|x)}{\partial x} = \Pr(Y = 2|x) (Y = 4|x)\beta_1 - \Pr(Y = 3|x)(Y = 4|x) \beta_2 - \Pr(1 - \Pr(Y = 4|x)) \Pr(Y = 4|x) \beta_3. \quad (14)$$

$$\frac{\partial \Pr(Y = 1|x)}{\partial x} = - \left(\frac{\partial \Pr(Y = 2|x)}{\partial x} + \frac{\partial \Pr(Y = 3|x)}{\partial x} + \frac{\partial \Pr(Y=4|x)}{\partial x} \right). \quad (15)$$

Data Analysis

To accomplish the research objectives, few data analysis techniques are applied including descriptive statistics used to explain respondent demographic information; Analytical Hierarchy Process (AHP) used at the process of customer pairwise-comparison on 4 servitization levels; Multiple Regression Analysis adopted for exploring significant factors; and Multinomial Logit Model (MNL) used to find the Marginal Effect of each significant factors in this research. The dependent variables are unordered choices of 4 servitization levels which will be compared by the customers. Research variables are acquired from literature review and can be defined as shown in Table 2.

Table 2 Variable Coding

No.	Variables	Variable Type	Measurement	Definition
1	LnY2Y1	Dependent	Ratio Scale	Natural logarithm of the probability of $Y = 2$ compared to $Y = 1$
2	LnY3Y1	Dependent	Ratio Scale	Natural logarithm of the probability of $Y = 3$ compared to $Y = 1$
3	LnY4Y1	Dependent	Ratio Scale	Natural logarithm of the probability of $Y = 4$ compared to $Y = 1$
4	Y1	Dependent	Ratio Scale	Probability of an event $Y = 1$, Product Only
1	Y2	Dependent	Ratio Scale	Probability of an event $Y = 2$, Service added to the product
2	Y3	Dependent	Ratio Scale	Probability of an event $Y = 3$, Service differential the product
3	Y4	Dependent	Ratio Scale	Probability of an event $Y = 4$, Service is the product
4	MeanPCP	Independent	Ratio Scale	Average score of chemical product only
5	MeanPCB	Independent	Ratio Scale	Average score of chemical blending
6	MeanPCK	Independent	Ratio Scale	Average score of chemical packaging
7	MeanPCS	Independent	Ratio Scale	Average score of chemical storage
8	MeanPCC	Independent	Ratio Scale	Average score of chemical container recycling
9	MeanPCT	Independent	Ratio Scale	Average score of chemical transportation
10	MeanSCD	Independent	Ratio Scale	Average score of chemical documentation
11	MeanSCI	Independent	Ratio Scale	Average score of chemical inventory
12	MeanSCW	Independent	Ratio Scale	Average score of chemical waste treatment

No.	Variables	Variable Type	Measurement	Definition
13	MeanKCH	Independent	Ratio Scale	Average score of chemical health risk assessment
14	MeanKES	Independent	Ratio Scale	Average score of environmental and safety program
15	MeanKWT	Independent	Ratio Scale	Average score of worker's training
16	Seg	Independent	Nominal	Segment type
17	Type	Independent	Nominal	Company type
18	Size	Independent	Nominal	Company size

Results

Demographic information of respondents was described by frequency and percentage. Table 3 below shows respondent demographic information.

Demographic Information

The majority of customer segment in the chemical supplier company was in industrial (68.5%) followed by consumer segment (24%), technology (6.5%), and resource (1%) varies in several types of company; for example, thinner (13%), food (13%), adhesive (11%), color (9.5%), petrochemical (9%), respectively. The size of customers was almost the same proportion between large (39.5%) and medium (36.5%) companies and the rest is small size (24%). Most of the customers' companies were located in Bangkok and perimeter (77%), and the rest is located in the East (15%), Central (5%), and others (3%) region of Thailand.

Table 3 Respondent Demographic Information

Category	Frequency	Percent (%)	Category	Frequency	Percent (%)
Industry Segment			Company Size		
Industrial	137	68.5	Small (<50)	48	24
Consumer	48	24	Medium (50-200)	73	36.5
Resource	2	1	Large (>200)	79	39.5
Technology	13	6.5	Total	200	100
Others	0	0			
Total	200	100			
Company Type			Year		
Adhesive	22	11	0-5 Years	20	10
Ink	8	4	6-10 Years	30	15
Packaging	15	7.5	10-15 Years	35	17.5
Color	19	9.5	> 15 Years	115	57.5
Petrochemical	18	9	Total	200	100
Resin	6	3			
Thinner	26	13	Location		
			Bangkok and Perimeter	154	77
Tyre (Wheel)	8	4	Central	10	5
Others (Industrial)	16	8	East	30	15
Cosmetic	16	8	North	4	2
Food	26	13	West	2	1
Medicine	3	1.5			

Others (Consumer)	3	1.5	South	0	0
Mining	0	0	Total	200	100
Others (Resource)	1	0.5			
Electronic	11	5.5			
Others (Electronic)	2	1			
Other Industry	0	0			
Total	200	100			

Table 4 explains customer companies by segment and size. It shows that the largest customer segment is the industrial segment dominated by large size companies (66 of 137 or 48%) followed by medium size companies (46 of 137 or 34%) and small size companies (25 of 137 or 18%).

Table 4 Respondent Demographic Information by Segment and Size

Segment / Size	Size			Total
	Small	Medium	Large	
Industrial	25	46	66	137
Consumer	21	19	8	48
Resource	0	1	1	2
Technology	2	7	7	13
Others	0	0	0	0
Total	48	73	79	200

After using AHP technique, probability of each choice of service level is calculated by pairwise comparison from the respondents. 0.1 consistency ratio is the requirement of the qualification of data from each respondent. The independent variables are selected by adopting multiple linear regression between independent variables and log odd value of each service level compared with the base of service level. In this study, product only is performed as the base of service level comparison. For example, the variable LnY2Y1 is natural logarithm of the probability of $Y = 2$ (service added to the product) compared to $Y = 1$ (product only). Multiple linear regression models of each log odd comparison were used to measure the significant level of the influence of independent variables. Only independent variables that meet the criteria of significant level will be carried further to calculate marginal effect in multinomial logit model in order to see the changes caused by these variables. As we have 3 groups of independent and dependent variables, 9 multiple regression models were run for the result. Independent variables from product, service and knowledge categories were plugged-in the model with dependent variables of natural logarithm of the probability of service added to the product compared to product only level (LnY2Y1), natural logarithm of the probability of service differential the product compared to product only level (LnY3Y1), and natural logarithm of the probability of service is the product compared to product only level (LnY4Y1) separately one at a time. Table 5 shows the results of 9 multiple regression models. As the result, independent variables that have significant level less than .05 or .1 will be selected and carried further in multinomial logistic models to find the marginal effects of independent variables toward those four dependent variables.

Table 5 Results of 9 Multiple Regression Models

Model	LnY2Y1		LnY3Y1		LnY4Y1	
	B	Std. Error	B	Std. Error	B	Std. Error
(Constant)	-.379	.620	-.334	.727	-.602	.769
MeanPCP	-.203**	.078	-.329**	.092	-.219**	.097
MeanPCB	-.013	.034	.030	.040	.099**	.043
MeanPCK	.095	.107	.092	.125	-.047	.133
MeanPCS	.020	.046	.049	.055	.162**	.058
MeanPCC	.042	.054	-.002	.063	-.141**	.067
MeanPCT	.164*	.100	.267*	.117	.312**	.124
$R^2 = .238$, Adjusted $R^2 = .057$, Sig. = .077* $R^2 = .275$, Adjusted $R^2 = .076$, Sig. = .018** $R^2 = .284$, Adjusted $R^2 = .081$, Sig. = .012**						

Model	LnY2Y1		LnY3Y1		LnY4Y1	
	B	Std. Error	B	Std. Error	B	Std. Error
(Constant)	-.671	.551	-1.068	.647	-1.600	.684
MeanSCD	.178**	.066	.246**	.078	.240**	.082
MeanSCI	-.005	.056	-.071	.066	.015	.070
MeanSCW	-.026	.053	.021	.062	.022	.065
$R^2 = .190$, Adjusted $R^2 = .036$, Sig. = .066*						
$R^2 = .230$, Adjusted $R^2 = .053$, Sig. = .014**						
$R^2 = .242$, Adjusted $R^2 = .059$, Sig. = .008**						
Model	LnY2Y1		LnY3Y1		LnY4Y1	
	B	Std. Error	B	Std. Error	B	Std. Error
(Constant)	.244	.497	.242	.588	-.633	.618
MeanKCH	-.055	.083	-.036	.098	.021	.103
MeanKES	.122*	.084	.151*	.100	.150*	.105
MeanKWT	-.016	.076	-.056	.090	.000	.094
$R^2 = .115$, Adjusted $R^2 = .013$, Sig. = .456						
$R^2 = .116$, Adjusted $R^2 = .014$, Sig. = .444						
$R^2 = .181$, Adjusted $R^2 = .033$, Sig. = .088*						

From the 1st to the 3rd multiple regression models, the independent variables that are considered statistically significant are MeanPCP and MeanPCT and have beta value of -.203 and .164 respectively in the first model, -.329 and .267 in the second model, and MeanPCP, MeanPCB, MeanPCS, MeanPCC, and MeanPCT have beta value of -.219, .099, .162, -.141, and .312 respectively in the third model. The adjusted R^2 value for the first to the third model was .057, .076, and .081 respectively meaning that less than 10% of the probability of service added to the product was explained by six predictors under product category. The 1st to the 3rd multiple regression models are:

$$\hat{Y}_1 = -.379 - .203PCP - .013PCB + .095PCK + .020PCS + .042PCC + .164PCT. \quad (16)$$

$$\hat{Y}_2 = -.334 - .329PCP + .030PCB + .092PCK + .049PCS - .002PCC + .267PCT. \quad (17)$$

$$\hat{Y}_3 = -.602 - .219PCP + .099PCB - .047PCK + .162PCS - .141PCC + .312PCT. \quad (18)$$

In the 4th to the 6th multiple regression models, the independent variable that is considered statistically significant is MeanSCD and has beta value of .178, .246, and .240 in the fourth, fifth, and sixth model, respectively. The adjusted R^2 value for the first to the third model was .036, .053, and .059 respectively meaning that less than 10% of the probability of service differential the product was explained by three predictors under service category. The 4th to the 6th multiple regression models are:

$$\hat{Y}_4 = -.671 + .178SCD - .005SCI - .026SCW. \quad (19)$$

$$\hat{Y}_5 = -1.068 + .246SCD - .071SCI - .021SCW. \quad (20)$$

$$\hat{Y}_6 = -1.600 + .240SCD - .015SCI - .022SCW. \quad (21)$$

While the 7th to the 9th multiple regression models, the independent variable that is considered statistically significant is MeanKES and has beta value of .122, .151, and .150 in the seventh, eighth, and ninth model, respectively. The adjusted R^2 value for the first to the third model was .013, .014, and .033 respectively meaning that less than 10% of the probability of service differential the product was explained by three predictors under knowledge category. The 7th to 9th multiple regression models are:

$$\hat{Y}_7 = .244 - .055KCH + .122KES - .016KWT. \quad (22)$$

$$\hat{Y}_8 = .242 - .036KCH + .151KES - .056KWT. \quad (23)$$

$$\hat{Y}_9 = -.633 + .021KCH + .150KES + .000KWT. \quad (24)$$

Based on the result of nine multiple regression models, 7 significant factors of the 4-category service levels are MeanPCP, MeanPCB, MeanPCS, MeanPCC, MeanPCT, MeanSCD, and MeanKES. These variables were used for finding the average marginal effects. The Average Marginal Effects (AMEs) was combined and convenient way to compute marginal effect of each dependent variable at every observed value of independent variable and average through the estimation of resulting effects (Leeper, 2017). Findings based upon the estimated equation (11) to (14) can be generated that 7 attributes were significant as presented in Table 6. This data indicates and distinguishes the 4-category service levels.

Data shown in Table 6 is the result of the average marginal effect of 7 significant factors calculated from equation (11) to (14). The 7 significant variables from 4-category service levels illustrated in Table 5

were chemical product only, chemical blending, chemical storage, chemical container recycling, transportation, chemical document, and environmental and safety programs.

Table 6 Logit Average Marginal Effects of Significant Factors of Four Categories Service Levels

No.	Significant Attributes	Logit average marginal effects			
		Product Only	Service Added to the Product	Service Differential the Product	Service is the Product
1	MeanPCP: Chemical Product Only	0.054	0.0003	-0.015	-0.039
2	MeanPCB: Chemical Blending	-0.008	-0.006	-0.006	0.020
3	MeanPCS: Chemical Storage	-0.013	-0.010	-0.009	0.033
4	MeanPCC: Chemical Container Recycling	0.012	0.009	0.008	-0.029
5	MeanPCT: Transportation	0.115	-0.008	-0.069	-0.069
6	MeanSCD: Chemical Documentation	-0.066	-0.008	.039	.035
7	MeanKES: Chemical Environmental and Safety Programs	-0.005	-.024	.015	.014

The marginal effect of the first variable, chemical product only, toward 4-category service levels shows that product only level is the service level that customers who focus on purchasing chemical product only should basically be concentrated compared to the others 3 service levels of service added to the product, service differential the product, and service is the product level. The marginal effect of 0.054 indicates that if there is an increase in the demand of chemical product only by one unit, the service of product only will be more likely to be selected at 5.4%. This research finding was consistent with the study of Eder, Delgado, Kortman, and Studies (2006). In terms of chemical product, traditional business models are focusing on selling chemical product by volume. Chemical suppliers do not have incentive to provide additional services, but they earn money by selling more amount of chemicals.

Secondly, for the chemical blending, service is the product was the preferable service customers want. The marginal effect of 0.02 can be explained that if there is an increase in the demand of chemical blending by one unit, the service level of service is the product will be more likely to be chosen by 2%. On the contrary, the marginal effect of the service level of product only is -0.008, this means the service level of product only will be less likely to be chosen by 0.8% if the demand of chemical blending increases by one unit. Moreover, the service level of service added to the product and service differential the product is also less likely to be selected by 6% if the level of chemical blending demand is increased by one unit because the marginal effect is -0.06. The good evident to support this finding is that chemical suppliers in developed countries, not only world leading companies for example Dow chemical but also local suppliers in North America, Europe, and Japan provide chemical blending service to their customer as bundle solution. They are concerning about safety and setting the highest priority when blending chemicals. With their highly equipped and experiences, this service is provided as custom solution to meet their customer requirement.

The third significant variable is chemical storage. The marginal effect shows that chemical supplier should provide service level of service as the product for customers who has requirement on chemical storage. The marginal effect of .033 indicates that when the demand of chemical storage increases by one unit, the service level of service is the product is more likely to be selected by 3.3%. This is opposite to the other three service levels that have negative marginal effects. From the result of marginal effect in table 6, it can be interpreted that when the demand of chemical storage increases by one unit, the service levels of chemical only, service added to the product, and service differential the product are less likely to be chosen by 1.3%, 1%, and 0.9%, respectively.

The next significant variable is chemical container recycling. The 0.012 marginal effect of product only level indicates that if the customer demand of chemical container recycling raises up one unit, the service level of product only is more likely to be selected at 1.2% of probability. Other two service levels are also having positive effects. Service added to the product and service differential the product are also more likely to be preferred at 0.8% and 0.9% respectively when the demand of chemical container recycling increases by one unit.

Transportation is another significant factor to be considered. The marginal effect of 0.115 can be explained that if the demand of transportation moves up one unit, the service level of product only is more

likely to be chosen by 11.5%. While the other three service levels have negative marginal effect. Service added to the product, service differential the product, and service is the product are less likely to be select by 0.8%, 6.9% and 6.9%, respectively, when the demand of transportation from customer shifts up one unit.

The sixth significant factor is chemical documentation. The positive value of the marginal effect relates to a positive impact of this factor toward service level of service differential the product and service is the product. This means service differential the product and service is the product are more likely to be selected with the probability of 3.9% and 3.5% respectively. This can also be explained that the product only, and service added to the product service levels have negative impact by -6.6% and -0.8% of probability respectively when the demand of chemical documentation increases by one unit. Therefore, customers are more intended to require differential services and service solution when they have more demand of chemical documentation.

The last significance for 4-category service level is chemical environmental and safety programs. The marginal effect sign explains that both service differential the product and service is the product will respond the request of customer on chemical environmental and safety programs. With marginal effect of 0.15 and 0.14, this implies that service differential and service is the product are more likely to be selected with probability of 1.5% and 1.4% respectively if the customer demand of chemical environmental and safety programs rises up one unit.

Discussion and Conclusion

The objectives of this paper were achieved. Firstly, chemical servitization framework was developed and consisted of two parts. The first part of the framework was composed of the three dimensions of customer segments, servitization levels and PSK system, and the second part was the suggestions for Thai chemical suppliers. The research explored the relationship between 4-category service levels and chemical customer requirements. The four service levels were product only, service added to the product, service differential the product and service is the product (Thoben, Eschenbacher, & Jagdev, 2001), and each service level has its own attractiveness of services to be composed of. The questionnaire was distributed to gather data, and descriptive statistics, Analytical Hierarchy Process (AHP), Multiple Linear Regression, and Multinomial Logit Model (MNL) were adopted for data analysis. Secondly, seven substantial factors were identified in order to analyze the service level of customer needs. These significant attributes were chemical product only, chemical blending, chemical storage, chemical container recycling, transportation, chemical documentation, and environmental and safety programs. With different component of services, each service level proposes its own character to meet customer requirement. The marginal effects explain better view which determinant should be focus to improve supplier service offerings for customers.

The research findings highlight the significant attributes of chemical service levels. There will be several guidelines for chemical suppliers to propose service offerings to their customer from this research. For chemical suppliers who propose chemical product only should offer not only selling chemical product in large volume for discount price, but also providing chemical container recycling and transportation services in order to facilitate their customers. Suppliers who have a business model of service differential the product should offer chemical documentation and environmental and safety programs services because their customers want special services rather than just the chemical products only. Suppliers who desire to change their business model from selling tangible products to providing chemical solutions should offer chemical blending, chemical storage, chemical documentation, and environmental and safety programs as bundle services along with chemical products to their customers. However, the study didn't have any suggestions for the suppliers who propose service with the product business model because the results did not show any chemical services that have enough impact to be included in this category. For the future study, researcher might examine some other attributes that have impacts on this service level. Future study may also investigate variables of these 4-category service levels in chemical industry in other industries in Thailand or the chemical industry in other countries. Substantial contributions for this paper are that this study is the first research proposing chemical servitization framework in Thailand and also providing guidance to chemical suppliers for different service level in order to meet their customers' requirements.

Results of the multiple regression models in Table 5 show low values of R-square in every model. The researchers worried about this issue, and were afraid that the low value of R-square would not be acceptable because the models were not well-defined. However, we found a book from Neter, Wasserman,

and Kutner (1985) explaining that R-square is not a measurement of fit, but it measures the explanatory of power. R-square could be low number because the researchers did not expect the model included all the relevant predictors to explain the dependent variables. Eventhough R-square is small, ranging from .012 to .081, but it is different from zero value. This can be indicated that the multiple regression models have statistically significant explanatory power with small effect size. In the social sciences where the models are difficult to specify, low R-square values are often expected.

Acknowledgement

I am thankful to Logistics and Supply Chain Management Interdisciplinary Program, Chulalongkorn University for the advice in this study and also thank you Thai-Nichi Institute of Technology for the financial support.

References

- Buschak, D., & Lay, G. (2014). Chemical Industry: Servitization in Niches. In *Servitization in Industry* (pp. 131-150).
- Chen, D., & Gusmeroli, S. (2015). Framework for Manufacturing Servitization in Virtual Enterprise Environment and Ecosystem. *IFAC-PapersOnLine*, 48, 2244-2249. <https://doi.org/10.1016/j.ifacol.2015.06.422>
- Eder, P., Delgado, L., Kortman, J., & Studies, I. f. P. T. (2006). *Chemical Product Services in the European Union*: Publications Office.
- Goedkoop, M. (1999). *Product Service systems, Ecological and Economic Basics*.
- Kanignant, T., Suthiwartnarueput, K., & Pornchaiwiseskul, P. (2018). Servitization framework for product transition in chemical distribution. In *2018 5th International Conference on Business and Industrial research (ICBIR)* (pp. 323-328). IEEE
- Kortman, J., Theodori, D., Ewijk, V. H., Verspeek, F., & Uitzinger, J. (2006). *Chemical product services in the European Union*.
- Leeper, T. (2017). *Interpreting Regression Results using Average Marginal Effects with R's margins*.
- McFadden, D. (1973). *Conditional Logit Analysis of Qualitative Choice Behaviour* (P. Zarembka ed.). Academic Press New York.
- Meyer, M. H., & Arthur, D. (1999). Product Development for Services. *The Academy of Management Executive* (1993-2005), 13(3), 64-76.
- Neter, J., Wasserman, W., & Kutner, M. H. (1985). *Applied linear statistical models : regression, analysis of variance, and experimental designs*. Homewood, Ill.: R.D. Irwin.
- Oliva, R., & Kallenberg, R. (2003). Managing the transition from products to services. *International Journal of Service Industry Management*, 14(2), 160-172. <https://doi.org/10.1108/09564230310474138>
- OSMEP. (2000). *Small and Medium Enterprise Promotion Act*. Office of Small and Medium Enterprise Promotion, Prime Minister's Office, Thailand
- Peduzzi, P., Concato, J., Kemper, E., Holford, T. R., & Feinstein, A. R. (1996). A simulation study of the number of events per variable in logistic regression analysis. *Journal of Clinical Epidemiology*, 49(12), 1373-1379. [https://doi.org/10.1016/S0895-4356\(96\)00236-3](https://doi.org/10.1016/S0895-4356(96)00236-3)
- Reiskin, E. D., White, A. L., Johnson, J. K., & Votta, T. J. (2000). Servitizing the Chemical Supply Chain. *Journal of Industrial Ecology*, 3(2-3).
- Robinson, T., Clarke-Hill, C., & Clarkson, R. (2002). Differentiation through Service: A Perspective from the Commodity Chemicals Sector. *Service Industries Journal - SERV IND J*, 22, 149-166. <https://doi.org/10.1080/714005092>
- Ryu, J., Rhim, H., Park, K., & Kim, H.-I. (2012). A Framework for Servitization of Manufacturing Companies. *Korea Production and Operative Management society*, 20(4), 145-163.
- Stoughton, M., & Votta, T. (2003). Implementing service-based chemical procurement: Lessons and results. *Journal of Cleaner Production*, 11, 839-849. [https://doi.org/10.1016/S09596526\(02\)00159-2](https://doi.org/10.1016/S09596526(02)00159-2)
- Strandhagen, J. O., Vallandingham, L. R., Fragapane, G., Strandhagen, J. W., Bertnum, A. B. H., & Sharma, N. (2017). Logistics 4.0 and emerging sustainable business models. *Advances in Manufacturing*, 5, 359-369.

- Thoben, K. D., Eschenbacher, J., & Jagdev, H. (2001). Extended products: evolving traditional product concepts. *In proceedings the 7th International Conference on Concurrent Enterprising* (pp. 27-29).
- Toffel, M. (2008). Contracting for Servicizing. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.1090237>
- Tukker, A. (2004). Eight types of product-service system: eight ways to sustainability? experiences from suspronet. *Business Strategy and the Environment*, 13, 246-260. <https://doi.org/10.1002/bse.414>
- Vandermerwe, S., & Rada, J. (1988). Servitization of Business: Adding Value by Adding Service. *European Management Journal*, 6(4), 314-324.

Research Article

Zero Inflated and Zero Truncated Negative Binomial Weighted Garima Distributions for Modeling Count Data

Pornpop Saengthong ^{1*}

¹Institute for Research and Academic Services, Assumption University, Bang Sao Thong, Samuthprakarn, 10570 Thailand

*E-mail: pornpopsng@au.edu

Received: 26/04/2021; Revised: 13/10/2021; Accepted: 15/12/2021

Abstract

In this paper, we introduced models for zero inflated and zero truncated based on the negative binomial-weighted Garima (NB-WG) distribution. The zero inflated negative binomial-weighted Garima (ZINB-WG) distribution is a discrete probability distribution for the excessive zero counts and overdispersion which is a mixture of Bernoulli distribution and negative binomial-weighted Garima (NB-WG) distribution. Meanwhile, a new zero truncated distribution named as the zero truncated negative binomial-weighted Garima (ZTNB-WG) distribution can be used when the response variable is the set of positive integers. Some properties of the two different versions of NB-WG distribution are discussed and the estimation of the parameters is derived by maximum likelihood method (MLE). In addition, the usefulness of the proposed distributions is illustrated by real data sets.

Keywords: zero inflated distribution, zero truncated distribution, negative binomial-weighted Garima distribution, count data, maximum likelihood estimation

Introduction

Modeling of count data is used to analyze non-negative integer outcomes in different fields of study such as insurance, medical, psychology, engineering, etc. The Poisson distribution is a classical tool for count data analysis but its underlying equidispersion condition, therefore, the negative binomial (NB) distribution is a commonly used alternative to the Poisson distribution for describing count data with overdispersion. Count data may either have an excess number of zeros (zero inflation: ZI) or the situation where the value zero do not exist (zero truncation: ZT). For ZI count data, a few common examples of such data are the number of claims for each policyholder, the number of people who tested positive for HIV in each cluster, and the number of day of hospitalizations. Several probabilistic models have been proposed to model count data with excess zeros. The zero inflated Poisson (ZIP) was first introduced by Lambert (1992) for handling excess of zero counts only if there is no overdispersion. However, the excessive zeros situation can cause overdispersion, for this data the zero inflated negative binomial (ZINB), zero-inflated generalized Poisson (ZIGP), and zero inflated double Poisson (ZIDP) are more appropriate (Greene, 1994; Ridout et al., 1998; Yang et al., 2009; Phang & Loh, 2013). In addition, zero inflated mixed NB distributions have been recently defined as alternative models which often have desirable properties for modeling count data with excess

zeros. The models contain a mixture of Bernoulli distribution and mixed NB distribution that include the zero inflated negative binomial-generalized exponential (ZINB-GE) distribution (Aryuyuen et al., 2014), zero inflated negative binomial-Crack (ZINB-CR) distribution (Saengthong et al., 2015), zero inflated negative binomial-Sushila (ZINB-S) distribution (Yamruboon et al., 2017), and zero inflated negative binomial-Erlang (ZINB-EL) distribution (Bodhisuwan et al., 2018).

On the other hand, the ZT model, a special case of the left-truncated model with cutpoint at zero, is applied for analyzing count data except for zero. Some examples of such data include: the number of ever born children, the length of stay in the intensive care unit (ICU) once a patient is admitted, the number of European red mites on apple leaves, and the number of items in a shopper's basket at a supermarket checkout line. For modeling the data as mentioned above, typically, zero truncated Poisson (ZTP) or zero truncated negative binomial (ZTNB) is assumed. Furthermore, many studies have shown that the zero truncated mixed/compound Poisson or negative binomial are more suitable for the ZT count data than ZTP and ZTNB such as zero truncated Poisson–Lindley (ZTPL) distribution (Ghitany et al., 2008), zero truncated Poisson–Garima (ZTPG) distribution (Shanker & Shukla, 2017), zero truncated Poisson–Ishita (ZTPI) distribution (Shukla et al., 2020), zero truncated negative binomial-generalized exponential (ZTNB-GE) distribution (Bodhisuwan, 2016), and zero truncated negative binomial-Erlang (ZTNB-EL) distribution (Bodhisuwan et al., 2017).

The negative binomial-weighted Garima (NB-WG) distribution is a new mixed NB distribution that includes three-parameters introduced by Bodhisuwan & Saengthong (2020). It represents generalization of distribution for count data, including the negative binomial-Garima (NB-G), and negative binomial-size biased Garima (NB-SBG). This new distribution has a heavy tail with positive skewness and may be considered as a preferable alternative for modeling over/underdispersed count data. Therefore, in this work, we proposed the alternative zero inflated and zero truncated models based on the NB-WG distribution for count data analysis.

For the rest of this paper, the constructions of zero inflated and zero truncated negative binomial – weighted Garima distributions, along with their some mathematical properties are provided in Sections 2 and 3. Sections 4 is reserved to the parameter estimation and the usefulness of the proposed distributions is illustrated in Section 5. Finally, the conclusions of the paper are presented in Section 6.

Methods

1) The zero inflated negative binomial-weighted Garima distribution

The probability mass function (pmf) of a zero inflated distribution is defined by

$$g(x) = \begin{cases} \omega + (1-\omega)h(0) & , x = 0 \\ (1-\omega)h(x) & , x = 1, 2, \dots, \end{cases} \quad (1)$$

where x is the count-valued random variable, ω is the zero inflation parameter ($0 < \omega < 1$), and $h(\cdot)$ is the pmf of the parent count model.

Now, we propose the zero inflated negative binomial-weighted Garima (ZINB-WG) distribution, which is a mixture of two distributions between Bernoulli and NB-WG distributions.

Definition 1. Let $X|\lambda$ be a random variable following a NB distribution with parameter $r > 0$ and $p = \exp(-\lambda)$, i.e., $X|\lambda \sim NB(r, p = \exp(-\lambda))$. If λ is distributed as the WG distribution with positive parameters θ and β , denoted by $\lambda \sim WG(\theta, \beta)$, then X is called a random variable of the $NB-WG(r, \theta, \beta)$ distribution.

Theorem 1. If $X \square NB-WG(r, \theta, \beta)$, then the pmf and the factorial moment of order k of the distribution are given as

$$h(x; r, \theta, \beta) = \frac{\theta^\beta}{(\theta + \beta + 1)\Gamma(\beta)} \binom{r+x-1}{x} \times \sum_{j=0}^x \binom{x}{j} (-1)^j \frac{(\theta+1)(\theta+r+j)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+r+j)^{\beta+1}}, \quad (2)$$

and

$$\mu_{[k]}(X) = \frac{\Gamma(r+k)}{\Gamma(r)} \sum_{j=0}^k \binom{k}{j} (-1)^j \frac{\theta^\beta [(\theta+1)(\theta-k+j)\Gamma(\beta) + \theta\Gamma(\beta+1)]}{(\theta+\beta+1)(\theta-k+j)^{\beta+1} \Gamma(\beta)}, \quad (3)$$

where $x = 0, 1, 2, \dots$, for r, θ and $\beta > 0$.

Proof: (see Bodhisuwan & Saengthong (2020)).

Definition 2. Let X be a random variable of NB-WG distribution with pmf $h(x; r, \theta, \beta)$ defined as in (2) and $g(x)$ in (1) is a pmf of the ZI distribution. Then $g(x)$ is a pmf of ZINB-WG distribution with parameters r, θ, β and ω .

Theorem 2. If $X \square ZINB-WG(r, \theta, \beta, \omega)$, the pmf of the distribution is given by

$$g(x; r, \theta, \beta, \omega) = \begin{cases} \omega + (1-\omega) \frac{\theta^\beta}{(\theta + \beta + 1)\Gamma(\beta)} \left(\frac{(\theta+1)(\theta+r)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+r)^{\beta+1}} \right), & x=0 \\ (1-\omega) \frac{\theta^\beta}{(\theta + \beta + 1)\Gamma(\beta)} \binom{r+x-1}{x} \times \sum_{j=0}^x \binom{x}{j} (-1)^j \frac{(\theta+1)(\theta+r+j)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+r+j)^{\beta+1}}, & x=1, 2, \dots, \end{cases} \quad (4)$$

where shape parameters $r, \beta > 0$ and scale parameters $\theta > 0, 0 < \omega < 1$.

Proof: By Definition 2, the pmf of the ZINB-WG distribution is obtained by substituting (2) into (1).

Figure 1 shows some specified parameters r, θ, β , and ω of the ZINB-WG distribution and their graphs. From the graphs it can be observed that the distribution has a positive skew and seems to be flat when θ is decreasing or ω is increasing. In addition, the pmf of the ZINB-GE distribution can take different shapes when values of r or β differ.

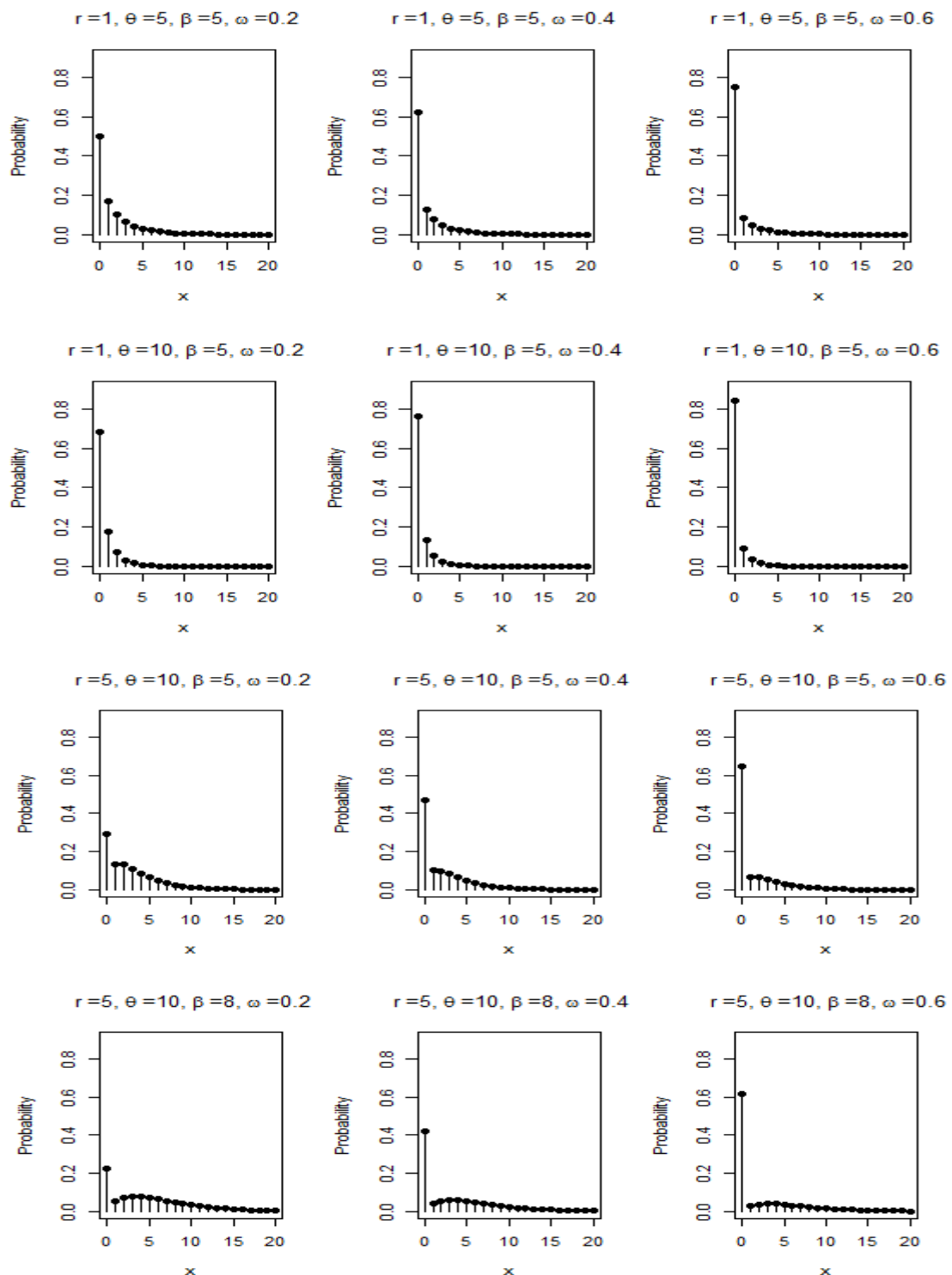


Figure 1 Some pmf plots of the ZINB-WG random variable with specified parameters r, θ, β and ω .

Theorem 3. If $X \square \text{ZINB-WG}(r, \theta, \beta, \omega)$, some basic properties are:

(a) The k^{th} moment of X is given by

$$E(X^k) = (1-\omega) \frac{\theta^\beta}{(\theta+\beta+1)\Gamma(\beta)} \sum_{x=1}^{\infty} x^k \binom{r+x-1}{x} \times \sum_{j=0}^x \binom{x}{j} (-1)^j \frac{(\theta+1)(\theta+r+j)\Gamma(\beta)+\theta\Gamma(\beta+1)}{(\theta+r+j)^{\beta+1}}, k=1,2,\dots, \quad (5)$$

(b) The mean and variance are given as

$$E(X) = (1-\omega)r \left(\frac{\theta^\beta ((\theta+1)(\theta-1)+\theta\beta)}{(\theta+\beta+1)(\theta-1)^{\beta+1}} - 1 \right), \text{ and}$$

$$\text{Var}(X) = E(X^2) - (E(X))^2$$

$$= (1-\omega) \left((r^2+r) \left(\frac{\theta^\beta ((\theta+1)(\theta-2)+\theta\beta)}{(\theta+\beta+1)(\theta-2)^{\beta+1}} \right) - (2r^2+r) \left(\frac{\theta^\beta ((\theta+1)(\theta-1)+\theta\beta)}{(\theta+\beta+1)(\theta-1)^{\beta+1}} \right) + r^2 \right) - \left((1-\omega)r \left(\frac{\theta^\beta ((\theta+1)(\theta-1)+\theta\beta)}{(\theta+\beta+1)(\theta-1)^{\beta+1}} - 1 \right) \right)^2.$$

2) The zero truncated negative binomial-weighted Garima distribution

The pmf of the zero truncated distribution is defined as

$$g(x) = \frac{h(x)}{1-h(0)}, x=1,2,\dots, \quad (6)$$

where $h(x)$ is a pmf of the non-truncated distribution.

Definition 3. Let X be a random variable of NB-WG distribution with pmf $h(x; r, \theta, \beta)$ defined as in (2) and $g(x)$ in (6) is a pmf of the ZT distribution. Then $g(x)$ is a pmf of zero truncated negative binomial – weighted Garima (ZTNB-WG) distribution with parameters r, θ and β .

Theorem 4. If $X \square \text{ZTNB-WG}(r, \theta, \beta)$, the pmf of the distribution is given by

$$g(x; r, \theta, \beta) = \frac{1}{1-h(0)} \left(\frac{\theta^\beta}{(\theta+\beta+1)\Gamma(\beta)} \right) \binom{r+x-1}{x} \times \sum_{j=0}^x \binom{x}{j} (-1)^j \frac{(\theta+1)(\theta+r+j)\Gamma(\beta)+\theta\Gamma(\beta+1)}{(\theta+r+j)^{\beta+1}}, \quad (7)$$

where $x = 1, 2, \dots$, for shape parameters $r, \beta > 0$, scale parameter $\theta > 0$ and

$$h(0) = \frac{\theta^\beta}{(\theta+\beta+1)\Gamma(\beta)} \left(\frac{(\theta+1)(\theta+r)\Gamma(\beta)+\theta\Gamma(\beta+1)}{(\theta+r)^{\beta+1}} \right).$$

Proof: By Definition 3, substitute (2) into (6), then the pmf for the ZTNB-WG distribution can be expressed as (7).

Some characteristics and graphs of the ZTNB-WG distribution with specified values of parameters r, θ and β are provided in Figure 2. It shows that the distribution has a positively

skewed and seems to be flat when θ is decreasing. Likewise, parameters r and β are those whose effect on the shape of this distribution.

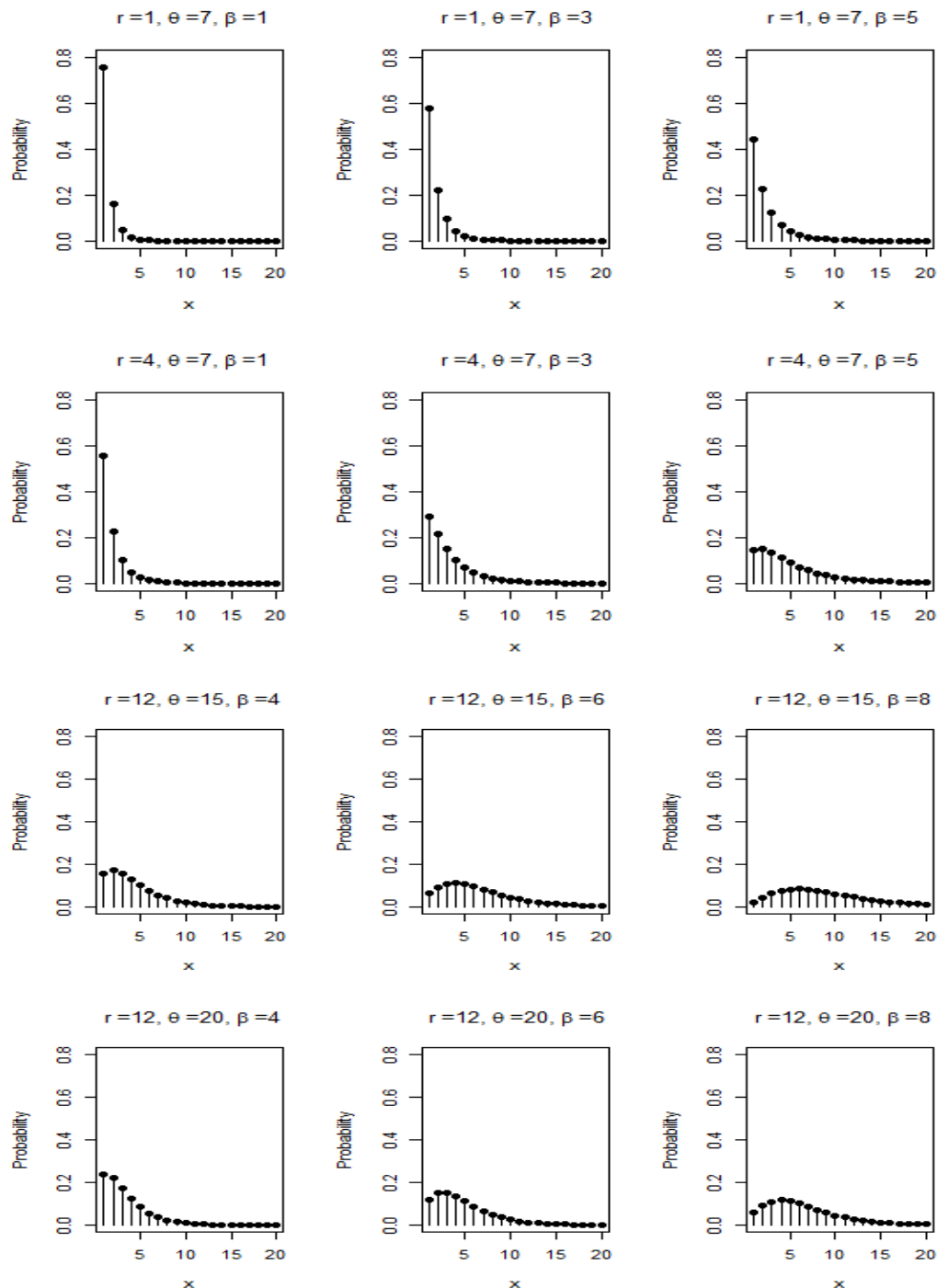


Figure 2 Some pmf plots of the ZTNB-WG random variable with specified parameters r, θ and β .

Theorem 5. If $X \square ZTNB-WG(r, \theta, \beta)$, some mathematical properties are:

(a) The factorial moment of order k of X is given by

$$\mu'_{[k]}(X) = \frac{1}{1-h(0)} \frac{\Gamma(r+k)}{\Gamma(r)} \sum_{j=0}^k \binom{k}{j} (-1)^j \frac{\theta^\beta [(\theta+1)(\theta-k+j)\Gamma(\beta) + \theta\Gamma(\beta+1)]}{(\theta+\beta+1)(\theta-k+j)^{\beta+1} \Gamma(\beta)}, \quad (8)$$

$$\text{where } k = 1, 2, \dots, \text{ for } r, \theta, \beta > 0 \text{ and } h(0) = \frac{\theta^\beta}{(\theta+\beta+1)\Gamma(\beta)} \left(\frac{(\theta+1)(\theta+r)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+r)^{\beta+1}} \right).$$

Proof: If $X|\lambda \square NB(r, p = \exp(-\lambda))$ which is weighted by $\frac{1}{1-h(0)}$ and $\lambda \square WG(\theta, \beta)$, then the factorial moment of order k of X can be derived from

$$\begin{aligned} \mu'_{[k]}(X) &= E_\lambda \left[\frac{1}{1-h(0)} \mu_k(X|\lambda) \right] \\ &= E_\lambda \left[\frac{1}{1-h(0)} \frac{\Gamma(r+k)}{\Gamma(r)} \frac{(1-e^{-\lambda})^k}{e^{-\lambda k}} \right] \\ &= \frac{1}{1-h(0)} \frac{\Gamma(r+k)}{\Gamma(r)} E_\lambda (e^\lambda - 1)^k. \end{aligned}$$

Applying the binomial expansion for the term $(e^\lambda - 1)^k$, we have

$$\begin{aligned} \mu'_{[k]}(X) &= \frac{1}{1-h(0)} \frac{\Gamma(r+k)}{\Gamma(r)} \sum_{j=0}^k \binom{k}{j} (-1)^j E_\lambda (e^{\lambda(k-j)}) \\ &= \frac{1}{1-h(0)} \frac{\Gamma(r+k)}{\Gamma(r)} \sum_{j=0}^k \binom{k}{j} (-1)^j M_\lambda(k-j). \end{aligned}$$

Based on the mgf of the WG distribution (see Bodhisuwan & Saengthong (2020), $\mu'_{[k]}(X)$ can be written as (8).

(b) The mean and variance are given by

$$E(X) = \frac{r \left[\frac{\theta^\beta ((\theta+1)(\theta-1) + \theta\beta)}{(\theta+\beta+1)(\theta-1)^{\beta+1}} \right] - r}{1-h(0)},$$

and

$$\begin{aligned} \text{Var}(X) &= \frac{1}{1-h(0)} \left[(r^2 + r) \left(\frac{\theta^\beta ((\theta+1)(\theta-2) + \theta\beta)}{(\theta+\beta+1)(\theta-2)^{\beta+1}} - \frac{2\theta^\beta ((\theta+1)(\theta-1) + \theta\beta)}{(\theta+\beta+1)(\theta-1)^{\beta+1}} + 1 \right) \right. \\ &\quad \left. + r \left(\frac{\theta^\beta ((\theta+1)(\theta-1) + \theta\beta)}{(\theta+\beta+1)(\theta-1)^{\beta+1}} \right) - r \right] - \frac{\left[r \left(\frac{\theta^\beta ((\theta+1)(\theta-1) + \theta\beta)}{(\theta+\beta+1)(\theta-1)^{\beta+1}} \right) - r \right]^2}{(1-h(0))^2}. \end{aligned}$$

Parameter estimation

In this part, the estimation of parameters for $ZINB-WG(r, \theta, \beta, \omega)$ and $ZTNB-WG(r, \theta, \beta)$ via the maximum likelihood estimation (MLE) are provided.

(a) MLE for the ZINB-WG distribution

In order to find a MLE, we first need to compute the likelihood function which can be obtained by

$$L(r, \theta, \beta, \omega) = \prod_{i=1}^n \left[I_{(x_i=0)} \left(\omega + (1-\omega) \frac{\theta^\beta ((\theta+1)(\theta+r)\Gamma(\beta) + \theta\Gamma(\beta+1))}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)} \right) \right. \\ \left. + I_{(x_i>0)} \left((1-\omega) \frac{\theta^\beta}{(\theta+\beta+1)\Gamma(\beta)} \binom{r+x-1}{x} \right. \right. \\ \left. \left. \times \sum_{j=0}^x \binom{x}{j} (-1)^j \frac{(\theta+1)(\theta+r+j)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+r+j)^{\beta+1}} \right) \right],$$

from which we calculate the log-likelihood function

$$\log L(r, \theta, \beta, \omega) = \ell(r, \theta, \beta, \omega) \\ = \sum_{i=1}^n \left[I_{(x_i=0)} \log \left(\omega + (1-\omega) \frac{\theta^\beta ((\theta+1)(\theta+r)\Gamma(\beta) + \theta\Gamma(\beta+1))}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)} \right) \right. \\ \left. + I_{(x_i>0)} \left(\log(1-\omega) + \beta \log(\theta) - \log(\theta+\beta+1) - \log \Gamma(\beta) + \log(r+x_i-1)! \right. \right. \\ \left. \left. - \log(r-1)! - \log(x_i)! + \log \sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1)(\theta+r+j)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+r+j)^{\beta+1}} \right) \right].$$

It can be verified that the first partial derivatives $\ell(r, \theta, \beta, \omega)$ with respect to r, θ, β and ω we then obtain the following differential equations:

$$\frac{\partial}{\partial r} \ell(r, \theta, \beta, \omega) \\ = \sum_{i=1}^n \left[I_{(x_i=0)} \left(\frac{(1-\omega) \left(\frac{\theta^\beta ((\theta+1)\omega - ((\theta+1)\omega + \theta\beta)(\beta+1))}{(\theta+\beta+1)\omega^{\beta+2}} \right)}{\omega + (1-\omega) \frac{\theta^\beta ((\theta+1)\omega\Gamma(\beta) + \theta\Gamma(\beta+1))}{(\theta+\beta+1)\omega^{\beta+1}\Gamma(\beta)}} \right) \right. \\ \left. + I_{(x_i>0)} \left(\psi(r+x_i) - \psi(r) \right) \right]$$

$$+ \frac{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \left\{ (\theta+1) w_j^{-\beta-1} \Gamma(\beta) - (\beta+1) w_j^{-\beta-2} [(\theta+1) w_j \Gamma(\beta) + \theta \Gamma(\beta+1)] \right\}}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1) w_j \Gamma(\beta) + \theta \Gamma(\beta+1)}{w_j^{\beta+1}}}} \right],$$

where $\psi(k) = \frac{\Gamma'(k)}{\Gamma(k)}$ is the digamma function.

$$\begin{aligned} & \frac{\partial}{\partial \theta} \ell(r, \theta, \beta, \omega) \\ &= \sum_{i=1}^n \left[I_{(x_i=0)} \left(\frac{(1-\omega) \left(\frac{\theta^\beta w(\beta+1) \Gamma(\beta) + \theta^\beta (\theta+1) \Gamma(\beta) + \theta^{\beta-1} w \Gamma(\beta+1) + \theta^\beta \Gamma(\beta+2)}{(\theta+\beta+1) w^{\beta+1} \Gamma(\beta)} \right)}{\omega + (1-\omega) \frac{\theta^\beta ((\theta+1) w \Gamma(\beta) + \theta \Gamma(\beta+1))}{(\theta+\beta+1) w^{\beta+1} \Gamma(\beta)}} \right. \right. \\ & \quad \left. \left. - \frac{\left(\frac{(\theta^\beta \Gamma(\beta) ((\theta+1) w + \theta \beta)) (w^\beta \Gamma(\beta) ((\theta+\beta+1)(\beta+1) + w))}{((\theta+\beta+1) w^{\beta+1} \Gamma(\beta))^2} \right)}{\omega + (1-\omega) \frac{\theta^\beta ((\theta+1) w \Gamma(\beta) + \theta \Gamma(\beta+1))}{(\theta+\beta+1) w^{\beta+1} \Gamma(\beta)}} \right) \right] \\ & \quad + I_{(x_i>0)} \left(\frac{\beta}{\theta} - \frac{1}{\theta+\beta+1} + \frac{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \left\{ w_j^{-\beta-1} [w_j \Gamma(\beta) + (\theta+1) \Gamma(\beta) + \Gamma(\beta+1)] \right\}}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1) w_j \Gamma(\beta) + \theta \Gamma(\beta+1)}{w_j^{\beta+1}}} \right. \\ & \quad \left. - \frac{(\beta+1) w_j^{-\beta-2} [(\theta+1) w_j \Gamma(\beta) + \theta \Gamma(\beta+1)]}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1) w_j \Gamma(\beta) + \theta \Gamma(\beta+1)}{w_j^{\beta+1}}} \right), \end{aligned}$$

$$\begin{aligned} & \frac{\partial}{\partial \beta} \ell(r, \theta, \beta, \omega) \\ &= \sum_{i=1}^n \left[I_{(x_i=0)} \left(\frac{(1-\omega) \left(\frac{\theta^\beta ((\theta+1) w (\psi(\beta) + \log(\theta)) + \theta \beta (\psi(\beta+1) + \log(\theta)))}{(\theta+\beta+1) w^{\beta+1}} \right)}{\omega + (1-\omega) \frac{\theta^\beta ((\theta+1) w \Gamma(\beta) + \theta \Gamma(\beta+1))}{(\theta+\beta+1) w^{\beta+1} \Gamma(\beta)}} \right) \right] \end{aligned}$$

$$\begin{aligned}
 & - \left(\frac{\theta^\beta ((\theta+1)w\Gamma(\beta) + \theta\Gamma(\beta+1))((\theta+\beta+1)\psi(\beta) + (\theta+\beta+1)\log w + 1)}{(\theta+\beta+1)^2 w^{\beta+1}\Gamma(\beta)} \right) \Bigg) \\
 & \quad \omega + (1-\omega) \frac{\theta^\beta ((\theta+1)w\Gamma(\beta) + \theta\Gamma(\beta+1))}{(\theta+\beta+1)w^{\beta+1}\Gamma(\beta)} \Bigg) \\
 & + I_{(x_i > 0)} \left(\log(\theta) - \frac{1}{\theta+\beta+1} - \psi(\beta) \right. \\
 & \quad + \left(\frac{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \{ w_j^{-\beta-1} [(\theta+1)w_j\Gamma(\beta)\psi(\beta) + \theta\Gamma(\beta+1)\psi(\beta+1)] \}}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1)w_j\Gamma(\beta) + \theta\Gamma(\beta+1)}{w_j^{\beta+1}}} \right. \\
 & \quad \left. \left. - \frac{w_j^{-\beta-1} [(\theta+1)w_j\Gamma(\beta) + \theta\Gamma(\beta+1)] \log w_j}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1)w_j\Gamma(\beta) + \theta\Gamma(\beta+1)}{w_j^{\beta+1}}} \right) \right],
 \end{aligned}$$

and

$$\begin{aligned}
 & \frac{\partial}{\partial \omega} \ell(r, \theta, \beta, \omega) \\
 & = \sum_{i=1}^n I_{(x_i=0)} \left[\frac{1 - \left(\frac{\theta^\beta ((\theta+1)(\theta+r)\Gamma(\beta) + \theta\Gamma(\beta+1))}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)} \right)}{\omega + (1-\omega) \frac{\theta^\beta ((\theta+1)(\theta+r)\Gamma(\beta) + \theta\Gamma(\beta+1))}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)}} \right] - I_{(x_i > 0)} \left(\frac{1}{1-\omega} \right),
 \end{aligned}$$

where $w = r + \theta$ and $w_j = r + \theta + j$.

These differential equations cannot be solved analytically, as they need to rely on Newton-Raphson method to approximate MLE. The ML estimates of $\hat{r}, \hat{\theta}, \hat{\beta}$ and $\hat{\omega}$ can be obtained by using numerical optimization with the nonlinear minimization (nlm) function in R package, namely stats (R Core Team, 2020). Therefore, we use the negative log-likelihood function as minimizing that is equivalent to maximizing the log-likelihood function.

(b) MLE for the ZTNB-WG distribution

The likelihood function of the ZTNB-WG(r, θ, β) distribution is given by

$$\begin{aligned}
 L(r, \theta, \beta) & = \frac{\theta^\beta}{(1-h(0))(\theta+\beta+1)\Gamma(\beta)} \prod_{i=1}^n \binom{r+x_i-1}{x_i} \\
 & \quad \times \sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1)(\theta+r+j)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+r+j)^{\beta+1}},
 \end{aligned}$$

where $x = 1, 2, \dots$, for $r, \theta, \beta > 0$ and $h(0) = \frac{\theta^\beta}{(\theta + \beta + 1)\Gamma(\beta)} \left(\frac{(\theta + 1)(\theta + r)\Gamma(\beta) + \theta\Gamma(\beta + 1)}{(\theta + r)^{\beta + 1}} \right)$.

The log-likelihood function can be calculated as

$$\begin{aligned} \log L(r, \theta, \beta) &= \ell(r, \theta, \beta) \\ &= \sum_{i=1}^n [\log(r + x_i - 1)! - \log x_i! - \log(r - 1)!] + n\beta \log(\theta) - n \log(\theta + \beta + 1) \\ &\quad - n \log \Gamma(\beta) + \sum_{i=1}^n \left[\log \sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta + 1)(\theta + r + j)\Gamma(\beta) + \theta\Gamma(\beta + 1)}{(\theta + r + j)^{\beta + 1}} \right] \\ &\quad - n \log \left(1 - \frac{\theta^\beta (\theta + 1)(\theta + r)\Gamma(\beta) + \theta\Gamma(\beta + 1)}{(\theta + \beta + 1)(\theta + r)^{\beta + 1} \Gamma(\beta)} \right). \end{aligned}$$

The first partial derivative of $\ell(r, \theta, \beta)$ with respect to r, θ and β are given as follows:

$$\begin{aligned} \frac{\partial}{\partial r} \ell(r, \theta, \beta) &= \sum_{i=1}^n \psi(r + x_i) - n\psi(r) \\ &\quad + \sum_{i=1}^n \left(\frac{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \{ (\theta + 1)w_j^{-\beta - 1} \Gamma(\beta) - (\beta + 1)w_j^{-\beta - 2} [(\theta + 1)w_j \Gamma(\beta) + \theta\Gamma(\beta + 1)] \}}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta + 1)w_j \Gamma(\beta) + \theta\Gamma(\beta + 1)}{w_j^{\beta + 1}}} \right) \\ &\quad + n \left(\frac{\frac{\theta^\beta ((\theta + 1)w - ((\theta + 1)w + \theta\beta)(\beta + 1))}{(\theta + \beta + 1)w^{\beta + 2}}}{1 - \frac{\theta^\beta (\theta + 1)(\theta + r)\Gamma(\beta) + \theta\Gamma(\beta + 1)}{(\theta + \beta + 1)(\theta + r)^{\beta + 1} \Gamma(\beta)}} \right), \end{aligned}$$

where $\psi(k) = \frac{\Gamma'(k)}{\Gamma(k)}$ is the digamma function.

$$\begin{aligned} \frac{\partial}{\partial \theta} \ell(r, \theta, \beta) &= \frac{n\beta}{\theta} - \frac{n}{\theta + \beta + 1} \\ &\quad + \sum_{i=1}^n \left(\frac{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \{ w_j^{-\beta - 1} [w_j \Gamma(\beta) + (\theta + 1)\Gamma(\beta) + \Gamma(\beta + 1)] \}}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta + 1)w_j \Gamma(\beta) + \theta\Gamma(\beta + 1)}{w_j^{\beta + 1}}} \right) \end{aligned}$$

$$\left. \begin{aligned} & - \frac{(\beta+1)w_j^{-\beta-2}[(\theta+1)w_j\Gamma(\beta)+\theta\Gamma(\beta+1)]}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1)w_j\Gamma(\beta)+\theta\Gamma(\beta+1)}{w_j^{\beta+1}}} \right\} \\ & + n \left[\frac{\left(\frac{\theta^\beta w(\beta+1)\Gamma(\beta)+\theta^\beta(\theta+1)\Gamma(\beta)+\theta^{\beta-1}w\Gamma(\beta+1)+\theta^\beta\Gamma(\beta+2)}{(\theta+\beta+1)w^{\beta+1}\Gamma(\beta)} \right)}{1 - \frac{\theta^\beta(\theta+1)(\theta+r)\Gamma(\beta)+\theta\Gamma(\beta+1)}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)}} \right. \\ & \left. - \frac{\left(\frac{(\theta^\beta\Gamma(\beta)((\theta+1)w+\theta\beta))(w^\beta\Gamma(\beta)((\theta+\beta+1)(\beta+1)+w))}{((\theta+\beta+1)w^{\beta+1}\Gamma(\beta))^2} \right)}{1 - \frac{\theta^\beta(\theta+1)(\theta+r)\Gamma(\beta)+\theta\Gamma(\beta+1)}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)}} \right], \end{aligned}$$

and

$$\begin{aligned} & \frac{\partial}{\partial \beta} \ell(r, \theta, \beta) \\ & = n \log(\theta) - \frac{n}{\theta+\beta+1} - n\psi(\beta) \\ & + \sum_{i=1}^n \left[\frac{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \left\{ w_j^{-\beta-1} [(\theta+1)w_j\Gamma(\beta)\psi(\beta)+\theta\Gamma(\beta+1)\psi(\beta+1)] \right\}}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1)w_j\Gamma(\beta)+\theta\Gamma(\beta+1)}{w_j^{\beta+1}}} \right. \\ & \quad \left. - \frac{w_j^{-\beta-1} [(\theta+1)w_j\Gamma(\beta)+\theta\Gamma(\beta+1)] \log w_j}{\sum_{j=0}^{x_i} \binom{x_i}{j} (-1)^j \frac{(\theta+1)w_j\Gamma(\beta)+\theta\Gamma(\beta+1)}{w_j^{\beta+1}}} \right] \\ & + n \left[\frac{\left(\frac{\theta^\beta ((\theta+1)w(\psi(\beta)+\log(\theta))+\theta\beta(\psi(\beta+1)+\log(\theta)))}{(\theta+\beta+1)w^{\beta+1}} \right)}{1 - \frac{\theta^\beta(\theta+1)(\theta+r)\Gamma(\beta)+\theta\Gamma(\beta+1)}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)}} \right] \end{aligned}$$

$$-\left(\frac{\theta^\beta ((\theta+1)w\Gamma(\beta) + \theta\Gamma(\beta+1))((\theta+\beta+1)\psi(\beta) + (\theta+\beta+1)\log w + 1)}{(\theta+\beta+1)^2 w^{\beta+1}\Gamma(\beta)} \right) \Bigg/ \left(1 - \frac{\theta^\beta (\theta+1)(\theta+r)\Gamma(\beta) + \theta\Gamma(\beta+1)}{(\theta+\beta+1)(\theta+r)^{\beta+1}\Gamma(\beta)} \right),$$

where $w = r + \theta$ and $w_j = r + \theta + j$.

The ML estimators of the ZTNB-WG distribution are obtained by solving the above partial derivative equations by using the nlm function in R package.

Results and discussion

In this section, we discuss applications of the ZINB-WG and ZTNB-WG distributions to four real data sets. The first and second data sets are considered to fit with the ZIP, ZINB, and ZINB-WG distributions, while the third and fourth data sets are used to illustrate the flexibility of the ZTP, ZTNB, and ZTNB-WG distributions.

The first data set is studied on the number of major derogatory reports (MDRs) in the credit history of individual credit card applicants (see Greene (1994)). As shown in Table 1, the percentage of zero in the number of MDRs is 80.36 and the variance-to-mean ratio or index of dispersion (ID) is 3.298, indicating that there is overdispersion ($ID > 1$) with a mean = 0.456, variance = 1.810 and a high percentage of zeros. Additionally, the second data set is due to Klugman et al. (2019) regarding the number of runs scored in each half-inning of World Series baseball games played from 1947 through 1960, which is displayed in Table 2. This data has 72.97% of zeros, mean = 0.484, and variance = 1.075. Likewise, it has overdispersion with an index of dispersion equal to 2.220.

The third data set contains of the number of households having at least one migrant from households according to the number of migrants (see Singh & Yadava (1981)). Based on mean = 1.547, variance = 0.819 and $ID = 0.529$, it can be indicated that this dataset presents an underdispersed data ($ID < 1$). The fitting results are shown in Table 3. Lastly, we consider the fourth data set in Table 4 which reports the number of mothers of the rural area with at least two live births by the number of infant and child deaths (see Shanker (2016)), in which the mean, variance, and ID are 1.403, 0.598, and 0.426 respectively (underdispersion).

According to the results of the distribution fitting, the best distribution is the distribution that corresponds to lower values of negative log-likelihood and a higher p-value of chi-square goodness of fit test. From the results in Table 1-4, we can conclude that the ZINB-WG and ZTNB-WG distributions are competing well against the others and gives a satisfactory fit.

Table 1 observed and expected frequencies for the consumer credit behavior data.

No. of MDRs	Observed	Fitting distributions		
		ZIP	ZINB	ZINB-WG
0	1060	1059.99	1060.08	1060.08
1	137	80.32	94.17	132.96
2	50	80.90	72.86	54.34
3	24	54.31	45.17	27.41
4	17	27.35	24.52	15.38
5	11	11.02	12.18	9.25
6	5	} 5.11	} 10.02	5.85
7	6			} 13.72
8-14	9			
Estimated parameters		$\hat{\lambda} = 2.0142$	$\hat{r} = 3.9566$	$\hat{r} = 8.0652$
		$\hat{\omega} = 0.7734$	$\hat{p} = 0.6878$	$\hat{\theta} = 5.0944$
			$\hat{\omega} = 0.7459$	$\hat{\beta} = 0.2664$
				$\hat{\omega} = 0.1932$
Negative Log-Likelihood		1140.513	1091.034	1055.526
Chi-squares		116.076	48.936	1.635
Degree of freedom		4	3	3
p-value		<0.0001	<0.0001	0.6514

Table 2 observed and expected frequencies for the runs scored data.

No. of Runs	Observed	Fitting distributions		
		ZIP	ZINB	ZINB-WG
0	1023	793.88	1023.01	1022.69
1	222	401.81	196.05	221.71
2	87	156.41	107.72	86.24
3	32	40.59	47.41	37.15
4	18	} 9.31	18.28	17.07
5	11		} 9.54	8.23
6	6			} 8.91
7+	3			

Estimated parameters	$\hat{\lambda} = 1.3064$ $\hat{\omega} = 0.6293$	$\hat{r} = 3.9622$ $\hat{p} = 0.7785$ $\hat{\omega} = 0.5703$	$\hat{r} = 12.4731$ $\hat{\theta} = 15.3552$ $\hat{\beta} = 0.6626$ $\hat{\omega} = 0.1981$
Negative Log-Likelihood	1311.282	1293.623	1283.016
Chi-squares	267.630	23.915	1.705
Degree of freedom	2	2	2
p-value	0	<0.0001	0.4264

Table 3 observed and expected frequencies for the migration data.

No. of Migrant	Observed	Fitting distributions		
		ZTP	ZTNB	ZTNB-WG
1	375	353.98	363.88	376.19
2	143	167.70	155.13	141.38
3	49	52.96	51.46	48.24
4	17	15.36	19.53	16.04
5	2			8.15
6	2			
7	1			
8	1			
Estimated parameters		$\hat{\lambda} = 0.9475$	$\hat{r} = 4.9786$ $\hat{p} = 0.8574$	$\hat{r} = 4.5082$ $\hat{\theta} = 32.8190$ $\hat{\beta} = 4.4986$
Negative Log-Likelihood		601.645	595.277	592.732
Chi-squares		8.987	2.022	0.659
Degree of freedom		2	1	1
p-value		0.0112	0.1534	0.4168

Table 4 observed and expected frequencies for the mortality data.

No. of Infant and Child Deaths	Observed	Fitting distributions		
		ZTP	ZTNB	ZTNB-WG
1	745	708.87	722.10	746.64
2	212	255.11	236.22	204.63
3	50	61.21	61.83	59.87
4	21	} 12.81	} 17.85	18.27
5	7			} 8.58
6	3			
Estimated parameters		$\hat{\lambda} = 0.7198$	$\hat{r} = 3.9940$ $\hat{p} = 0.8690$	$\hat{r} = 0.7688$ $\hat{\theta} = 78.2835$ $\hat{\beta} = 28.4131$
Negative Log-Likelihood		889.844	878.842	872.491
Chi-squares		37.003	15.169	2.537
Degree of freedom		2	1	1
p-value		<0.0001	<0.0001	0.1112

Conclusion

This paper presents zero modified versions of the NB-WG distribution, which are ZINB-WG and ZTNB-WG distributions. Some mathematical properties involving mean, variance, and k^{th} factorial moment of the proposed distributions have been derived. The parameter estimation via the maximum likelihood method is also implemented. Finally, the efficiency of the ZINB-WG and ZTNB-WG distributions compared to other traditional probability distributions are illustrated by real data sets.

Acknowledgement

The authors are thankful to the editor and all referees for their valuable suggestions.

References

- Aryuyuen, S., Bodhisuwan, W., & Supapakorn, T. (2014). Zero inflated negative binomial-generalized exponential distribution and its applications. *Songklanakarin Journal of Science & Technology*, 36(4), 483-491.
- Bodhisuwan, R. (2016). A mathematical model of count data analysis based on ZTNB-GE distribution. *The 12th International Conference on Mathematics, Statistics, and Their Applications (ICMSA)* (pp.113-116). Banda Aceh, Indonesia.
- Bodhisuwan, W., Pudprommarat, C., Bodhisuwan, R., & Saothayanun, L. (2017). Zero-truncated negative binomial - Erlang distribution. *The 13th IMT-GT International Conference on Mathematics, Statistics and Their Applications (ICMSA)*. Malaysia.

- Bodhisuwan, W., Samutwachirawong, S., & Payakkapong, P. (2018). The zero-inflated negative binomial-Erlang distribution: An application to highly pathogenic avian influenza H5N1 in Thailand. *Songklanakarin Journal of Science & Technology*, 40(6), 1428-1436.
- Bodhisuwan, W., & Saengthong, P. (2020). The Negative Binomial – Weighted Garima Distribution: Model, Properties and Applications. *Pakistan Journal of Statistics and Operation Research*, 16(1), 1-10. <https://doi.org/10.18187/pjsor.v16i1.3013>
- Ghitany, M. E., Al-Mutairi, D. K., & Nadarajah, S. (2008). Zero-truncated Poisson-Lindley distribution and its Applications. *Mathematics and Computers in Simulation*, 79(3), 279–287.
- Greene, W. (1994). *Accounting for excess zeros and sample selection in Poisson and negative binomial regression models*. Working Paper EC-94-10, New York University, New York, U.S.A.
- Klugman, S. A., Panjer, H. H., & Willmot, G. E. (2019). *Loss Models: From Data to Decision* (5th ed.). New York, USA: Wiley.
- Lambert, D. (1992). Zero-inflated Poisson regression, with an application to defects in manufacturing. *Technometrics*, 34(1), 1-14.
- Phang, Y. N., & Loh, E. F. (2013). Zero Inflated Models for Overdispersed Count Data. *International Journal of Health and Medical Engineering*, 7(8), 1331-1333.
- R Core Team. (2020). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria.
- Ridout, M. S., Demetrio, C. G. B., & Hinde J. P. (1998). Models for counts data with many zeros. *The 19th International Biometric Conference* (pp. 179-192). Cape Town.
- Saengthong, P., Bodhisuwan, W., & Thongteeraparp, A. (2015). The zero inflated negative binomial-Crack distribution: some properties and parameter estimation. *Songklanakarin Journal of Science & Technology*, 37(6), 701-711.
- Shanker, R., & Fesshaye, H. (2016). On zero-truncation of poisson, poisson-lindley and poisson-sujatha distributions and their applications. *Biometrics & Biostatistics International Journal*, 3(4), 125–135.
- Shanker, R., & Shukla, K. K. (2017). Zero-Truncated Poisson-Garima Distribution and its Applications. *Biostatistics and Biometrics Open Access Journal*, 3(1), 555605. <https://doi.org/10.19080/BBOAJ.2017.03.555605>
- Shukla, K. K., Shanker, R., & Tiwari, M. (2020). Zero-truncated Poisson-Ishita Distribution and Its Applications. *Journal of scientific research*, 64(2), 287-294.
- Singh, S. N., & Yadava, K. N. (1981). Trends in rural out-migration at household level. *Rural Demography*, 8(1), 53–61.
- Yamrubboon, D., Thongteeraparp, A., Bodhisuwan, W., & Jampachaisri, K. (2017). Zero inflated negative binomial-Sushila distribution and its application. *The 13th IMT-GT International Conference on Mathematics, Statistics and Their Applications (ICMSA)*. Universiti Utara MalaysiaKedah; Malaysia.
- Yang, Z., Hardin, J. W., & Addy C. L. (2009). Testing overdispersion in the zero-inflated Poisson model. *Journal of Statistical Planning and Inference*, 139(9), 3340-3353. <https://doi.org/10.1016/j.jspi.2009.03.016>

Research Article

Finite Dimensional Simple Poisson Modules of a Poisson Algebras A

Nagornchat Chansuriya^{1,*} and Supaporn Fongchanta²

¹ Faculty of Science, Energy and Environment, King Mongkut's University of Technology North Bangkok (Rayong Campus), Rayong 21120, Thailand

² Department of Mathematics and Statistics, Faculty of Science and Technology, Chiang Mai Rajabhat University, Chiang Mai 50300, Thailand

*E-mail: nagornchat.c@sciee.kmutnb.ac.th

Received: 11/06/2021; Revised: 18/09/2021; Accepted: 24/03/2022

Abstract

This study focus on a Poisson algebra $A = \mathbb{K}[x, y, z]$ with a Poisson bracket $\{x, y\} = yx + z$, $\{y, z\} = zy + x$, $\{z, x\} = xz + y$. There are precisely five Poisson maximal ideals as following.

$$J_1 = xA + yA + zA,$$

$$J_2 = (x+1)A + (y+1)A + (z+1)A,$$

$$J_3 = (x+1)A + (y-1)A + (z-1)A,$$

$$J_4 = (x-1)A + (y+1)A + (z-1)A, \text{ and}$$

$$J_5 = (x-1)A + (y-1)A + (z+1)A.$$

In this research, we determine the finite dimensional simple Poisson modules annihilated by J_i ; $i = 1, 2, \dots, 5$.

Keywords: Poisson algebra, Poisson module, simple Poisson module, finite-dimensional simple Poisson module

Introduction

The Poisson algebra is an algebra with the brackets, called Poisson bracket, introduced in 1809 by Joseph-Louis Lagrange and his student, Siméon-Denis de Poisson, as an algorithm useful to produce solutions of motion. It is the important tool to study in Mathematics and Physics. There are many mathematicians study a Poisson algebra as a Poisson geometry, quantum group and deformation of commutative associative algebras. In Physics, a Poisson algebra is crucial for study the deformation quantization, Hamiltonian mechanic and topological field theories.

There are many authors who interested in a Poisson algebra. The study on a Poisson algebra as an operator algebras are widely study. Sasom (2006) used the direct method

constructed the Poisson algebra from the algebra T and the quantized enveloping algebra $U_q(sl_2)$. She found the Poisson maximal ideals of each Poisson algebra and determined the Poisson modules annihilated by each Poisson maximal ideal. After that, Jordan (2010) applied the method in Erdmann and Wildon (2006) (low-dimensional Lie algebra) to the Poisson algebra constructed from the quantized enveloping algebra $U_q(sl_2)$, a Poisson algebra B , presented in Sasom (2006). He showed that there was a unique d -dimensional simple Poisson module over B annihilated by each Poisson maximal ideals for each $d \geq 1$. Afterwards, Sasom (2012) studied the Poisson algebra $A = \square[x, y, z]$ with a Poisson bracket $\{x, y\} = yx + z$, $\{y, z\} = zy + x$ and $\{z, x\} = xz + y$. She found that there were precisely five Poisson maximal ideals as the following.

$$\begin{aligned} J_1 &= xA + yA + zA, \\ J_2 &= (x+1)A + (y+1)A + (z+1)A, \\ J_3 &= (x+1)A + (y-1)A + (z-1)A, \\ J_4 &= (x-1)A + (y+1)A + (z-1)A, \\ J_5 &= (x-1)A + (y-1)A + (z+1)A. \end{aligned}$$

She used the direct method constructed the Poisson modules which annihilated by each $J_i; i=1, \dots, 5$. Later, Changtong and Sasom (2018) studied the Poisson algebra $A = \square[x, y, z]$ with a Poisson bracket $\{x, y\} = yx + x + y + z$, $\{y, z\} = zy + x + y + z$ and $\{z, x\} = xz + x + y + z$. They found that there were only two Poisson maximal ideals $J_1 = xA + yA + zA$ and $J_2 = (x+3)A + (y+3)A + (z+3)A$. They used the same method in Jordan (2010) showed that every finite-dimensional simple Poisson module over A annihilated by J_1 was one-dimensional and for each $d \geq 1$, there was a unique d -dimensional simple Poisson module over B annihilated by J_2 .

In this present paper, we focus on the Poisson algebra $A = \square[x, y, z]$ constructed by Sasom (2012) with the Poisson maximal ideals as above. We use the same method in Jordan (2010) and Changtong and Sasom (2018) determine the finite dimensional simple Poisson modules annihilated by each Poisson maximal ideal. This Poisson module is very important to the study in the Poisson geometry and the quantum mechanic.

Materials and methods

In this section, it contains some of the materials that will be used throughout this research. The main topics are Lie algebras, Low-dimensional Lie algebra, Poisson algebra and Poisson module.

Lie algebras

Definition 1. [Henderson (2012)] Let L be a vector space over a field F . A *Lie algebra* is a vector space L with a map $[-, -]: L \times L \rightarrow L$ satisfying:

(i) the map $[-, -]$ is bilinear, i.e., for all $a, b \in \square$ and $x, y, z \in L$

$$[ax + by, z] = a[x, z] + b[y, z],$$

$$[x, ay + bz] = a[x, y] + b[x, z];$$

(ii) the map $[-, -]$ is skew-symmetric, i.e. $[y, x] = -[x, y]$ for all $x, y \in L$;

(iii) the map $[-, -]$ satisfies the Jacobi identity, i.e.

$$[x, [y, z]] + [y, [z, x]] + [z, [x, y]] = 0 \text{ for all } x, y, z \in L.$$

So, the map $[-, -]$ is called *Lie bracket* and (ii) implies $[x, x] = 0$ for all $x \in L$.

Next, we will introduce the special Lie algebras that useful in this study. Firstly, we introduce the *special linear Lie algebra* $sl(n, \square)$ of $n \times n$ matrices with trace 0. It has dimension $n^2 - 1$. The most important is a Lie algebra $sl(2, \square)$ of 2×2 matrices with trace 0. It is a three-dimensional Lie algebra with the standard basis

$$e = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, \quad h = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad f = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}.$$

The Lie bracket of $sl(2, \square)$ is given by $[e, f] = h$, $[h, e] = 2e$, $[h, f] = -2f$.

Theorem 2. [Henderson (2012)] For all $n \geq 2$, $sl(n, \square)$ is simple.

For more details of $sl(2, \square)$ see in Henderson (2012).

Another one is the *special unitary Lie algebra* $su(2, \square)$ of 2×2 skew-Hermitian matrices with trace 0. It has the basis

$$x = \frac{1}{2} \begin{pmatrix} 0 & i \\ i & 0 \end{pmatrix}, \quad y = \frac{1}{2} \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}, \quad z = \frac{1}{2} \begin{pmatrix} i & 0 \\ 0 & -i \end{pmatrix}.$$

The Lie bracket of $su(2, \square)$ is given by $[x, y] = z$, $[y, z] = x$, $[z, x] = y$. So, $su(2, \square)$ has dimension 3 which implies that $su(2, \square) \cong sl(2, \square)$.

For more details of $su(2, \square)$ see in Pfeifer (2003) and Kirillov (2008).

Definition 3. Let L be a Lie algebra. A subalgebra H of L is a *Lie subalgebra* if $[x, y] \in H$, for all $x, y \in H$.

Definition 4. An *ideal* of a Lie algebra L is a subspace I of L such that $[x, y] \in I$ for all $x \in L$, $y \in I$.

Let I, J be ideals of the Lie algebra L . The product $[I, J]$ is the span of the commutators of elements of I, J , that is $[I, J] := \text{Span}\{[x, y] \mid x \in I, y \in J\}$. One of the tools for investigating the results is the derived algebra. It is defined as following.

Definition 5. Let I, J be ideals of a Lie algebra L . If $I = J = L$, then the ideal $[L, L]$ is called the *derived algebra* of L .

Low-dimensional Lie algebra

The low-dimension is the basic way to find how many non-isomorphic Lie algebras in order to classify them. The reason to study on the low-dimensional Lie algebras is that there often appear as subalgebras of the larger Lie algebras. We will look at the Lie algebras of dimensions 1, 2 and 3.

It is easily to see that every 1 dimensional Lie algebra is abelian. For dimension 2, we can study by using the following theorem.

Theorem 6. [Erdmann & Wildon (2006)] Let F be any field. Up to isomorphism, there is a unique two-dimensional non-abelian Lie algebra over F . This Lie algebra has a basis $\{x, y\}$ such that its Lie bracket is described by $[x, y] = x$. The centre of this Lie algebra is 0.

For 3-dimensional Lie algebras, we focus for 2 cases. Firstly, the 3 dimensional Lie algebra L with its derived algebra $[L, L]$ has dimension 1 and $[L, L] \subseteq Z(L)$, the center of L , appears uniquely and it has a basis $\{f, g, h\}$ where $[f, g] = h \in Z(L)$. This Lie algebra is known as the *Heisenberg algebras*. For another one, suppose that L is a complex Lie algebra of dimension 3 such that $[L, L] = L$. Up to isomorphism, $sl(2, \mathbb{C})$ is the unique 3-dimensional Lie algebras L with $[L, L] = L$.

For other cases and more details of low-dimensional Lie algebras, one can study in Erdmann & Wildon (2006).

Poisson algebra and Poisson module

Definition 7. A *Poisson algebra* A is a commutative algebra over a field F together with a bilinear map $\{-, -\}: A \times A \rightarrow A$ such that $(A, \{-, -\})$ is a Lie algebra and satisfies a Leibniz identity

$$\{xy, z\} = x\{y, z\} + \{x, z\}y, \text{ for all } x, y, z \in A.$$

We call $\{-, -\}$ a *Poisson bracket* on A . A subalgebra B of A is a *Poisson subalgebra* if $\{b, c\} \in B$ for all $b, c \in B$.

Definition 8. Let A be a Poisson algebra. An ideal I of A is a *Poisson ideal* if $\{i, a\} \in I$ for all $i \in I, a \in A$.

Moreover, the factor A/I is also a Poisson algebra with a Poisson bracket $\{a+I, b+I\} = \{a, b\} + I$ for all $a, b \in A$.

Definition 9. Let A be a Poisson algebra. A Poisson ideal I of A , which $I \neq A$, is said to be a *Poisson maximal ideal* if it is also a maximal ideal of A .

Definition 10. Let A be a commutative Poisson algebra with a Poisson bracket $\{-, -\}$ and M be a module over A . An A -module M is called a *Poisson module* if there is a bilinear form $\{-, -\}_M : A \times M \rightarrow M$ such that

- (i) $\{a, bm\}_M = \{a, b\}m + b\{a, m\}_M$,
- (ii) $\{ab, m\}_M = a\{b, m\}_M + b\{a, m\}_M$,
- (iii) $\{\{a, b\}, m\}_M = \{a, \{b, m\}_M\}_M + \{b, \{a, m\}_M\}_M$,

for all $a, b \in A$ and $m \in M$.

A submodule N of a Poisson module M is called a *Poisson submodule* if $\{a, n\}_M \in N$ for all $a \in A$ and $n \in N$.

Definition 11. Let M be a Poisson A -module and $J \subseteq M$. The *annihilator* of J in A , in the module sense, is denoted by $\text{ann}_A(J)$ and the set $\text{Pann}_A(J)$ is defined by:

$$\text{Pann}_A(J) = \{a \in A \mid \{a, m\}_M = 0, \forall m \in J\}.$$

Theorem 12. [Jordan (2010)] Let A be a finitely generated Poisson algebra and M be a Poisson A -module. Let $J = \text{ann}_A(M)$. The following statements hold.

- (i) J is a Poisson ideal of A .
- (ii) If M is finite-dimensional and simple then J is a maximal ideal of A .
- (iii) $\square + J^2 \subseteq \text{Pann}_A(M)$.

Let $(A, \{-, -\})$ be Poisson algebra and I, J be Poisson ideals of A . Then IJ is a Poisson ideals of A . Of course I and J are Lie subalgebras of A under $\{-, -\}$. If $I \subseteq J$, then I is a Lie ideal of J and J/I is a Lie algebra. In particular, J/J^2 is always a Lie algebra. For the study on the Poisson modules, we can find that the factor J/I is a Poisson A -module with $\{a, j+I\}_{J/I} = \{a, j\}_J + I$ where I, J are Poisson ideals of A with $I \subseteq J$ and $a \in A, j \in J$. By the same argument, J/I is also a Lie algebra. Every Poisson subalgebra of J/I is a Lie ideal, so if J/I is simple as a Lie algebra, then it is simple as a Poisson module. If A is a finitely generated Poisson algebra and J is a Poisson maximal ideal, so that

$A = J + \square$, then the converse is also true because every Lie ideal of J/I is then a Poisson A -submodule.

The next theorem is the main theorem by Jordan (2010). That is a method to determine the finite-dimensional simple Poisson modules over any affine Poisson algebra.

Theorem 13. [Jordan (2010)] Let A be a finitely generated Poisson algebra.

(i) Let M be a finite-dimensional simple Poisson A -module and $J = \text{ann}_A(M)$. There is a simple module M^* for the Lie algebra $g(J)$ such that $M^* = M$, as a vector space, and $[j + J^2, m]_{M^*} = \{j, m\}_M$ for all $j \in J$ and $m \in M$.

(ii) Let J be a Poisson maximal ideal of A and N be a finite-dimensional simple $g(J)$ -module. There exist a simple Poisson A -module N' and a Lie homomorphism $f: A \rightarrow g(J)$ such that $N' = {}^f N$ as a Lie module over A and $J = \text{ann}_A(N')$.

(iii) For all finite-dimensional simple Poisson modules M , $M^{**} = M$. For all Poisson maximal ideals J of A and all finite-dimensional simple $g(J)$ -modules N , $N' = N$.

(iv) The procedure in (i) and (ii) establish a bijection Γ from the set of isomorphism classes of finite-dimensional simple Poisson module over A to the set of pairs $(J, \hat{N})$, where J is a Poisson maximal ideal of A , N is a finite-dimensional simple $g(J)$ -module and $\hat{N}$ denote its isomorphism class, given by $\Gamma(\hat{M}) = (\text{ann}_A(M), \hat{M}^*)$.

For more details of Poisson algebra and Poisson modules see in Jordan (2010).

Results

In this section, we study on the Poisson algebra A with Poisson maximal ideals $J_i; i=1,2,\dots,5$. We use the method in Erdmann & Wildon (2006) to determine the finite dimensional simple Poisson modules annihilated by Poisson maximal ideals $J_i; i=1,2,\dots,5$. The results of the study are follows.

Lemma 14. Let A be a Poisson algebra with a Poisson maximal ideal $J_1 = xA + yA + zA$. For $d \geq 1$, there is a unique d -dimensional simple Poisson A -module annihilated by J_1 .

Proof. Let A be a Poisson algebra with a Poisson maximal ideal $J_1 = xA + yA + zA$. We consider the Lie algebra $g(J_1) = J_1/J_1^2$. So, $g(J_1)$ has a basis x, y, z and the bracket $\{x, y\}, \{y, z\}$ and $\{z, x\}$ became $[x, y] = z, [y, z] = x, [z, x] = y$. Then $g(J_1) \cong su(2, \square) \cong sl(2, \square)$. It follows that, for each $d \geq 1$, A has a unique d -dimensional simple Poisson module annihilated by J_1 .

Lemma 15. Let A be a Poisson algebra with a Poisson maximal ideal $J_2 = (x+1)A + (y+1)A + (z+1)A$. For $d \geq 1$, there is a unique d -dimensional simple Poisson A -module annihilated by J_2 .

Proof. Let A be a Poisson algebra with a Poisson maximal ideal

$J_2 = (x+1)A + (y+1)A + (z+1)A$. We consider the Lie algebra $g(J_2) = J_2 / J_2^2$. Let $m = x+1$, $n = y+1$, $p = z+1$. Then $J_2 = mA + nA + pA$. So, in J_2 ,

$$\begin{aligned}\{m, n\} &= \{x+1, y+1\} = \{x, y\} \\ &= yx + z \\ &= (n-1)(m-1) + (p-1) \\ &= nm - n - m + 1 + p - 1 \\ &= nm - n - m + p.\end{aligned}$$

Similarly, we have $\{n, p\} = pn - p - n + m$ and $\{p, m\} = mp - m - p + n$.

Thus, in $g(J_2)$, the bracket $\{m, n\}$, $\{n, p\}$, $\{p, m\}$ became

$$[m, n] = -n - m + p, \quad [n, p] = -p - n + m, \quad [p, m] = -m - p + n.$$

Next, we will show that $\{[m, n], [n, p], [p, m]\}$ is linearly independent. Let β_1, β_2 and β_3 be scalars such that $\beta_1[m, n] + \beta_2[n, p] + \beta_3[p, m] = 0$. So,

$$\begin{aligned}\beta_1(-n - m + p) + \beta_2(-p - n + m) + \beta_3(-m - p + n) &= 0 \\ (-\beta_1 - \beta_2 + \beta_3)n + (-\beta_1 + \beta_2 - \beta_3)m + (\beta_1 - \beta_2 - \beta_3)p &= 0.\end{aligned}$$

Then we obtain that

$$\begin{aligned}-\beta_1 - \beta_2 + \beta_3 &= 0, \\ -\beta_1 + \beta_2 - \beta_3 &= 0, \\ \beta_1 - \beta_2 - \beta_3 &= 0.\end{aligned}$$

This system can be written as following.

$$\begin{pmatrix} -1 & -1 & 1 \\ -1 & 1 & -1 \\ 1 & -1 & -1 \end{pmatrix} \begin{pmatrix} \beta_1 \\ \beta_2 \\ \beta_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}.$$

Thus it has the augmented matrix $\left(\begin{array}{ccc|c} -1 & -1 & 1 & 0 \\ -1 & 1 & -1 & 0 \\ 1 & -1 & -1 & 0 \end{array} \right)$, which has non-zero

determinant. This implies that the system has exactly one solution that is $\beta_1 = \beta_2 = \beta_3 = 0$. Hence $\{[m, n], [n, p], [p, m]\}$ is linearly independent. So, the derived

algebra $[g(J_2), g(J_2)]$ has dimension 3, which implies that $g(J_2) \cong sl(2, \mathbb{C})$. By Handerson (2012), it is a well-known that $sl(2, \mathbb{C})$ has a unique d -dimensional simple Poisson module annihilated by J_2 for each $d \geq 1$ and $sl(2, \mathbb{C})$ is 1-homogeneous. By Theorem 13, the result hold for A .

Lemma 16. Let A be a Poisson algebra with a Poisson maximal ideal

$J_3 = (x+1)A + (y-1)A + (z-1)A$. For $d \geq 1$, there is a unique d -dimensional simple Poisson A -module annihilated by J_3 .

Proof. Let A be a Poisson algebra with a Poisson maximal ideal

$J_3 = (x+1)A + (y-1)A + (z-1)A$. We consider the Lie algebra $g(J_3) = J_3 / J_3^2$. Let $m = x+1$, $n = y-1$, $p = z-1$. Then $J_3 = mA + nA + pA$. So, in J_3 ,

$$\{m, n\} = nm - n + m + p$$

$$\{n, p\} = pn + p + n + m$$

$$\{p, m\} = mp + m - p + n.$$

Thus, in $g(J_3)$, the bracket $\{m, n\}$, $\{n, p\}$, $\{p, m\}$ became
 $[m, n] = -n + m + p$, $[n, p] = p + n + m$, $[p, m] = m - p + n$.

The proof is similar to Lemma 15 that $\{[m, n], [n, p], [p, m]\}$ is linearly independent. Thus, the derived algebra $[g(J_3), g(J_3)]$ has dimension 3, which implies that $g(J_3) \cong sl(2, \mathbb{C})$. Hence A has a unique d -dimensional simple Poisson module annihilated by J_3 for each $d \geq 1$.

Lemma 17. Let A be a Poisson algebra with a Poisson maximal ideal

$J_4 = (x-1)A + (y+1)A + (z-1)A$. For $d \geq 1$, there is a unique d -dimensional simple Poisson A -module annihilated by J_4 .

Proof. Let A be a Poisson algebra with a Poisson maximal ideal

$J_4 = (x-1)A + (y+1)A + (z-1)A$. We consider the Lie algebra $g(J_4) = J_4 / J_4^2$. Let $m = x-1$, $n = y+1$, $p = z-1$. Then $J_4 = mA + nA + pA$. So, in J_4 ,

$$\{m, n\} = nm + n - m + p,$$

$$\{n, p\} = pn - p + n + m, \text{ and}$$

$$\{p, m\} = mp + m + p + n.$$

Thus, in $g(J_4)$, the bracket $\{m, n\}$, $\{n, p\}$, $\{p, m\}$ became

$$[m, n] = n - m + p, \quad [n, p] = -p + n + m, \quad [p, m] = m + p + n.$$

The proof is similar to Lemma 15 that $\{[m, n], [n, p], [p, m]\}$ is linearly independent. Thus, the derived algebra $[g(J_4), g(J_4)]$ has dimension 3, which implies that $g(J_4) \cong sl(2, \mathbb{C})$. Hence A has a unique d -dimensional simple Poisson module annihilated by J_4 for each $d \geq 1$.

Lemma 18. Let A be a Poisson algebra with a Poisson maximal ideal

$J_5 = (x-1)A + (y-1)A + (z+1)A$. For $d \geq 1$, there is a unique d -dimensional simple Poisson A -module annihilated by J_5 .

Proof. Let A be a Poisson algebra with a Poisson maximal ideal

$J_5 = (x-1)A + (y-1)A + (z+1)A$. We consider the Lie algebra $g(J_5) = J_5/J_5^2$. Let $m = x-1$, $n = y-1$, $p = z+1$. Then $J_5 = mA + nA + pA$. So, in J_5 ,

$$\{m, n\} = nm + n + m + p,$$

$$\{n, p\} = pn + p - n + m, \text{ and}$$

$$\{p, m\} = mp - m + p + n.$$

Thus, in $g(J_5)$, the bracket $\{m, n\}$, $\{n, p\}$, $\{p, m\}$ became $[m, n] = n + m + p$, $[n, p] = p - n + m$, $[p, m] = p - m + n$.

The proof is similar to Lemma 15 that $\{[m, n], [n, p], [p, m]\}$ is linearly independent. Thus, the derived algebra $[g(J_5), g(J_5)]$ has dimension 3, which implies that $g(J_5) \cong sl(2, \mathbb{C})$. Hence A has a unique d -dimensional simple Poisson module annihilated by J_5 for each $d \geq 1$.

Theorem 19. Let A be a Poisson algebra with Poisson bracket

$$\{x, y\} = yx + z, \quad \{y, z\} = zy + x, \quad \{z, x\} = xz + y$$

and the Poisson maximal ideals J_i ; $i=1, 2, \dots, 5$. There is a unique d -dimensional simple Poisson A -module annihilated by J_i for each $d \geq 1$.

Proof. The proof of this theorem is directly obtained from Lemma 14 - Lemma 18.

Conclusion

The study on the Poisson algebra with some bracket could bring the difficulty to find its finite-dimensional simple Poisson modules. In this research, we apply the method shown in Jordan (2010) to determine the finite-dimensional simple Poisson module over a Poisson

algebra A . We can found that there is a unique d -dimensional simple Poisson A -module annihilated by J_i for each $d \geq 1$.

Acknowledgement(s)

The authors would also like to thanks the referees for their careful reading of the paper and their useful comments.

References

- Changtong, K., & Sasom, N. (2018). Poisson Maximal Ideals and the Finite-dimensional Simple Poisson Modules of a Certain Poisson Algebra. *KKU Science Journal*, 46(3), 606-613.
- Erdmann, K., & Wildon, M. J. (2006). *Introduction to Lie Algebras* (1st Ed.). London: Springer-Verlag.
- Henderson, A. (2012). *Representations of Lie Algebras: An Introduction Tthrough gl_n* (1st Ed.). UK: Cambridge University Press.
- Jordan, D. A. (2010). Finite-dimensional Simple Poisson Modules. *Algebras and Representation Theory*, 13, 79-101.
- Kirillov, A. J. (2008). *Introduction to Lie groups and Lie algebras* (1st Ed.). New York: Cambridge University Press.
- Pfeifer, W. (2003). *The Lie Algebras $su(N)$: An Introduction* (3^{ed} Ed.). Switzerland: Birkhäuser Basel.
- Sasom, N. (2006). *Reversible skew Laurent polynomial rings, rings of invariants and related ring* [Unpublished doctoral's thesis]. University of Sheffield.
- Sasom, N. (2012). Finite-dimensional Simple Poisson Modules. *Chiang Mai Journal of Science*, 39(4), 678-687.

Research Article

An approximation of average run length using numerical integration methods on EWMA control chart for MAX(q,r) process

Piyaphon Paichit ¹ and Kanita Petcharat ^{2*}

¹ Faculty of Science, Silpakorn University, Sanamchandra palace campus, NakhonPathom 73000, Thailand

² Faculty of Applied Science, King Mongkut's University of Technology North Bangkok, Bangsue, Bangkok 10800, Thailand

*E-mail: kanita.p@sci.kmutnb.ac.th

Received: 27/01/2022; Revised: 21/04/2022; Accepted: 10/05/2022

Abstract

This research aims to measure a method for estimating the average run length (ARL) of an exponentially weighted moving average control chart using numerical integration when the data is a moving average model with exogenous variables. We compare the average run lengths achieved using three distinct methodologies. The midpoint, trapezoid, and Gaussian rules are all used. Additionally, we compare the CPU time used by the ARL assessment. The simulations show that the ARL values obtained from using the midpoint and Gaussian rules are similar. The result obtained from using the trapezoid approach, on the other hand, is less different at roughly 1%. Additionally, the midpoint and trapezoid approaches were the fastest when CPU time was considered, requiring between 4-6 seconds. On the other hand, Gaussian's rule requires around 37-43 seconds.

Keywords: EWMA, control chart, average run length, numerical integral, moving average

Introduction

Quality control is an essential tool for reducing the many deficiencies throughout the manufacturing process. Statistical Process Control (SPC) is used to design and enhance processes. Control charts are a robust process of monitoring or regulating manufacturing processes in real-time. In 1924, Walter A. Shewhart invented the control chart to help manufacturers minimize waste and enhance quality. Control charts have since been extensively utilized for identifying and monitoring effects on the quality of processes in various applications, including industrial production, public health, computer networking and telecommunications, finance and economics, environmental research, and others. Nowadays, Shewhart control charts, cumulative sum charts (CUSUM) and exponentially weighted moving average control charts (EWMA) are popularly used in industrial processes. Robert (1959) developed the EWMA control chart to detect a slight change in the mean, suitable for normal distribution. It is now well accepted that a strong control chart is essential for identifying small and moderate changes in the process. Control chart performance is often measured by the average run length (ARL). In most cases, the ARL_0 indicates that the process is under control, whereas the ARL_1 suggests that the process is out of control and should be small. The ARL of EWMA control

chart may be estimated using the Monte Carlo Simulations (MC) technique, which is the conventional way to test the validity and compare it to the other methods. However, processing takes a long time, the Martingale (Sukparangsri & Navikov, 2008), Markov chain (Brook and Evan, 1972; Lucus & Saccicci, 1990) and integral equation (See Srivastava & Wu, 1997; Areepong, 2009; Mittitelu et al., 2010; Petcharat et al., 2013; Paichit, 2017; Preerajit et al., 2018).

The serially correlated data may occur in the real world of the data set when the data are AR(1) and ARMA(1,1) processes. Lu & Reynolds (1999) employed the integral equation technique to derive the ARL. Some literature try to evaluate ARL when data set are time series model. Petcharat et al. (2013) proposed the analytical expression for the ARL of the EWMA control chart for the moving average of order q (MA(q)) process. For the seasonal AR(P)_L process. Later, Petcharat (2016) derived the explicit formula of ARL for the EWMA chart and compared it to the CUSUM control chart and found that the EWMA chart was more sensitive to detecting small shifts. Paichit (2017) proved the explicit formula of ARL for the EWMA control chart base on autoregressive with an explanatory variable model. Sunthornwa et al. (2018) proved exact solution and approximate the numerical of ARL for EWMA control charts using the seasonal ARFIMA with exponential white noise. Moreover, Suriyakart (2020) studied the sensitivity of the EWMA control chart for detecting the mean change in processes when processes are AR(1), ARMA(1,1) and IMA(1) with exponential white noise and found that the EWMA control chart is good performance for detecting the small change of the process mean. Sunthornwat & Areepong (2020) derived analytical ARL and approximately the numerical ARL for the CUSUM control chart for seasonal and non-seasonal moving averages processes. Recently, Petcharat (2021) presented the exact formula of ARL for the MAX(1, r) process with exponential white noise, with its existence and uniqueness being proved by Banach's fixed-point theory.

According to this literature review, the ARL for measuring the efficiency of the EWMA control chart is very important for comparing control chart performance. In this research, we extend Petcharat (2021) to evaluate the ARL of the EWMA control chart by using numerical integration methods for data is the model moving average model with exogenous variables (MAX(q , r)).

Methods

Exponentially weighted moving average (EWMA) control chart

Robert (1959) initially suggested the EWMA control chart. The recursive equation below would be used to express the EWMA control chart.

$$Y_t = (1-\lambda)Y_{t-1} + \lambda Z_t, \quad t=1,2,\dots, \quad (1)$$

where λ is exponential smoothing parameter, $0 < \lambda < 1$. Z_t is sequence of process observations. The control limit of control charts are upper control limit (UCL) and lower control limit (LCL) as follows,

$$LCL/UCL = \mu \pm L\sigma \sqrt{\frac{\lambda}{2-\lambda}}, \quad (2)$$

where μ is average or target value, L is the width of the control limit and σ is standard deviation of the process.

In this research, we present EWMA control chart for the process observations is the sequence of a moving average process with explanatory variables (MAX(q,r)) with exponential white noise defined as

$$Z_t = \mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \theta_2 \varepsilon_{t-2} - \dots - \theta_q \varepsilon_{t-q} + \sum_{i=1}^r \beta_i X_{it}, \quad (3)$$

where ε_t is exponential white noise : $\varepsilon_t \sim \text{Exp}(\alpha)$, θ_i is moving average coefficient, $i=1,2,\dots,q$,

$-1 \leq \theta_i \leq 1$, X_{it} is exogenous variable and β_i is a coefficient of X_{it} . $\varepsilon_t \sim \sum_{i=1}^r \beta_i X_{it}$

The corresponding stopping time for (1) define as

$$\tau = \inf \{t > 0; Y_t > b\}, H_0 = u, \quad b > y. \quad (4)$$

where b denote control limit.

Let $E_\omega(\cdot)$ denote the expectation under probability density function $f(y, \alpha)$ that the change-point occurs at point ω , where $\omega < \infty$. Thus by definition, the ARL for MAX(q,r) process with an initial value $H_0 = u$ is as follow

$$ARL = H(u) = E_\infty(\tau_b) < \infty. \quad (5)$$

Numerical integral equation (NIE) of average run length for MAX(q,r) on EWMA control chart

Let $H(u)$ be the average run length for the EWMA control chart as defined in equation (5). Assume that $C_0 = u$ be the initial process under in-control. The integral equation $H(u)$ can be derived by Fredholm integral equation of the second kind denotes as follows:

$$H(u) = 1 + \frac{1}{\lambda} \int_0^b h(y) f\left(\frac{y - (1-\lambda)u}{\lambda} + \mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \theta_2 \varepsilon_{t-2} - \dots - \theta_q \varepsilon_{t-q} + \sum_{i=1}^r \beta_i X_{it}\right) d(y),$$

Such that

$$H(u) = 1 + \frac{1}{\lambda \alpha} \int_0^b h(y) e^{\frac{y}{\lambda \alpha}} e^{-\frac{(1-\lambda)u}{\lambda \alpha} + \frac{\mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \theta_2 \varepsilon_{t-2} - \dots - \theta_q \varepsilon_{t-q} + \sum_{i=1}^r \beta_i X_{it}}{\alpha}} d(y), \quad (6)$$

The integral equation $H(u)$ may be approximated using numerical quadrature procedures. The Gaussian rule, the midpoint rule, and the trapezoidal rule are all employed in this research.

I. Gaussian Rule

The quadrature rule is used to approximate an integral, and the Gauss-Legendre quadrature rule is used to approximate the solution. Suppose the $\{\varpi_j, j = 1, \dots, m\}$ is a point on

the interval $[0, b]$ and set $\{k_j, j = 1, \dots, m\}$ as $k_j = b/m \geq 0; j = 1, 2, \dots, m$. The quadrature rule evaluates an integral's approximation as follows:

$$\int_0^b K(y) f(y) dy \approx \sum_{j=1}^m \varpi_j f(a_j),$$

where $\varpi_j = \frac{b}{m} \left(j - \frac{1}{2} \right); j = 1, 2, \dots, m$.

It can be written $R_{m \times m}$ be matrix form as

$$[R]_{jl} = \frac{1}{\lambda} \sum_{j=1}^m k_j f \left(\frac{\varpi_k - (1-\lambda)\varpi_m}{\lambda} + \left(\mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q} + \sum_{i=1}^r \beta_i X_{it} \right) \right), j, l = 1, 2, \dots, m.$$

Let $H_G(u)$ be the numerical approximation for an integral equation solution in equation (6) which can be found as the linear equations as follows:

$$H_G(u) = 1 + \frac{1}{\lambda} \sum_{j=1}^m k_j H(\varpi_k) f(\mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q} + \sum_{i=1}^r \beta_i X_{it}), \quad (7)$$

where $\varpi_j = \frac{b}{m} \left(j - \frac{1}{2} \right); j = 1, 2, \dots, m$.

II. Midpoint Rule

By using the midpoint rule, Let m be subinterval on $[0, b]$ and let

$$f(A_j) = f \left(\frac{\varpi_k - (1-\lambda)u}{\lambda} + \left(\mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q} + \sum_{i=1}^r \beta_i X_{it} \right) \right).$$

The approximation for the integral equation solution in equation (6) is given by

$$H_M(u) = 1 + \frac{1}{\lambda} \sum_{j=1}^m k_j L(\varpi_j) f(A_j), \quad (8)$$

where $\varpi_j = \frac{b}{m} \left(j - \frac{1}{2} \right)$ and $k_j = \frac{b}{m}; j = 1, 2, \dots, m$.

III. Trapezoidal rule

By using the trapezoidal rule, suppose m be subinterval on $[0, b]$. let

$$f(A_j) = f \left(\frac{\varpi_k - (1-\lambda)u}{\lambda} + \left(\mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q} + \sum_{i=1}^r \beta_i X_{it} \right) \right).$$

The approximation for the integral equation solution in equation (6) is given by

$$H_T(u) = 1 + \frac{1}{\lambda} \sum_{j=1}^{m+1} k_j L(\varpi_j) f(A_j), \quad (9)$$

where $\varpi_j = jk_j$ and $k_j = \frac{b}{m}$; $j = 1, 2, \dots, m-1$, in other case, $k_j = \frac{b}{2m}$.

Results and discussion

In this part, we compare ARL values for EWMA control chart on $MAX(q,r)$ process with exponential white noise obtain from NIE methods which are Gaussian Rule, midpoint rule and Trapezoidal rule according to equation (7) to (9) with $m=800$ subintervals. We set $H_G(u)$ is ARL from NIE method using the Gaussian rule from explicit formula, $H_M(u)$ is ARL from NIE method using the midpoint rule and $H_T(u)$ is ARL from NIE method using the trapezoidal rule. Additionally, we compare the computing times of three approaches. The computational time for the three methods is estimated in seconds by the central processing unit (CPU) time (Windows 8 OEM, Intel(R) core(TM)i5-8265U CPU@1.60GHz 1.80GHz RAM 8.00 GB (7.89 GB usable)).

In Table 1, the parameters value b for EWMA control chart was selected by setting $\lambda = (0.05, 0.10, 0.20)$, $ARL_0 = 370$ and $\alpha_0 = 1$ in the case of $MAX(1,1)$ with parameter $\beta = 1.0$, $\theta_1 = (0.15, 0.25, 0.30)$ and $\theta_2 = (0.10, 0.25, 0.50)$, respectively.

Table 1 Comparison of ARL_0 between using numerical integration for $MAX(2,2)$ process with parameter $\alpha_0 = 1$ for $ARL_0 = 370$

λ	θ_1	θ_2	b	$H_M(u)$ (Time: seconds)	$H_G(u)$ (Time: seconds)	$H_T(u)$ (Time: seconds)	%Diff
0.05	0.15	0.10	0.019481940	370.577 (4.765)	370.577 (39.703)	370.286 (5.344)	0.0785
		0.25	0.022672201	370.514 (4.781)	370.514 (38.516)	370.349 (5.344)	0.0445
		0.50	0.029210100	370.435 (4.875)	370.435 (41.36)	370.208 (5.218)	0.0613
0.10	0.25	0.10	0.043581900	370.382 (5.25)	370.382 (41.047)	370.156 (5.328)	0.0610
		0.25	0.050820500	370.627 (4.875)	370.627 (41.75)	370.402 (5.219)	0.0607
		0.50	0.065747070	370.268 (4.953)	370.268 (41.203)	370.045 (5.266)	0.0602
0.20	0.3	0.10	0.093940200	370.287 (4.875)	370.287 (41.734)	370.067 (5.500)	0.0594
		0.25	0.110015960	370.527 (4.813)	370.527 (41.265)	370.309 (5.297)	0.0588
		0.50	0.143624000	370.392 (4.796)	370.392 (41.188)	370.177 (5.430)	0.0580

As shown in Table 1, from NIE methods using Gaussian and midpoint rules are similar values, but using trapezoidal rules yield slightly different results on $m=800$ subintervals. If we take into account the absolute percentage error computed by $\%Diff = \frac{H_M(u) - H_T(u)}{H_M(u)} \times 100$. It can

note that ARL_0 values using the trapezoid rule $\%Diff$ are less than 0.08% compared to other methods. However, the CPU time using midpoint and trapezoidal rules are less than the CPU time using the Gaussian rule.

The ARL values indicate the performance of EWMA control chart for detection change in the process mean using Gaussian, midpoint, and trapezoidal rules for $m=800$ subintervals shown in Table 2 to Table 4. We set the parameter $\alpha_0=1$ when process is in-control and the parameter $\alpha_1 = \alpha_0(1 + \delta)$ when process is out of control where shift size (δ) = 0.005, 0.02, 0.04, 0.06, 0.08, 0.1, 0.5, and 1.0. In Table 2, the initial $ARL_0 = 370$ for MAX(2,1) with parameters $\theta_1 = 0.35$, $\theta_2 = 0.6$ with ARL_0 , the initial In Table 3 = 0.01761086. $b = 0.05$ and $\lambda = 2.0$, $\beta_1 = 0.05$ and $b = 0.022672201$. In Table 4, the initial $ARL_0 = 370$ for MAX(3,2) with parameters $\theta_1 = 0.15$, $\theta_2 = 0.25$, $\theta_3 = 0.45$ with ARL_0 , the initial In Table 3 = 0.03579842. $b = 0.05$ and $\lambda = 0.7$, $\beta_2 = 0.5$, β_1

Table 2 ARL values for MAX(2,2) with parameters $\theta_1 = 0.15$, $\theta_2 = 0.25$, with $\lambda = 0.7$, $\beta_2 = 0.5$, $\beta_1 = 0.05$ and $b = 0.022672201$ band $= 0.05$ and $\lambda = 370$

shift (δ)	$H_M(u)$ (Time: seconds)	$H_G(u)$ (Time: seconds)	$H_T(u)$ (Time: seconds)
0	370.315 (4.172)	370.315 (37.062)	370.086 (4.703)
0.005	68.831 (4.219)	68.831 (38.047)	68.789 (4.563)
0.02	20.6369 (4.125)	20.6369 (38.578)	20.6248 (4.781)
0.04	11.0660 (4.266)	11.0660 (40.781)	11.0598 (4.906)
0.06	7.75846 (4.328)	7.75846 (41.531)	7.75435 (5.031)
0.08	6.08201 (4.187)	6.08201 (41.406)	6.07893 (5.156)
0.10	5.06887 (4.187)	5.06887 (41.719)	5.06641 (5.031)
0.5	1.7953 (4.250)	1.7953 (37.531)	1.79483 (5.047)
1.0	1.38795 (4.157)	1.38795 (42.422)	1.38772 (5.125)

Table 3 ARL values for max (2,1) using numerical integral equation method given $\theta_1 = 0.35$, $\theta_2 = 0.6$ with $= ARL_0$ for $=0.01761086$ $b = 0.05$, $\lambda = 2.0$, $\beta = 370$

shift (δ)	$H_M(u)$ (Time: seconds)	$H_G(u)$ (Time: seconds)	$H_T(u)$ (Time: seconds)
0	370.5144 (4.781)	370.5144 (37.422)	370.286 (4.89)
0.005	73.8584 (4.735)	73.8584 (37.984)	73.8135 (4.906)
0.02	22.3826 (4.657)	22.3826 (38.578)	22.3695 (5.219)
0.04	12.0145 (4.735)	12.0145 (39.515)	12.0078 (4.985)
0.06	8.4196 (4.984)	8.4196 (39.172)	8.4151 (4.875)
0.08	6.59482 (4.781)	6.59482 (38.9999)	6.59145 (4.969)
0.10	5.49099 (4.625)	5.49099 (40.235)	5.48829 (5.234)
0.5	1.91019 (4.718)	1.91019 (39.391)	1.909669 (5.469)
1.0	1.45598 (4.828)	1.45598 (40.594)	1.45572 (5.343)

Table 4 ARL values for max (2,2) using numerical integral equation method given $\theta_1 = 0.15$, $\theta_2 = 0.25$, $\theta_3 = 0.45$ with $= ARL_0$ for $=0.03579842$. $b=0.05$ and $\lambda = 0.7$, $\beta_2 = 0.5$, $\beta_1 = 370$

Shift (δ)	$H_M(u)$ (Time: seconds)	$H_G(u)$ (Time: seconds)	$H_T(u)$ (Time: seconds)
0	370.584 (4.781)	370.584 (37.36)	370.358 (5.094)
0.005	85.2404 (4.735)	85.2404 (37.984)	85.1888 (5.172)
0.02	26.4916 (4.657)	26.4916 (38.922)	26.4761 (5.016)
0.04	14.2629 (4.735)	14.2629 (38.922)	14.2549 (5.312)
0.06	9.98991 (4.984)	9.98991 (39.860)	9.9845 (5.719)
0.08	7.8138 (4.781)	7.81380 (37.265)	7.80973 (5.515)
0.10	6.49476 (4.625)	6.49476 (37.157)	6.49149 (5.407)
0.5	2.18525 (4.718)	2.18525 (37.578)	2.18458 (5.437)
1.0	1.62088 (4.828)	1.62088 (37.375)	1.62053 (5.203)

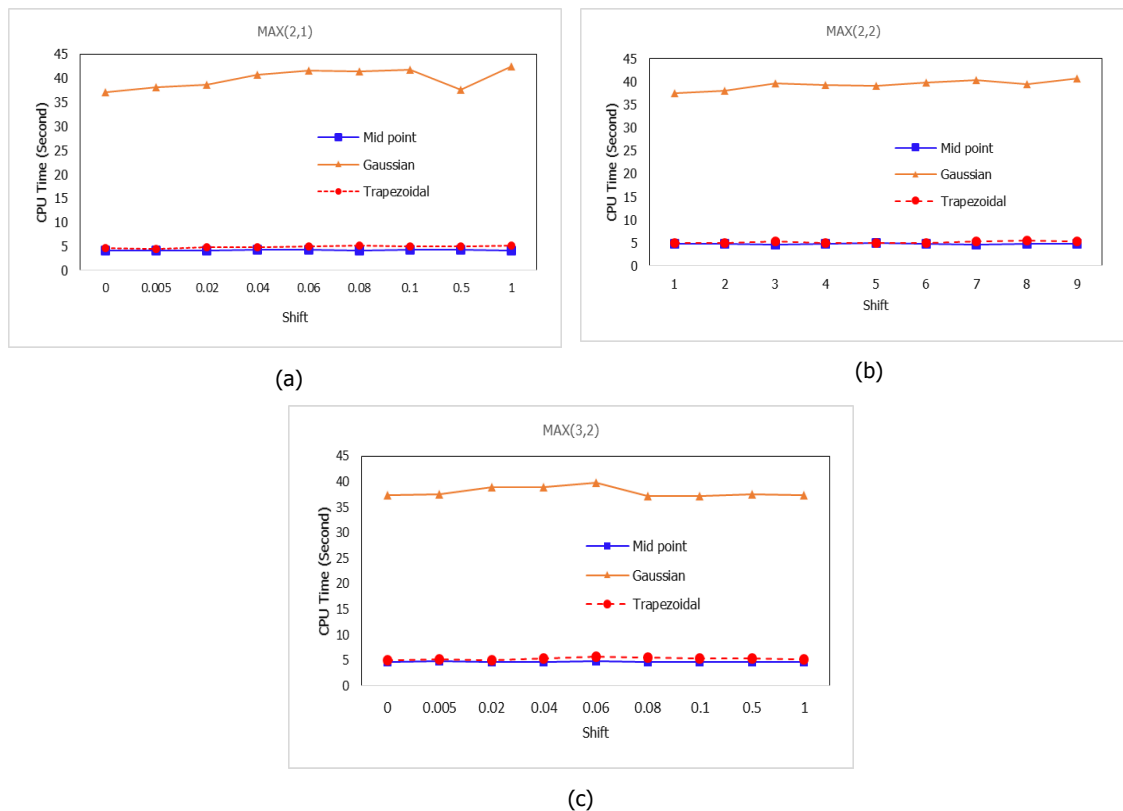


Figure 1 The CPU times for evaluating ARL values of EWMA control chart : (a) MAX(2,1) (b) MAX(2,2) (c) MAX(3,2)

Table 2 to Table 4 show that ARL_0 from the NIE method using Gaussian, midpoint and trapezoid rules on $m=800$ subintervals were equally effective in detecting process changes. In terms of CPU time, we discovered that the midpoint rule takes the least time to calculate, followed by the trapezoid rule and the Gaussian rule. Refer to Figures 1.

Conclusion

This paper proposes the numerical integration method using Gaussian, midpoint and trapezoid rules of ARL for moving average process with explanatory variables (MAX(q,r)) on the EWMA control chart. The results found that ARL values from the three methods are close together. The computational times for using midpoint and trapezoid rules take less than 6 seconds. Besides, using the Gaussian rule takes more CPU times, around 37 - 43 seconds. As a result, employing midpoint and trapezoid rules may significantly reduce computing time compared to the Gaussian rule.

Acknowledgement

This research was funded by Faculty of Applied Science, King Mongkut's University of Technology North Bangkok Contract no. 641081.

References

- Areepong, Y. (2009). *An integral equation approach for analysis of control charts* [Unpublished doctoral thesis]. University of Technology.
- Brook, D., & Evans, D. A. (1972). An approach to the probability distribution of CUSUM run lengths. *Biometrika*, 59, 539-548.
- Shewhart, W. A. (1931). *Economic control of quality of manufactured product*. Van Nostrand, New York.
- Lu, C. W., & Reynolds, M. R. (1999). EWMA control charts for monitoring the mean of autocorrelated processes. *Journal of Quality Technology*, 31, 166-188.
- Lucas, J. M., & Saccucci, M. S. (1990). Exponentially weighted moving average control schemes: properties and enhancements. *Technometrics*, 32(1), 1-29.
- Paichit, P. (2017). Exact expression for average run length of control chart of ARX(p) procedure. *KKU Science Journal*, 45(4), 948-958.
- Peerajit, W., Areepong, Y., & Sukparungsee, S. (2018). Numerical integral equation method for ARL of CUSUM chart for long-memory process with non-seasonal and seasonal ARFIMA models. *Thailand Statistician*, 16(1), 26-37.
- Petcharat, K., Areepong, Y., Sukparungsee, S., & Mititelu, G. (2013). Exact solution of average run length of EWMA chart for MA(q) processes. *Far East Journal of Mathematical Sciences*, 78(2), 291-300.
- Petcharat, K. (2016). Explicit formula of ARL for SMA(Q)_L with exponential white noise on EWMA chart. *International Journal of Applied Physics and Mathematics*, 6(4), 218-255. <https://doi.org/10.17706/ijapm.2016.6.4.218-225>
- Petcharat, K. (2021). The performance of EWMA control chart for MAX(1,r) process. *Proceedings of The International Multi Conference of Engineers and Computer Scientists, 20-22 October, 2021, Hong Kong*.
- Robert, S. W. (1959). Control chart tests based on geometric moving average. *Technometrics*, 411-430.
- Sukparungsee, S., & Novikov, A. A. (2008). Analytical approximations for detection of a change-point in case of light-tailed distributions. *Journal of Quality Measurement and Analysis*, 4, 49-56. https://www.ukm.my/jqma/v4_2/JQMA-4-2-05-abstractrefs.pdf
- Sunthornwat, R., Areepong, Y., & Sukparungsee, S. (2018). Average run length of cumulative sum control charts for SARMA(1,1)_L models. *Thailand Statistician*, 15(2), 184-195.
- Sunthornwat, R., & Areepong, Y. (2020). Average run length on CUSUM control chart for seasonal and non-seasonal moving average processes with exogenous variables. *Symmetry*, 12(1), 173. <https://doi.org/10.3390/sym12010173>
- Suriyakat, W. (2020). On sensitivity of control chart for monitoring serially correlated data. *Interdisciplinary Research Review*, 15(3), 44-47.

Research Article

Non – Stationary Booking Demand in Two-Class Overbooking Model

Piyachat Leelasilapasart¹ and Murati Somboon^{1*}

¹Department of Applied Statistics, Faculty of Applied Science, King Mongkut's University of Technology North Bangkok, 10800 Thailand

*E-mail: murati.s@sci.kmutnb.ac.th

Received: 21/01/2022; Revised: 29/03/2022; Accepted: 27/04/2022

Abstract

Accurate demand forecasting is a crucial component of airline revenue management. However, it is cumbersome to find exact forecasting demand since the historical data, the censor data, do not reflect the actual demand. In order to obtain estimated demand, the Expectation-Maximization (EM) method is used to calculate the estimated demand. Furthermore, the previous research mainly focuses on the stationary demand, which does not represent an actual situation. So, this research engages in finding the optimal overbooking and booking limit of the Two-Class overbooking with non-stationary demand that maximizes the expected profit using numerical experiments. In addition, the performance of the Two-Class overbooking model with non-stationary demand is compared with the booking limit of the airline's policy. The optimal number of update booking limits with non – stationary demand using real data are concerned.

The simulation finds that the optimal booking calculated mostly is the booking limit. However, the booking limit is suited to the overbooking limit in some situations. Moreover, the expected profit from the optimal overbooking limit is greater than or equal to the airline's booking policy in all study situations. The Two-Class overbooking model yielded up to 21% more the expected profit margins than airline booking policies. In real data study, the Two-Class overbooking model is taken to calculate the optimal number of update booking limits when the demand is non – stationary and it is found that the 16-point update booking limit gives the maximum profit.

Keywords: Revenue Management, Overbooking, Two-Class overbooking, Non-stationary demand, Expectation-Maximization.

1. Introduction

Revenue management (RM) is a well-known application of mathematical modeling mostly employed in the airline industry to improve profitability with high fixed costs and low marginal costs. RM is defined as "selling the correct product to the correct customer at the correct time for the proper price" (Cross,1997). The aviation industry is an intensely competitive industry that constantly engages in overbooking to reinforce its revenues, taking advantage of the fact that customers cancel airfare reservations or do not show up for flights (Suzuki, 2006). The concept is straightforward. Airlines sell seats by creating different fare classes carrying different restrictions and prices.

The Air Transport Information Division of Airport of Thailand Public Company Limited (AOT) reported that the total air traffic, aircraft movement, and passenger numbers in Thailand increased by 21, 22, and 21 percent from 2015 to 2019. So, Thailand can increase revenue by roughly \$200 million US.

Unfortunately, after the World Health Organization (WHO) declared the COVID-19 outbreak a pandemic since March 2020, daily life across the world has changed. The aviation industry has been one of the hardest-hit industries globally since the beginning of the outbreak. Some airlines would evolve during a crisis like no others bringing the industry into survival mode, diminished by a loss of traffic and revenues. While positive signs for recovery, COVID-19 remains an existential crisis for airports, airlines, and business sectors. The airline companies in Thailand are also affected by the pandemic outbreak; however, they could benefit from using RM strategies when the pandemic is better.

The Two-Class model, which concentrates on the booking control problem, dates back to Littlewood (1972). All booking requests show up at the time of service. Littlewood (1972) does not allow overbooking, which means there are no cancellations or no-shows. The Two-Class overbooking model proposed by Somboon and Amaruchkul (2016) is taken into account in this study. The model combines two strategies, overbooking and seat inventory, and proposes a static Two-Class overbooking model, in which low fare (Class-2) passengers arrive prior to high fare (Class-1) passengers. The airline incurs a penalty cost, any of which rejected booking requests. Also, the penalties are different for the two classes. The airline may overbook Class-2 customers. The two fare classes may have different show-up rates. We refer an interested reader to Somboon and Amaruchkul (2016) and also the literature cited therein for more details on the combined Two-Class overbooking model.

Accurate demand forecasting is a vital component in RM. The booking data of departure flights are used to forecast the demand for future departing flights. However, the previous research generates only the stationary demand, which does not represent an actual situation. A pattern in which demand is not constant for each time but varies due to seasonality, trend or other factors, the non-stationary demand, is considered. For example, a booking demand changes according to the number of days close to the travel date. The closer the travel date is, the more demand for booking. This research mainly focuses on finding the optimal overbooking and booking limit that maximizes the expected profit using numerical experiments. Moreover, the performance of the Two-Class overbooking model with non-stationary demand is compared with the airline's policy booking limit. Finally, find the optimal number of update booking limits when the demand is non-stationary with real data.

The paper is organized, as follows: Section 2 briefly reviews the Two-Class overbooking model. Section 3 details the numerical setting and the procedure for simulation. Section 4 presents numerical results of the numerical study. Real data is applied in Section 5, and Section 6 presents conclude.

2. Two-Class Overbooking Model

In 2016, Somboon and Amaruchkul propose the two-class overbooking model by the following these information.

Let $\mathcal{C}_+$ and $\mathcal{I}$ be the set of non-negative integers and the real number set respectively. Also, given t is the number of days until departure and suppose that $(y)^+ = \max(0, y)$ for $y \in \mathcal{I}$. Let $D(t)$ be the distribution function of a random variable at t days before departure. The quantile function of $D(t)$ can be written as

$$F_{D(t)}^{-1}(a) = \inf\{x: P(D(t) \leq x) \geq a\}$$

Class-2 reservations are assumed to begin a booking before Class-1 reservations in this model. When a Class- i customer is accepted, the airline earns revenue p_i , $p_1 > p_2 > 0$, for each $i = 1, 2$. If the airline rejects the request, it will incur a penalty cost g_i where $g_1 > g_2 > 0$. In this model, the penalty cost includes both opportunity cost and loss-of-goodwill. Loss of business reputation or the loss-of-goodwill determines the customer satisfaction that can be intangible and cumbersome to gauge in practice.

The opportunity cost is a measurement of future revenue loss based on what happens after the lost sales occur. If a customer is prone to return to make a request, the opportunity cost is the expected loss of revenue from this situation. On the other hand, if a customer never returns to make another reservation with the airline, the opportunity cost includes all future revenues.

In practice, an optimal booking limit is determined on a proportion of refund cost over time and change in show-up probability. It can be affected by overbooking limits that vary over time. Assume that $x(t) \in \mathbb{Z}_+$ is a Class-2 booking limit at the update booking limit t days prior to departure. It implies that the reservations for Class-2 are accepted up to $x(t)$. Overbooking is permitted if $x(t)$ is greater than capacity $k(t)$ at the update booking limit point t . Let $D_i(t)$ be the demand of Class- i booking at the update booking limit point t for each $i = 1, 2$. The Class-2 reservations and rejected quantity are $\min(x(t), D_2(t))$ and $(D_2(t) - x(t))^+$ at the update booking limit point t , respectively.

After all Class-2 reservations have arrived, Class-1 customers begin their booking so that the available capacity for Class-1 customers will be $(k(t) - \min(x(t), D_2(t)))^+$. However, due to high penalty costs and high priority, we do not overbook Class-1 passengers, so they are accepted up to the remaining capacity. The number of Class- i reservations at the updated booking limit point t can be written as:

$$B_1(x(t)) = \min(k(t) - B_2(x(t)))^+, D_1(t), B_2(x(t)) = \min(x(t), D_2(t)).$$

Some customers may not show up for or cancel their reservations before departure. Furthermore, we considered that both no-show and cancellation customers are the same in this model. Let $B_i(x(t)) = y_i$ be the number of Class- i reservations and $W_i(y_i)$ be the number of Class- i show-ups. Assume $W_i(y_i)$ is a Class- i show-up probability with a binomial distribution with parameters y_i and q_i , where $q_i \in (0, 1]$ (Sawaki, 1989; Ringbom and Shy, 2002; and Somboon and Amaruchkul, 2016). Tasman Empire Airways demonstrated that using the binomial distribution is a suitable model for the show-up distribution (Thompson, 1961). Assume that r_i is a refund for no-show Class- i reservation for $i = 1, 2$, where $g_i \in (0, 1)$; $r_i = g_i p_i$, which g_i is a proportion of the revenue cost.

Some passengers are denied at the departure if the number of show-up passengers exceeds capacity. As mentioned above that we do not overbook Class-1 customers, so all denied boarding passengers are Class-2. The compensation that the airline has to pay all passengers who were denied boarding is h , where $h > p_2$. However, it could include a higher fare class ticket for the next flight, cash vouchers, or hotel accommodations.

An optimal booking limit at the update booking limit point $x^*(t)$ that maximizes its expected profit is preferred:

$$p(x(t)) = E \sum_{i=1}^{\kappa-2} [p_i B_i(x(t)) - r_i (B_i(x(t)) - W_i(B_i(x(t))))] - E \sum_{i=\kappa}^{\kappa+1} (W_2(B_2(x(t))) - k(t))^+ - E \sum_{i=\kappa+1}^{\kappa+2} (D_2(t) - x(t))^+ + g_1(D_1(t) - (k(t) - B_2(x(t))))^+.$$

From the equation above, it can be seen that there are two uncertainty parts which are demand and the number of show-up customers. However, we are capable find a closed form of the optimal booking limit. Finding a closed form of the optimal booking limit can use the theorem below (Somboon and Amaruchkul, 2016).

Theorem 1. For $x = 0, 1, \dots, \kappa - 2$ and $x = \kappa, \kappa + 1, \dots$, the expected profit function $\pi(x)$ is piecewise is unimodal in each piece. The expected profit $\pi(x)$ has a local maximum point x' on $x = 0, 1, \dots, \kappa - 2$ as

$$x' = \begin{cases} 0 & ; 0 \leq \tau < P(D_1 > \kappa - 1) \\ \kappa - F_{D_1}^{-1}(1 - \tau) & ; P(D_1 > \kappa - 1) \leq \tau \leq P(D_1 > 0) \\ \kappa - 2 & ; P(D_1 > 0) < \tau \leq 1 \end{cases} \quad (1)$$

On the other hand, if $0 < \alpha_2 / (h\theta_2) < \bar{F}(\kappa - 1; \kappa, \theta_2)$, then the expected profit $\pi(x)$ has a local maximum point x'' for $x = \kappa, \kappa + 1, \dots$. It can be written as

$$x'' = \arg \min \{x \in \{\kappa, \kappa + 1, \dots\} : \bar{F}(\kappa - 1; \kappa, \theta_2) > \frac{\alpha_2}{h\theta_2}\}. \quad (2)$$

Otherwise, the expected profit function is increasing.

3. Numerical Illustration

As mentioned in the introduction, accurate demand forecasting is a key to gauging optimal booking and overbooking. It is cumbersome to find exact forecasting demand since the historical data, which is the censor data, do not reflect the actual demand. That is, the booking limit determines how many seats can be sold on a flight. A booking in a fare class is accepted by an airline until the booking limit is reached. The airline then stops selling seats in that fare class. It also ceases to collect valuable data. Demand for travel in that fare class may exceed the booking limit, but the data does not reflect this, as the booking limit is censored or "constrained". From this cause, true demand cannot obtain directly, which can only obtain estimated demand. In order to obtain estimated demand, the unconstraining method is used to calculate the estimated demand. In quantity-based RM, the Expectation – Maximization (EM) method is the most commonly used for correcting for constrained data with four basic steps: (i) replace missing values with estimated values, (ii) estimate parameters, (iii) re-estimate the missing values assuming the new parameter estimates are correct, (iv) re-estimate the parameters, and iterating until convergence (Talluri et al, 2004: 474-475).

Like Somboon and Amaruchkul (2016), we assume that the demand of seats for each class in a year follows Poisson distribution. The Poisson variable is assumed in numerical experiments; however, the proof in Theorem 1 does not need to assume Poisson distribution but holds for any non-negative random variable.

In this experiment, we generate non-stationary booking demand by dividing the booking demand into four quarters. Let λ_1 and λ_2 be initial parameters referring the seat demand of customers class- i ; $i = 1, 2$ increasing 5 percent each quarter. We assume that $\lambda_1 = 40, 50, 60, 70$, and $\lambda_2 = 90, 100, 110, 120$, respectively. The plane's capacity is given by $\kappa = 162$ and the fare for class- i ; $i = 1, 2$ is denoted by $p_1 = 3,043$ and $p_2 = 945$, respectively. r_i is the refund for

passenger in each class where $r_i = 0.5p_i$ and $r_i = 0.8p_i$. The show-up probability θ_i for customer class- i ; $i = 1, 2$ are 0.7 and 0.9, respectively. Also, assume that the compensation cost $h = 2,000$ Baht must be paid to all passenger when the number of passengers exceeds capacity.

The EM method is used to estimate the censored demand data then uncensored it using the unconstraining method. In this setting, we partition the demand into two classes by employing the demand function proposed by Komsan Suriya (2009) for airline industry in Thailand. Given that is the demand function.

Class-1 demand function

$$p_1 = 3,588 - 188.605q_1. \quad (3)$$

Class-2 demand function

$$p_2 = 1,763 - 188.605q_2. \quad (4)$$

where q_i is the demand for class- i , for $i = 1, 2$.

After substituting $p_1 = 3,043$ in (3) and $p_2 = 945$ in (4), we obtain the demand for Class-1 and Class-2 as, $q_1 = 4.6$ and $q_2 = 6.9$ million customers, respectively. The proportion of Class-1 (ν_1) and Class-2 (ν_2) passengers is

$$\nu_1 = 0.4 \text{ and } \nu_2 = 0.6. \quad (5)$$

We considered the proportions in (5) to divide the number of reservations into two classes.

Next, the booking and overbooking limit is gauged using the Two-Class overbooking model, then calculated the expected profit from the booking and overbooking limit. Repeat 1,000 iterations for all the steps and average the profit. The average profit between the Two-Class overbooking model and the airline's policy are compared.

4. Numerical Result

In this experiment, we simulate non-stationary booking demand detailed in the section 3. Then, we estimate censored booking demand data using EM method, and the Two-Class overbooking model is used to find the optimal overbooking for non-stationary demand. The expected profit in each situation can be summarized below.

4.1 The optimal overbooking and the expected profit of the Two-Class overbooking model for non-stationary booking demand.

Table 1 shows that the optimal booking demand when $\lambda_1 = 40$ and $\lambda_2 = 90$ for all r_i and θ_i is the booking limit. Said that the airlines should accept Class-2 customers to make bookings as booking limit to achieve the highest expected profit.

Likewise, the optimal booking limit for $\lambda_1 = 50$ and $\lambda_2 = 100$ for all r_i and θ_i is the booking limit. The airlines have to open the reservation for Class-2 customers as booking limit to exceed the highest expected profit. The result as shown in Table 2.

Table 1 The optimal overbooking and the expected profit when $\lambda_1 = 40$ and $\lambda_2 = 90$

r_1	r_2	θ_1	θ_2	Optimal booking limit	Expected profit
1521.5	472.5	0.7	0.7	103	209,278.85
1521.5	472.5	0.7	0.9	104	217,153.85
1521.5	472.5	0.9	0.7	103	226,188.63
1521.5	472.5	0.9	0.9	103	234,063.63

Table 1 The optimal overbooking and the expected profit when $\lambda_1=40$ and $\lambda_2=90$ (Cont.)

r_1	r_2	θ_1	θ_2	Optimal booking limit	Expected profit
1521.5	756	0.7	0.7	103	202,201.59
1521.5	756	0.7	0.9	104	214,801.59
1521.5	756	0.9	0.7	102	219,111.37
1521.5	756	0.9	0.9	103	231,711.37
2434.4	472.5	0.7	0.7	104	194,299.68
2434.4	472.5	0.7	0.9	105	202,174.68
2434.4	472.5	0.9	0.7	103	221,355.33
2434.4	472.5	0.9	0.9	104	229,230.33
2434.4	756	0.7	0.7	103	187,222.42
2434.4	756	0.7	0.9	105	199,822.42
2434.4	756	0.9	0.7	102	214,278.07
2434.4	756	0.9	0.9	103	226,878.07

While the result in both Table 1 and Table 2 are booking limit, the results in Table 3 show that the optimal booking limit for $\lambda_1 = 60$ and $\lambda_2 = 110$ for all r_i and θ_i also booking limit almost all the situations. There is one situation when $r_1 = 2,434.4$, $r_2 = 472.5$, $\theta_1 = 0.9$ and $\theta_2 = 0.7$ that is the overbooking. As a result, to reach the highest expected profit, the airline should open booking for Class-2 customers to book as same as optimal overbooking limit.

Table 2 The optimal overbooking and the expected profit when $\lambda_1=50$ and $\lambda_2=100$

r_1	r_2	θ_1	θ_2	Optimal booking limit	Expected profit
1521.5	472.5	0.7	0.7	94	238,360.92
1521.5	472.5	0.7	0.9	95	247,094.70
1521.5	472.5	0.9	0.7	93	258,224.33
1521.5	472.5	0.9	0.9	94	266,985.67
1521.5	756	0.7	0.7	93	231,836.35
1521.5	756	0.7	0.9	95	245,376.37
1521.5	756	0.9	0.7	93	251,370.72
1521.5	756	0.9	0.9	94	265,418.54
2434.4	472.5	0.7	0.7	95	221,204.43
2434.4	472.5	0.7	0.9	96	230,281.05
2434.4	472.5	0.9	0.7	94	252,450.86
2434.4	472.5	0.9	0.9	94	261,541.23
2434.4	756	0.7	0.7	94	214,174.52
2434.4	756	0.7	0.9	96	228,424.91
2434.4	756	0.9	0.7	93	245,561.12
2434.4	756	0.9	0.9	94	259,974.11

Table 3 The optimal overbooking and the expected profit when $\lambda_1 = 60$ and $\lambda_2 = 110$

r_1	r_2	θ_1	θ_2	Optimal booking limit	Expected profit
1521.5	472.5	0.7	0.7	88	257,480.99
1521.5	472.5	0.7	0.9	89	267,740.80
1521.5	472.5	0.9	0.7	87	278,028.95
1521.5	472.5	0.9	0.9	88	288,288.76
1521.5	756	0.7	0.7	87	253,252.38
1521.5	756	0.7	0.9	89	269,292.18
1521.5	756	0.9	0.7	86	274,082.26
1521.5	756	0.9	0.9	88	290,122.06
2434.4	472.5	0.7	0.7	89	239,066.52
2434.4	472.5	0.7	0.9	90	249,321.08
2434.4	472.5	0.9	0.7	233	271,935.23
2434.4	472.5	0.9	0.9	88	282,015.11
2434.4	756	0.7	0.7	88	234,555.98
2434.4	756	0.7	0.9	90	250,593.68
2434.4	756	0.9	0.7	87	267,996.56
2434.4	756	0.9	0.9	88	283,848.41

Table 4 The optimal overbooking and the expected profit when $\lambda_1 = 70$ and $\lambda_2 = 120$

r_1	r_2	θ_1	θ_2	Optimal booking limit	Expected profit
1521.5	472.5	0.7	0.7	88	257,397.69
1521.5	472.5	0.7	0.9	89	267,316.26
1521.5	472.5	0.9	0.7	87	278,209.29
1521.5	472.5	0.9	0.9	88	288,127.85
1521.5	756	0.7	0.7	87	253,217.64
1521.5	756	0.7	0.9	89	268,709.34
1521.5	756	0.9	0.7	86	274,312.74
1521.5	756	0.9	0.9	88	289,804.44
2434.4	472.5	0.7	0.7	89	238,719.58
2434.4	472.5	0.7	0.9	90	248,638.14
2434.4	472.5	0.9	0.7	233	272,462.99
2434.4	472.5	0.9	0.9	180	281,909.05
2434.4	756	0.7	0.7	88	234,256.03
2434.4	756	0.7	0.9	90	249,747.73
2434.4	756	0.9	0.7	87	268,121.58
2434.4	756	0.9	0.9	88	283,424.28

However, Table 4 shows that the optimal booking limit for $\lambda_1 = 70$ and $\lambda_2 = 120$ have both booking limit and overbooking limit. In this case, if the airline wants to succeed the maximum expected profit, the airline have to adjust the optimal booking for all r_i and θ_i ; especially, in the situation that the booking limit is greater than the capacity. It means that the airline can only open booking for Class-2 customers.

The numerical study investigates that booking limit and overbooking limit provide the same expected profit in some situations. In this case, we recommend that the airline use a booking limit to sell both classes' flight tickets since the show-up probability of each class is uncertain in real situations. Generally, the show-up probability of Class-2 passengers is less than Class-1 passengers because of a cheaper flight fare, i.e., booking during promotion time. Class-2 passengers are most likely to terminate the fare without notice. Because of this reason, the airline should allow both Class-1 and Class-2 passengers to book flight tickets.

4.2 The comparison between the Two-class overbooking for both classes passengers to the airline's policy when the booking demand is non-stationary

For performance comparison, we assume that the expected profit of the airline's policy has a fixed booking limit for all periods, which are 9, 17, 41, 81, 122, and 171, which are 5%, 10%, 25%, 50%, 75%, and 105% of capacity, respectively. Otherwise, we investigate that the expected profit of the Two-Class overbooking model is greater than or equal to the airline's booking policy for all situations. The airline will get less advantage when the airline booking limit (ABL) is 9, followed by 17, 41, 81, 122, and 171, respectively.

Table 5 Loss Profit per Flight when $\lambda_1 = 40$ and $\lambda_2 = 90$

r_1	r_2	θ_1	θ_2	ABL	Loss Profit per Flight	Percentage of Loss Profit per Flight
2434.4	472.5	0.9	0.7	9	34,810.07	18.66
				17	31,030.07	16.30
				41	19,690.07	9.76
				81	1,903.07	0.87
				122	0.00	0.00
				171	0.00	0.00
1521.5	472.5	0.9	0.9	9	34,810.07	17.90
				17	31,030.07	15.66
				41	19,690.07	9.40
				81	1,903.07	0.84
				122	0.00	0.00
				171	0.00	0.00
2434.4	756	0.9	0.7	9	13,842.88	6.91
				17	12,330.88	6.11
				41	7,794.88	3.78
				81	680.08	0.32
				122	0.00	0.00
				171	0.00	0.00
2434.4	756	0.9	0.9	9	13,842.88	6.50
				17	12,330.88	5.75
				41	7,794.88	3.56
				81	680.08	0.30
				122	0.00	0.00
				171	0.00	0.00

Table 5 shows the loss profit per flight when $\lambda_1 = 40$ and $\lambda_2 = 90$ for all r_i and θ_i . It can be seen that the airline earns less advantage when the airline booking limit is 9, followed by 17, 41, 81, 122, and 171, respectively. Likewise, when $\lambda_1 = 50$ and $\lambda_2 = 100$ for all r_i , θ_i , and the airline booking limit, the airline losses its profit as shown in A.1 in Appendix (see others situation result in Appendix)

It concludes that if the airline open the observation for Class-2 customers in small amounts, the airline will get less profit. However, when the Two-Class overbooking model is applied, the model calculates the optimal booking limit suited for the situations. Consider the difference between the expected profit from the Two-Class overbooking model and the airline's booking policy; it can be seen that the differences in the profit are approximately 0-21%.

Practically, most airlines update their booking limit every month before departure for at least six months and then update every day in the last week before departure (Phillips, 2005). The following section will find the optimal number of update booking limits when real data demand is non – stationary.

5. Real Data Study

Passengers' information from two flights was received from the airline. The flights are from Bangkok to Phuket on every Sunday in 2014. Flight A is scheduled to depart for Sunday morning and Flight B is set for Sunday evening. The data includes both the number of reservations and the number of passengers who showed up at the departure time. Table 6 shows that the plane has 162 seats and that the airline's revenue is divided into eleven classes.

Table 6 Percentage of passengers from 2014 reclassified by revenue.

Class	Revenue	Percent ^a
Y	4,675	10.64
M	4,350	4.97
K	3,850	13.82
N	3,500	3.72
T	3,050	20.47
L	2,800	5.27
H	2,550	8.02
Q	2,100	4.47
V	1,900	9.23
G	1,690	12.58
B	945	6.71

Note: a. The Percentage of passengers in Table 6 was received from the airline company in Thailand where the name has not been disclosed to protect its and interviewees' privacy.

In this setting, the eleven classes were reclassified into two classes in order to use the Two-Class overbooking model. So, the revenue for Class-1 and Class-2 after rearrange is $p_1 = 3,043$ and $p_2 = 945$.

We established four different cases of updated booking limit points. Firstly, the booking limit is updated every day in the final week before departure; in this case, there is seven points update booking limit. Similarly, the booking limit is updated every month for the next six months and every day in the final week before departure. It can be assumed that there are thirteen points for updating the booking limit. The booking limit is then updated every three months before departure more than six months in advance, and the updated booking limit is considered the second case, giving this case fourteen updating points. Finally, we consider that the booking

limit is updated every month and every day in the last week before departure, which can be counted as sixteen points to update the booking limit. Let b_i be the number of class- i booking at the departure time for $i = 1, 2$ and $W_i(b_i)$ be the number of class- i show-up customers. The profit can be written as

$$\hat{\pi} = \sum_{i=1}^2 [p_i b_i - r_i (b_i - W_i(b_i))] - h(W_2(b_2) - \kappa)^+ . \quad (6)$$

Assume that t is the number of days before departure and x_t^* be the optimal booking limit at the updated booking limit points t before departure. The updated booking limit points are shown in Table 7.

Table 7 The Update Booking Limit Points for the Four Cases

Case	t
7 points	6, 5, 4, 3, 2, 1, 0
13 points	180, 150, 120, 90, 60, 30, 6, 5, 4, 3, 2, 1, 0
14 points	270, 180, 150, 120, 90, 60, 30, 6, 5, 4, 3, 2, 1, 0
16 points	330, 300, 270, 180, 150, 120, 90, 60, 30, 6, 5, 4, 3, 2, 1, 0

From the previous section, the Two-Classes overbooking can be used with the non – stationary demand using simulation technique. In this section, the real data is used to find the optimal number of update booking limits follow this steps:

- 1) Using real data, forecast demand ($\hat{\lambda}$) for the next flight using an exponential smoothing technique with a smoothing constant determined by minimizing the sum of squared errors.
- 2) Use the proportion of Class-1 and Class-2 passengers in (5) to divide the forecasted demand into two classes. The average demand of Class-1 and Class-2 are $\hat{\lambda}_{1,360}$ and $\hat{\lambda}_{2,360}$.
- 3) Generate the non – stationary demand for both Class-1 and Class-2 over 360 days, which is assumed to be a Poisson distribution with means $\hat{\lambda}_{1,360}$ and $\hat{\lambda}_{2,360}$, respectively.
- 4) Compute the optimal initial booking limit x_{360}^* using Theorem 1.
- 5) Compute the number of every day Class-1 and Class-2 bookings using x_{360}^* .
- 6) At update booking limit point t , recomputed the optimal booking limit x_t^* and forecast the two classes' demand ($\hat{\lambda}_{1,t}, \hat{\lambda}_{2,t}$) over the remaining days before departure.
- 7) Compute a new optimal booking limit x_t^* using the forecasted demand
- 8) Generate the number of show-up customers following binomial distribution with a mean equal to the number of class- i bookings and the show-up probability is θ_i , $i = 1, 2$, then compute the profit using (6).
- 9) Repeat 1,000 iterations for all steps and average the profit.

In this experiment, the best solution of the optimal update booking limit point is 16, which gives the profit approximately 32% – 42% greater than 14 update booking limit points. Moreover, 7 and 13 points provide the same profit for all the situations. The optimal update booking limit of 13 point gives the highest profit in stationary demand situations, while the update

booking limit of 16 point is more suitable in non-stationary demand situations. The result is shown in table 8.

Table 8 Profit of the Number of Update Booking Limit Points in different cases

θ_1	θ_2	7 points	13 points	14 points	16 points
0.7	0.7	234,015.76	234,015.76	236,970.50	323,450.99
0.7	0.8	238,740.76	238,740.76	241,448.38	325,017.75
0.7	0.9	241,859.89	241,859.89	244,772.03	326,587.10
0.7	0.95	244,222.39	244,222.39	246,992.78	327,351.43
0.8	0.7	251,067.90	251,067.90	254,572.05	354,704.28
0.8	0.8	255,822.95	255,822.95	259,015.01	356,275.45
0.8	0.9	260,547.95	260,547.95	263,492.42	357,835.87
0.8	0.95	262,437.95	262,437.95	265,458.02	358,598.11
0.9	0.7	268,274.24	268,274.24	272,197.11	385,984.22
0.9	0.8	272,526.74	272,526.74	276,399.52	387,536.86
0.9	0.9	277,251.74	277,251.74	280,876.93	389,096.20
0.9	0.95	277,874.48	277,874.48	281,841.54	389,847.83
0.95	0.7	279,104.43	279,104.43	282,587.31	401,629.10
0.95	0.8	281,928.39	281,928.39	285,674.59	403,193.53
0.95	0.9	286,653.39	286,653.39	290,152.00	404,747.25
0.95	0.95	288,543.39	288,543.39	292,117.60	405,500.98

6. Conclusion

Simulating the non-stationary booking demand data to calculate the optimal booking limit using the Two-Class overbooking model finds that the optimal booking calculated mostly is the booking limit. It means that seats are open to reserve for Class-2 customers as booking limit, while the rest of the seats were reserved for Class-1 customers. Otherwise, the booking limit is suited to the overbooking limit in some situations, that is, the open to Class-2 customers overbooking (in which case, there will be no seats left for Class-1 customers).

For comparison, the expected profit from the optimal overbooking limit is greater than or equal to the airline's booking policy in all study situations. The Two-Class overbooking model yielded up to 21% more the expected profit margins than airline booking policies. Using real data, the Two-Class overbooking model is then taken to gauge the optimal number of update booking limits when the demand is non – stationary. It finds that the update booking limit 16 point gives the maximum profit.

Acknowledgement

This research was funded by Faculty of Applied Science, King Mongkut's University of Technology North Bangkok Contract no. 6345104.

References

- Cross, R. G. (1997). *Revenue Management : Hard-Core Tactics for Market Domination*. Crown Business.
- Littlewood, K. (1972). Special issue papers: Forecasting and control of passenger bookings. *Journal of Revenue and Pricing Management*, 4(2), 111-123.
- Phillips, R. L. (2005). *Pricing and revenue optimization*. Stanford University Press.
- Ringbom, S., & Shy, O. (2002). The adjustable-curtain strategy : Overbooking of multiclass service. *Journal of Economics*, 77(1), 73-90.
- Sawaki, K. (1989). An analysis of airline seat allocation. *Journal of Operations Research Society of Japan*, 32(4), 411- 419.
- Somboon, M., & Amaruchkul, K. (2016). Combined overbooking and seat inventory control for two-class revenue management model. *Songklanakarin Journal of Science & Technology*, 38(6), 657-665.
- Suriya, K. (2009). The Impact of Low Cost Airlines to Airline Industry : An Experience of Thailand. *Jurnal Ekonomi Malaysia*, 43, 3-25.
- Suzuki, Y. (2006). The net benefit of airline overbooking. *Transportation Research Part E: Logistics and Transportation Review*, 42(1), 1-19. <https://doi.org/10.1016/j.tre.2004.09.001>
- Talluri, K. T., Van Ryzin, G., & Van Ryzin, G. (2004). *The theory and practice of revenue management*. Kluwer Academic Publishers.
- Thompson, H. R. (1961). Statistical problems in airline reservation control. *Journal of the Operational Research Society*, 12(3), 167-185. <https://doi.org/10.2307/3006774>

Appendix

Table A.1 Loss Profit per Flight when $\lambda_1 = 50$ and $\lambda_2 = 100$

r_1	r_2	θ_1	θ_2	ABL	Loss Profit per Flight	Percentage of Loss Profit per Flight
2434.4	472.5	0.9	0.7	9	38,647.25	18.08
				17	34,867.25	16.02
				41	23,527.25	10.28
				81	4,681.06	1.89
				122	361.33	0.14
				171	361.33	0.14
1521.5	472.5	0.9	0.9	9	38,647.25	17.34
				17	34,867.25	15.38
				41	23,527.25	9.88
				81	4,681.06	1.82
				122	361.33	0.14
				171	361.33	0.14
2434.4	756	0.9	0.7	9	15,444.45	6.71
				17	13,932.45	6.01
				41	9,396.45	3.98
				81	1,857.97	0.76
				122	1,590.72	0.65
				171	1,590.72	0.65

Table A.1 Loss Profit per Flight when $\lambda_1 = 50$ and $\lambda_2 = 100$ (Cont.)

r_1	r_2	θ_1	θ_2	ABL	Loss Profit per Flight	Percentage of Loss Profit per Flight
2434.4	756	0.9	0.9	9	15,312.84	6.26
				17	13,800.84	5.61
				41	9,264.83	3.70
				81	1,726.36	0.67
				122	1,459.11	0.56
				171	1,459.11	0.56

Table A.2 Loss Profit per Flight when $\lambda_1 = 60$ and $\lambda_2 = 110$

r_1	r_2	θ_1	θ_2	ABL	Loss Profit per Flight	Percentage of Loss Profit per Flight
2434.4	472.5	0.9	0.7	9	37,029.68	15.76
				17	33,249.68	13.93
				41	21,909.68	8.76
				81	3,009.68	1.12
				122	0.00	0.00
				171	0.00	0.00
1521.5	472.5	0.9	0.9	9	37,319.63	15.25
				17	33,539.63	13.50
				41	22,199.63	8.54
				81	3,299.63	1.18
				122	289.94	0.10
				171	289.94	0.10
2434.4	756	0.9	0.7	9	14,739.90	5.82
				17	13,227.90	5.19
				41	8,691.90	3.35
				81	1,131.90	0.42
				122	4,991.58	1.90
				171	4,991.58	1.90
2434.4	756	0.9	0.9	9	14,927.85	5.55
				17	13,415.85	4.96
				41	8,879.85	3.23
				81	1,319.85	0.47
				122	5,179.53	1.86
				171	5,179.53	1.86

Table A.3 Loss Profit per Flight when $\lambda_1 = 70$ and $\lambda_2 = 120$

r_1	r_2	θ_1	θ_2	ABL	Loss Profit per Flight	Percentage of Loss Profit per Flight
2434.4	472.5	0.9	0.7	9	37,488.86	15.95
				17	33,708.86	14.12
				41	22,368.86	8.94
				81	3,468.86	1.29
				122	0.00	0.00
				171	0.00	0.00
1521.5	472.5	0.9	0.9	9	37,488.86	15.34
				17	33,708.86	13.58
				41	22,368.86	8.62
				81	3,468.86	1.25
				122	0.00	0.00
				171	0.00	0.00
2434.4	756	0.9	0.7	9	14,742.00	5.82
				17	13,230.00	5.19
				41	8,694.00	3.35
				81	1,134.00	0.42
				122	4,440.79	1.68
				171	4,440.79	1.68
2434.4	756	0.9	0.9	9	14,931.00	5.56
				17	13,419.00	4.97
				41	8,883.00	3.24
				81	1,323.00	0.47
				122	4,629.79	1.66
				171	4,629.79	1.66

Research Article

การเปรียบเทียบตัวแบบจำแนกเสียงเปิดวาล์วของอุปกรณ์ดักไอน้ำ

Comparison of classification model for steam trap valve opening sound

ปณณนุช ดำนงค์^{1*} พิมพภา ตานินพงศ์¹ อนุชา พรหมวังขวา² และ จักรเมธ บุตรกระจำง³

Punyanut Damnong^{1*} Phimphaka Taninpong¹ Anucha Promwungkwa and Jakramate Bootkrajang³

¹ภาควิชาสถิติ คณะวิทยาศาสตร์ มหาวิทยาลัยเชียงใหม่ 239 ถ.ห้วยแก้ว ต.สุเทพ อ.เมือง เชียงใหม่ 50200 ประเทศไทย

²ภาควิชาวิศวกรรมเครื่องกล คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเชียงใหม่ 239 ถ.ห้วยแก้ว ต.สุเทพ อ.เมือง เชียงใหม่ 50200 ประเทศไทย

³ภาควิชาวิทยาการคอมพิวเตอร์ คณะวิทยาศาสตร์ มหาวิทยาลัยเชียงใหม่ 239 ถ.ห้วยแก้ว ต.สุเทพ อ.เมือง เชียงใหม่ 50200 ประเทศไทย

¹Department of Statistics, Faculty of Science, Chiang Mai University, 239 Huay Kaew Road, Suthep, Muang district, Chiang Mai 50200 Thailand

²Department of Mechanical Engineering, Faculty of Engineering, Chiang Mai University, 239 Huay Kaew Road, Suthep, Muang district, Chiang Mai 50200 Thailand

³Department of Computer Science, Faculty of Science, Chiang Mai University, 239 Huay Kaew Road, Suthep, Muang district, Chiang Mai 50200 Thailand

*E-mail: punyanutdamnong@gmail.com

Received: 29/06/2021; Revised: 30/10/2021; Accepted: 16/11/2021

บทคัดย่อ

การวิจัยนี้มีวัตถุประสงค์เพื่อสร้างตัวแบบการจำแนกเสียงเปิดวาล์วของอุปกรณ์ดักไอน้ำและเปรียบเทียบประสิทธิภาพของตัวแบบที่ใช้วิธีซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) และวิธีหน่วยความจำระยะสั้นแบบยาว (Long Short-Term Memory: LSTM) ทำการบันทึกเสียงการทำงานของอุปกรณ์ดักไอน้ำ และแยกคุณลักษณะของเสียง โดยใช้วิธีอัตราค่าตัดศูนย์ (zero crossing rate) วิธี spectral centroid วิธี spectral rolloff วิธีสัมประสิทธิ์เซปสตรัมบนความถี่แบบเมล (mel-frequency cepstral coefficient) และการแปลงฟูเรียร์ระยะสั้น (short-time fourier transform) นอกจากนี้ยังทำการเปรียบเทียบประสิทธิภาพของตัวแบบโดยใช้ชุดข้อมูล 2 ชุด ได้แก่ ชุดข้อมูลที่ข้อมูลไม่สมดุล กับชุดข้อมูลสมดุลที่ทำการแก้ไขปัญหาคือข้อมูลไม่สมดุลด้วยวิธีผสมผสาน (hybrid sampling) ผลการวิจัยพบว่าตัวแบบ SVM เมื่อใช้ Kernel function เป็น polynomial 3rd degree และตัวแบบ LSTM ที่

เรียนรู้จากชุดข้อมูลที่สมดุลให้ค่า F1 สูงกว่าการเรียนรู้จากชุดข้อมูลที่ไม่สมดุล โดยมีค่า F1 เท่ากันเท่ากับ 66.67% แต่ตัวแบบ SVM มีค่าความเที่ยงสูงกว่าตัวแบบ LSTM ซึ่งมีค่าเท่ากับ 63.64% และ 52.94% ในขณะที่ตัวแบบ LSTM มีค่าการเรียกคืนสูงกว่าตัวแบบ SVM ซึ่งมีค่าเท่ากับ 90.00% และ 70.00% ตามลำดับ

คำสำคัญ: การจำแนกเสียง, อุปกรณ์ดักไอน้ำ, ซัพพอร์ทเวกเตอร์แมทรีน, หน่วยความจำระยะสั้นแบบยาว

Abstract

This research aims to create steam trap valve opening sound classification models using two classification methods including support vector machine (SVM) and long short-term memory (LSTM), and to compare the performance of the models. This study employs five feature extraction methods including zero-crossing rate, spectral centroid, Mel-frequency cepstral coefficient, spectral rolloff, and short-term Fourier transform. In addition, this study compares the performance of the models using two datasets including imbalanced dataset and balanced dataset which is resolved by using hybrid sampling method. The results show that SVM with Polynomial kernel function and LSTM provide higher F1 score when learning from balanced dataset than using imbalanced dataset. Moreover, SVM and LSTM provide the same F1 score of 66.67%. However, SVM provides higher precision than LSTM with value of 63.64% and 52.94%. In addition, LSTM gives higher recall than SVM with value of 90.00% and 70.00%, respectively.

Keywords: sound classification, steam trap, support vector machine, long short-term memory

บทนำ

อุปกรณ์ดักไอน้ำ (steam trap) เป็นวาล์วอัตโนมัติที่ช่วยในการระบายไอน้ำควบแน่น (condensate) ติดตั้งในหม้อไอน้ำซึ่งใช้ในกระบวนการของการผลิตในโรงงานอุตสาหกรรมบางประเภท หม้อไอน้ำเป็นหนึ่งในอุปกรณ์ที่สำคัญของระบบไอน้ำ ระบบไอน้ำเป็นระบบที่ผลิตไอน้ำประกอบด้วยอุปกรณ์และระบบย่อยต่าง ๆ ได้แก่ หม้อไอน้ำ ระบบส่งจ่ายไอน้ำ และระบบนำกลับไอน้ำควบแน่น (Ministry of Energy, 2010) ซึ่งระบบนำกลับไอน้ำควบแน่นนั้นจำเป็นต้องระบายไอน้ำที่เหลือจากการผลิตและไอน้ำควบแน่นให้ทัน หากระบายไอน้ำควบแน่นออกไม่ทันจะส่งผลให้มีไอน้ำควบแน่นค้างในระบบไอน้ำ ซึ่งอาจส่งผลต่อคุณภาพของผลิตภัณฑ์ ดังนั้นอุปกรณ์ดักไอน้ำ จึงช่วยในการระบายไอน้ำควบแน่นโดยสามารถปรับปริมาณการระบายไอน้ำควบแน่นได้ตามปริมาณไอน้ำควบแน่นที่เกิดขึ้นจริง (Department of Industrial Works, 2010) นอกจากนี้ยังทำหน้าที่ป้องกันการรั่วซึมของน้ำ และระบายก๊าซและอากาศออกจากระบบโดยไม่สูญเสียไอน้ำ หากอุปกรณ์ดักไอน้ำทำงานไม่มีประสิทธิภาพจะส่งผลต่อระบบไอน้ำ

น้ำอย่างมาก (Ministry of Energy, 2010) ดังนั้นจึงจำเป็นที่จะต้องทำการตรวจสอบคุณภาพในการทำงานของอุปกรณ์ดักไอน้ำอย่างสม่ำเสมอ วิธีการตรวจสอบคุณภาพการทำงานของอุปกรณ์ดักไอน้ำมีหลายวิธี เช่น การทดสอบด้วยการใช้น้ำรดอุปกรณ์ดักไอน้ำ การใช้เครื่องวัดอุณหภูมิทดสอบอุปกรณ์ดักไอน้ำ การเผาดูการระบายของอุปกรณ์ดักไอน้ำ การใช้เครื่องทดสอบอัลตราโซนิกทดสอบอุปกรณ์ดักไอน้ำ และการใช้เครื่องตรวจวัดอุปกรณ์ดักไอน้ำ (TM5) ข้อดีของการใช้เครื่องตรวจวัดอุปกรณ์ดักไอน้ำ คือ สามารถตรวจวัดและประมวลผลได้อย่างถูกต้องแม่นยำ เชื่อถือได้ และสามารถทำซ้ำได้ (Ministry of Energy, 2011) แต่ข้อเสียคือ มีค่าใช้จ่ายสูง ดังนั้นวิธีที่มีค่าใช้จ่ายต่ำกว่าคือ ตรวจสอบการทำงานของอุปกรณ์ดักไอน้ำโดยการฟังเสียงการเปิดวาล์วเพื่อระบายไอน้ำความดันของอุปกรณ์ดักไอน้ำ ซึ่งการตรวจจับเสียงการเปิดวาล์วของอุปกรณ์ดักไอน้ำโดยใช้เทคนิคการรู้จำเสียงอาจเป็นทางเลือกหนึ่งในการตรวจสอบคุณภาพของอุปกรณ์ดักไอน้ำแทนการซื้อเครื่องตรวจวัดอุปกรณ์ดักไอน้ำ

จากงานวิจัยที่เกี่ยวข้องกับการจำแนกเสียงหรือการตรวจจับความผิดปกติของสัญญาณเสียง พบว่าตัวแบบที่นิยมใช้มีหลายวิธี ได้แก่ วิธี Support Vector Machine (SVM) (Kaewtai, 2009; Giannakopoulos, 2015) วิธี k-Nearest Neighbor (k-NN) (Kaewtai, 2009; Giannakopoulos, 2015) วิธีโครงข่ายประสาทเทียม (Artificial Neural Network: ANN) (Kaewtai, 2009; Galván-Tejada, et al., 2016) วิธี Hidden Markov Model (HMM) (Giannakopoulos, 2015) วิธี Random Forest (RF) (Galván-Tejada, et al., 2016) วิธีโครงข่ายประสาทเทียมแบบวนกลับ (Recurrent Neural Network: RNN) (Lim et al., 2019) และ วิธี Extreme Gradient Boosting (XGBoost) (Li et al., 2018) โดยพบว่าวิธี SVM จำแนกเพลงไทยเดิมได้ดีที่สุดให้ค่าความถูกต้องเท่ากับ 84.95% และสามารถจำแนกเสียงเพลงกับเสียงพูดได้ค่าความถูกต้องถึง 94.6% (Kaewtai, 2009; Giannakopoulos, 2015) สำหรับวิธี RNN สามารถจำแนกเสียงกรนได้ค่าความถูกต้องสูงเช่นเดียวกัน โดยมีค่าความถูกต้องถึง 98.9% (Lim et al., 2019) นอกจากนี้ในงานวิจัยที่มีการเปรียบเทียบประสิทธิภาพของวิธี SVM และ Naïve Bayes สำหรับการจำแนกกลุ่มข้อความ (Hassan et al., 2012) พบว่าตัวแบบ SVM ให้ค่า micro-average and macro-average F-measure สูงกว่าตัวแบบ Naïve Bayes ดังนั้นงานวิจัยนี้จึงสนใจเปรียบเทียบประสิทธิภาพของวิธี SVM และ LSTM ซึ่ง LSTM เป็นตัวแบบ RNN ชนิดพิเศษที่สามารถเรียนรู้อนุกรมเวลาระยะยาวในการจำแนกเสียงเปิดวาล์วของอุปกรณ์ดักไอน้ำ และในการสร้างตัวแบบการจำแนกเสียงจำเป็นต้องทำการแยกคุณลักษณะของเสียง (feature extraction) จากเสียงที่บันทึกก่อน ซึ่งวิธีแยกคุณลักษณะของเสียงที่นิยมใช้ได้แก่ spectral centroid (Kaewtai, 2009; Giannakopoulos, 2015) อัตราค่าตัดศูนย์ (zero crossing rate: ZCR) (Kaewtai, 2009; Giannakopoulos, 2015; Lim et al., 2019) สัมประสิทธิ์เซปสตรัมบนความถี่แบบเมล (mel-frequency cepstral coefficient: MFCC) (Kaewtai, 2009; Giannakopoulos, 2015; Galván-Tejada, et al., 2016; Lim et al., 2019) และการแปลงฟูเรียร์ระยะสั้น (short-time fourier transform: STFT) (Lim et al., 2019) นอกจากนี้ยังมีอีกหลายวิธี ได้แก่ ค่าสัมประสิทธิ์จากการประมาณค่าเชิงเส้น (linear predictive coefficient: LPC) (Kaewtai, 2009) power spectrum (Kaewtai, 2009) energy (Giannakopoulos, 2015) entropy of energy (Giannakopoulos, 2015) spectral spread (Giannakopoulos, 2015) spectral entropy (Giannakopoulos, 2015) spectral flux (Giannakopoulos,

2015) spectral rolloff (Giannakopoulos, 2015) chroma vector (Giannakopoulos, 2015) chroma deviation (Giannakopoulos, 2015) statistical features (Galván-Tejada, et al., 2016) และ wavelet packet energy (Li et al., 2018)

งานวิจัยนี้มีวัตถุประสงค์เพื่อสร้างตัวแบบจำแนกเสียงเปิดวาล์วของอุปกรณ์คักไอน้ำจากเครื่องกำเนิดไอน้ำ และเปรียบเทียบประสิทธิภาพของตัวแบบของวิธี SVM กับวิธีหน่วยความจำระยะสั้นแบบยาว (long short-term memory: LSTM) โดยใช้เกณฑ์ค่า F1 (F1-score) ที่สูงที่สุด นอกจากนี้งานวิจัยนี้ทำการเปรียบเทียบประสิทธิภาพของตัวแบบทั้งสองโดยใช้ชุดข้อมูล 2 ชุด ได้แก่ ชุดข้อมูลที่ข้อมูลเสียงที่ไม่สมดุลซึ่งเป็นข้อมูลที่มีเสียงเปิดวาล์วน้อยกว่าเสียงที่ไม่ใช่เสียงเปิดวาล์ว กับชุดข้อมูลสมดุลที่ทำการแก้ไขปัญหาค่าข้อมูลเสียงที่ไม่สมดุลด้วยวิธีผสมผสาน (hybrid sampling) เพื่อหาตัวแบบที่เหมาะสมที่สุดสำหรับจำแนกเสียงเปิดวาล์วของอุปกรณ์คักไอน้ำ

ระเบียบวิธีวิจัย

1. ข้อมูลและการเก็บรวบรวมข้อมูล

ข้อมูลเสียงที่ใช้ในการศึกษาได้จากการบันทึกเสียงการทำงานของอุปกรณ์คักไอน้ำ โดยสร้างท่อใส่อุปกรณ์บันทึกเสียงและนำไปติดกับท่อที่ปล่อยคอนเดนเสท อากาศ และก๊าซ ออกจากระบบไอน้ำซึ่งเป็นบริเวณที่อยู่ใกล้กับอุปกรณ์คักไอน้ำ แล้วทำการบันทึกเป็นระยะเวลา 1 วัน (24 ชั่วโมง)

2. การเตรียมข้อมูล

การเตรียมข้อมูลซึ่งเป็นสัญญาณเสียงความยาว 24 ชั่วโมงเพื่อใช้ในการสร้างตัวแบบการจำแนกเสียง มีกระบวนการดังนี้ 1) ทำการปรับลดค่า sampling rate ของเสียง ซึ่งก็คือ การสุ่มตัวอย่างสัญญาณอนาล็อกซึ่งเป็นสัญญาณไฟฟ้ามาแปลงเป็นสัญญาณดิจิทัลที่คอมพิวเตอร์เข้าใจ ซึ่งสัญญาณเสียงที่ได้จากการบันทึกมีค่า sampling rate เท่ากับ 44,100 Hz จะทำการปรับลดเหลือ 16,000 Hz 2) เปลี่ยนสัญญาณเสียงจากระบบสเตอริโอเป็นระบบโมโน และกรองค่าความถี่ให้อยู่ในช่วง 3,000 Hz ถึง 7,000 Hz 3) ทำการตัดแบ่งไฟล์เสียงออกเป็นไฟล์เสียงที่มีความยาวไฟล์เสียงละ 2.5 วินาที บันทึกไฟล์เสียงให้อยู่ในรูปแบบไฟล์สัญญาณเสียงนามสกุล wav (.wav) ได้จำนวนไฟล์ทั้งหมด 24,510 ไฟล์ 4) ทำการแยกคุณลักษณะของไฟล์เสียงที่ตัดแบ่งโดยใช้วิธีที่อธิบายในหัวข้อ 3. จะได้ชุดข้อมูลที่มีข้อมูลทั้งหมด 24,510 แถว ในจำนวนนี้เป็นเสียงเปิดวาล์วจำนวน 31 แถวและไม่ใช่เสียงเปิดวาล์วจำนวน 24,479 แถว 4) สุ่มตัวอย่างข้อมูลโดยใช้วิธีสุ่มตัวอย่างแบบชั้นภูมิ โดยสุ่มตัวอย่างจากข้อมูลเสียงเปิดวาล์วและข้อมูลที่ไม่ใช่เสียงเปิดวาล์ว คิดเป็นร้อยละ 70 ของจำนวนข้อมูลในแต่ละกลุ่ม โดยการสุ่มอย่างง่าย รวมกันได้ชุดข้อมูลสำหรับการเรียนรู้ (training set) ส่วนที่เหลืออีกร้อยละ 30 ของข้อมูลแต่ละกลุ่มจะเป็นชุดข้อมูลสำหรับการทดสอบ (testing set)

นำชุดข้อมูลสำหรับการเรียนรู้มาแบ่งออกเป็น 2 ส่วน ส่วนที่หนึ่งสุ่มตัวอย่างร้อยละ 80 ของข้อมูลชุดนี้เป็นชุดข้อมูลเรียนรู้ (training set) และที่เหลือร้อยละ 20 ของข้อมูลชุดนี้เป็นชุดข้อมูลตรวจสอบ (validation set) จำนวนข้อมูลของแต่ละชุดข้อมูลแสดงดังตารางที่ 1

ตารางที่ 1 จำนวนข้อมูลตัวอย่างของแต่ละชุดข้อมูล

กลุ่มข้อมูล	Training Set	Validation Set	Testing Set	รวม
เสียงเปิดวาล์ว	17	4	10	31
ไม่ใช่เสียงเปิดวาล์ว	13,709	3,428	7,342	24,479
รวม	13,726	3,432	7,352	24,510

3. การแยกคุณลักษณะของเสียง

งานวิจัยนี้ใช้วิธีการแยกคุณลักษณะของเสียงทั้งหมด 5 วิธี ได้แก่ ZCR spectral centroid spectral rolloff MFCC และ STFT โดยมีรายละเอียดของแต่ละวิธีดังนี้

3.1 อัตราค่าตัดศูนย์ (zero crossing rate: ZCR)

อัตราค่าตัดศูนย์ (ZCR) (Tzanetakis & Cook, 2002) คืออัตราการข้ามผ่านศูนย์ภายในเฟรม t กล่าวคือ อัตราค่าตัดศูนย์ของสัญญาณเสียง คือ อัตราของสัญญาณข้ามเส้นศูนย์แอมพลิจูดโดยเปลี่ยนจากค่าบวกเป็นค่าลบ หรือเปลี่ยนจากค่าลบเป็นค่าบวก อัตราค่าตัดศูนย์ของแต่ละเฟรมคำนวณโดยใช้สมการ (1)

$$ZCR_t = \sum_{i=1}^N \frac{|\text{sign}(x_i) - \text{sign}(x_{i-1})|}{2N} \quad (1)$$

โดยที่ $\text{sign}(x_i)$ คือ ฟังก์ชันที่ให้ค่าสัญญาณโดเมนเวลา x_i ซึ่งมีค่าเป็น 1 เมื่อ x_i อยู่เหนือเส้นสัญญาณศูนย์ หรือมีค่าเป็น -1 ถ้า x_i อยู่ต่ำกว่าเส้นสัญญาณศูนย์ นอกจากนี้มีค่าเป็น 0 ถ้า x_i อยู่บนเส้นสัญญาณศูนย์

N คือ ความยาวของสัญญาณภายในเฟรม t

i คือ หมายเลขช่วงความถี่ ซึ่งมีค่าอยู่ในช่วง 1 ถึง N

3.2 Spectral Centroid

Spectral Centroid (Tzanetakis & Cook, 2002) เป็นการวัดตำแหน่งและรูปร่างของสเปกตรัมอย่างง่าย ที่แสดงถึงค่าศูนย์กลางของสเปกตรัมจากการแปลงฟูเรียร์ระยะสั้น (Short-Time Fourier Transform) Spectral Centroid (C_t) ของเฟรม t คำนวณโดยใช้สมการ (2)

$$C_t = \frac{\sum_{i=1}^N M_t[i] \times i}{\sum_{i=1}^N M_t[i]} \quad (2)$$

โดยที่ $M_t[i]$ คือ การแปลงฟูเรียร์ที่เฟรม t และช่วงความถี่ที่ i

N คือ ความยาวของสัญญาณในเฟรม t

3.3 Spectral rolloff

Spectral rolloff คือ การวัดความเบ้ของรูปร่างสเปกตรัมซึ่งมีลักษณะเบ้ขวา สำหรับสเปกตรัมที่เบ้ขวาจะเป็นส่วนที่มีค่า spectral rolloff สูง กำหนดให้เป็นองค์ประกอบทางความถี่รอง ภายใต้อัตรา 85 เปอร์เซ็นต์ของการกระจายขนาดของสเปกตรัม (Tzanetakis & Cook, 2002; Burred & Lerch, 2004) คำนวณโดยใช้สมการ (3)

$$\sum_{i=1}^K |X_t[i]| \leq 0.85 \sum_{i=1}^{N/2} |X_t[i]| \quad (3)$$

โดยที่ $X_t[i]$ คือ แมกนิจูดของสเปกตรัมที่สอดคล้องกัน

i คือ หมายเลขช่วงความถี่

K คือ ความถี่ spectral rolloff ในเฟรม t

N คือ ความยาวของสัญญาณในเฟรม t

3.4 สัมประสิทธิ์เซปสตรัมบนความถี่แบบเมล (mel-frequency cepstral coefficient: MFCC)

เซปสตรัม (cepstral) เป็นการแปลงโคไซน์แบบไม่ต่อเนื่อง (discrete cosine transform) ของลอการิทึมจากสเปกตรัมสัญญาณในช่วงสั้น ๆ สัมประสิทธิ์เซปสตรัมบนความถี่แบบเมลเป็นเทคนิคที่ปรับปรุงจากเซปสตรัม ด้วยการปรับสเกลของสเปกตรัมให้อยู่บนสเกลที่เหมาะสมสำหรับการฟังของมนุษย์โดยสังเกตจากลักษณะของสัญญาณเสียง สัญญาณเสียงในช่วงความถี่ต่ำจะมีความสำคัญมากกว่าช่วงความถี่สูง จึงออกแบบสเกลของสเปกตรัมให้สามารถเก็บรายละเอียดของสัญญาณเสียงช่วงความถี่ต่ำได้มากกว่า เรียกการออกแบบนี้ว่าสเกลเมล (mel scale) โดยมีขั้นตอนในการคำนวณค่าสัมประสิทธิ์เซปสตรัมบนสเกลเมล (Junto, 2015; Kopanyapipat, 2015) ดังนี้

ในการคำนวณค่า MFCC จะเริ่มจาก 1) แปลงสัญญาณเสียงให้อยู่ในโดเมนของความถี่ โดยวิธีการแปลงฟูเรียร์แบบเร็ว (fast fourier transform: FFT) 2) คำนวณหาพลังงานสเปกตรัมที่ผ่านตัวกรอง (mel scale filtering) ขั้นตอนนี้นำความถี่ที่ได้จากการคำนวณสเปกตรัม (ค่า FFT) มาหาขนาดกำลังสอง ได้ $|x(k)|^2$ ส่งผ่านชุดตัวกรองแบบสามเหลี่ยมในสเกลเมล เพื่อเน้นความสำคัญของความถี่ที่อยู่ในช่วงกลางของชุดตัวกรองแต่ละตัวกรองตามสมการ (4) โดยที่ความถี่กลางของตัวกรองแต่ละชุดนั้นเกิดจากการแปลงค่าความถี่ปกติ (f) ให้อยู่บนสเกลเมล ($Mel(f)$) ตามสมการ (5) 3) การคำนวณสัมประสิทธิ์เซปสตรัมบนสเกลเมล (MFCC) จะนำลอการิทึม ($\log(.)$) ของพลังงานมาผ่านการแปลงโคไซน์แบบไม่ต่อเนื่อง (discrete cosine transform) ทำให้ได้ค่าสัมประสิทธิ์เซปสตรัมบนสเกลเมล c ลำดับที่ m ตามสมการ (6)

$$E_j = \sum_{k=0}^{n/2-1} \phi_j(k) |x(k)|^2; 0 \leq j \leq J \quad (4)$$

เมื่อ ϕ_j คือ ค่าประจำตัวกรองที่ j

$x(k)$ คือ สเปกตรัม

และ E_j คือ ค่าพลังงานสเปกตรัมที่ผ่านตัวกรอง

$$\text{Mel}(f) = 1125 \times \ln\left(1 + \frac{f}{700}\right) \quad (5)$$

โดยที่ $\text{Mel}(f)$ คือ ความถี่แบบเมล (Mel-Frequency)

และ f คือ ความถี่ปกติ หน่วยเป็น เฮิรตซ์ (Hz)

$$c_m = w_t(m) \sum_{j=1}^J \log_{10}(E_j) \cos\left(\frac{\pi}{J}(j - 0.5)m\right); m = 0, 1, \dots, J - 1 \quad (6)$$

$$\text{เมื่อ } w_t(m) = \begin{cases} \frac{1}{\sqrt{J}}, m = 0 \\ \frac{\sqrt{2}}{J}, 1 \leq m < J \end{cases}$$

โดยที่ $w_t(m)$ คือ ฟังก์ชันเงื่อนไข

j คือ อันดับของตัวกรอง

J คือ จำนวนตัวกรอง

3.5 การแปลงฟูเรียร์ระยะสั้น (short-time fourier transform: STFT)

การแปลงฟูเรียร์ระยะสั้น เป็นอีกหนึ่งวิธีที่มีประสิทธิภาพสำหรับการประมวลผลสัญญาณเสียง ซึ่งจะแปลงสัญญาณเสียงจากโดเมนเวลาไปเป็นโดเมนเวลาและความถี่ โดยอาศัยฟังก์ชันหน้าต่าง สามารถระบุแอมพลิจูดที่ซับซ้อนเทียบกับเวลา และความถี่ของสัญญาณเสียง (Colen, 1995) ในการคำนวณ STFT จะเริ่มต้นจากการแบ่งสัญญาณเสียงออกเป็นช่วงสั้น ๆ แล้วคูณแต่ละส่วนตามฟังก์ชันหน้าต่าง เนื่องจากการแปลงฟูเรียร์แบบไม่ต่อเนื่องทางเวลา (discrete-time fourier transform: DTFT) จะได้สูตรคำนวณ STFT (Ahmadizadeh, 2014) ดังสมการ (7)

$$X(\omega, n_0) = \sum_{n=-N}^{N-1} w[n]x[n + n_0]\exp(-j\omega n) \quad (7)$$

โดยที่ $(x[n_0 - N], \dots, x[n_0 + N - 1])^T$ คือ ส่วนของสัญญาณเสียง

$w[n]$ คือ ฟังก์ชันหน้าต่าง

และ $X(\omega, n_0)$ คือ STFT เป็นฟังก์ชันของทั้งความถี่ ω และเวลา n_0

4. การแก้ไขปัญหาข้อมูลไม่สมดุล

เนื่องจากข้อมูลเสียงในชุดข้อมูลเรียนรู้เป็นข้อมูลที่ไม่สมดุลในแต่ละกลุ่มจึงทำการแก้ปัญหาด้วยวิธีการปรับข้อมูลตัวอย่างแบบผสมผสาน (hybrid sampling) เป็นการนำเทคนิควิธีสุ่มเกิน (Oversampling) ซึ่งเป็นการสุ่มเพิ่มจำนวนข้อมูลในกลุ่มส่วนน้อยให้มีจำนวนใกล้เคียงหรือเท่ากับจำนวนของกลุ่มส่วนมาก และวิธีสุ่มลด (Undersampling) ซึ่งเป็นวิธีการสุ่มลดจำนวนข้อมูลในกลุ่มส่วนมากให้มีจำนวนใกล้เคียงหรือเท่ากับจำนวนในกลุ่มส่วนน้อย มาทำงานร่วมกัน (Kesornsit, et al., 2018) โดยมีขั้นตอนในการปรับข้อมูลตัวอย่างดังต่อไปนี้

- 1) หาค่ากลางของจำนวนข้อมูลเพื่อเป็นจำนวนในการสุ่มตัวอย่างโดยใช้จำนวนข้อมูลในกลุ่มส่วนมาก และจำนวนข้อมูลในกลุ่มส่วนน้อย ซึ่งสามารถคำนวณได้จากสมการที่ (8)

$$\text{center} = \frac{(\text{class}_0 + \text{class}_1)}{2} \quad (8)$$

เมื่อ class_0 คือ จำนวนข้อมูลของกลุ่มส่วนมาก ซึ่งเป็นกลุ่มข้อมูลเสียงที่ไม่ใช่เสียงเปิดวาล์ว

และ class_1 คือ จำนวนข้อมูลของกลุ่มส่วนน้อย ซึ่งเป็นกลุ่มข้อมูลเสียงเปิดวาล์ว

- 2) ใช้เทคนิควิธีการสุ่มเกินเพื่อเพิ่มจำนวนข้อมูลในกลุ่มเสียงเปิดวาล์วโดยใช้วิธีการสุ่มตัวอย่างแบบง่ายในการสุ่มข้อมูลเสียงเปิดวาล์วให้มีจำนวนเพิ่มขึ้นจนเท่ากับค่ากลางที่คำนวณได้จากข้อ 1
- 3) ใช้เทคนิควิธีการสุ่มลดเพื่อลดข้อมูลในกลุ่มเสียงที่ไม่ใช่เสียงเปิดวาล์วโดยใช้วิธีการสุ่มตัวอย่างแบบง่ายในการสุ่มข้อมูลที่ไม่ใช่เสียงเปิดวาล์วให้มีจำนวนเท่ากับค่ากลางที่คำนวณได้จาก 1
- 4) รวมข้อมูลที่ได้ในข้อ 2) และข้อ 3) ข้อมูลที่ได้จะเป็นข้อมูลที่มีจำนวนข้อมูลเสียงเปิดวาล์วและข้อมูลที่ไม่ใช่เสียงเปิดวาล์วเท่ากัน เพื่อใช้ในการสร้างตัวแบบและทดสอบตัวแบบต่อไป

5. เทคนิคการจำแนกเสียง

งานวิจัยนี้สร้างตัวแบบการจำแนกเสียงและทำการเปรียบเทียบตัวแบบจากวิธีซัพพอร์ตเวกเตอร์แมทชีน และวิธีหน่วยความจำระยะสั้นแบบยาว โดยมีรายละเอียดแต่ละวิธีดังนี้

5.1 วิธีซัพพอร์ตเวกเตอร์แมทชีน (support vector machine: SVM)

ซัพพอร์ตเวกเตอร์แมทชีน เป็นการเรียนรู้แบบมีผู้สอน ซึ่งนำมาช่วยแก้ปัญหาการจำแนกกลุ่มของข้อมูล โดยอาศัยหลักการของการหาค่าเหมาะที่สุด (optimization) เพื่อประมาณค่าน้ำหนักของตัวแบบเชิงเส้น เพื่อให้ได้เส้นตรงที่ใช้แบ่งกลุ่มข้อมูล (hyperplane) ที่เหมาะสมที่สุด แนวคิดของวิธี SVM คือการแปลงค่าของข้อมูลใน input space ไปยังพื้นที่ที่คุ้นลักษณะ (feature space) โดยอาศัย kernel function แล้วหาเส้นตรงที่ใช้แบ่งข้อมูลทั้งสองออกจากกัน โดยเส้นตรงที่แบ่งข้อมูลสองกลุ่มที่ดีที่สุดเป็นเส้นตรงที่มีระยะขอบมากที่สุด (Pinmuang & Thonhkam, 2017) แสดงดังรูปที่ 1 และสมการสำหรับการตัดสินใจ (Thamsiri & Meesad, 2011) แสดงดังสมการที่ (9)

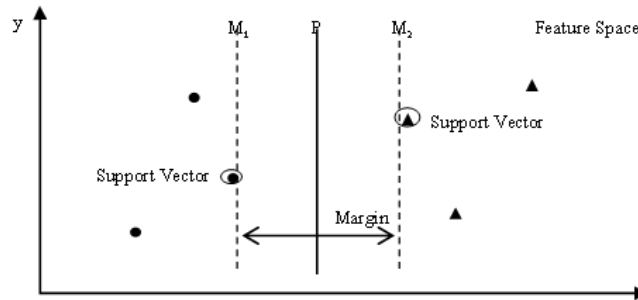
$$y = \text{sign}(\sum_{i=1}^{n_s} w_i \phi_i(x) \phi_i(x_i) + b) \quad (9)$$

$$\phi(x) = [\phi_1(x_1), \phi_2(x_2), \dots, \phi_n(x_n)]^T \quad (10)$$

โดย y คือ ค่ากลุ่มของข้อมูลมีค่าเป็น 1 เมื่อข้อมูลอยู่ในกลุ่มที่สนใจ (เป็นเสียงเปิดวาล์ว) และมีค่าเป็น -1 เมื่อข้อมูลอยู่ในกลุ่มที่ไม่สนใจ (ไม่ใช่เสียงเปิดวาล์ว)

w_i คือ ค่าน้ำหนัก (weight)

x คือ ค่าคุณลักษณะหรือข้อมูลนำเข้า
 x_i คือ ซัพพอร์ตเวกเตอร์
 n_s คือ จำนวนซัพพอร์ตเวกเตอร์
 และ b คือ ค่าเอนเอียง (bias)



รูปที่ 1 เส้นตรงในการจำแนกข้อมูลและระยะขอบ (Margin)

แต่ข้อมูลจริงที่พบส่วนใหญ่ นั้น ข้อมูลทั้ง 2 กลุ่มไม่สามารถแบ่งได้ด้วยสมการเส้นตรง ดังนั้นจึงต้องมีเครื่องมือมาช่วยให้ข้อมูลเหล่านั้นเรียงตัวใหม่ในพื้นที่หลักมิติ ซึ่งเครื่องมือเหล่านั้นคือ kernel function ดังสมการ (11)

$$K(x_i, x_j) = \phi(x_i)\phi(x_j) \quad (11)$$

ทั้งนี้ kernel function ที่นิยมใช้มีอยู่ 3 ชนิดด้วยกัน (Kaewtai, 2009) ได้แก่

- โพลีโนเมียล (polynomial)

$$K(x_i, x_j) = (\langle x_i^T x_j \rangle + 1)^d \quad (12)$$

เมื่อ d คือ ค่าเลขยกกำลัง

- เรเดียลเบสฟังก์ชัน (radial basis function: RBF)

$$K(x_i, x_j) = \exp\left(-\frac{\|x_i - x_j\|^2}{2\sigma^2}\right) \quad (13)$$

เมื่อ σ คือ ค่าพารามิเตอร์

- ซิกมอยด์ (sigmoid)

$$K(x_i, x_j) = \tanh(\alpha x_i^T x_j + c) \quad (14)$$

เมื่อ α และ c คือ ค่าพารามิเตอร์

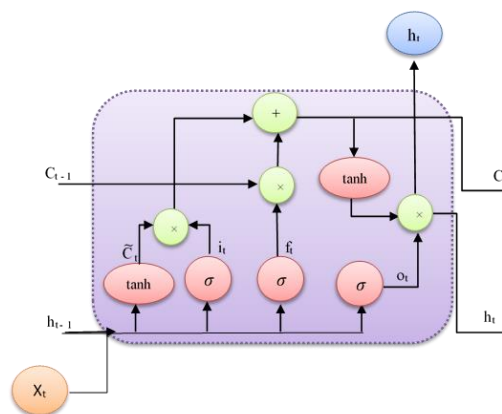
ดังนั้นสมการสำหรับการจำแนกข้อมูลสามารถแสดงดังสมการที่ (15)

$$f(x) = \text{sign}\left(\sum_{i=1}^{n_s} w_i K(x_i, x_j) + b\right) \quad (15)$$

เมื่อ $f(x)$ คือ ฟังก์ชันที่ใช้ในการจำแนกกลุ่มของ SVM โดยหาก $f(x)$ มีค่าเป็นบวก จะจำแนกเสียงเป็นเสียงเปิดวาล์ว และหาก มีค่าเป็นลบ จะจำแนกเสียงเป็นไม่ใช่เสียงเปิดวาล์ว

5.2 วิธีหน่วยความจำระยะสั้นแบบยาว (long short-term memory: LSTM)

หน่วยความจำระยะสั้นแบบยาว (LSTM) เป็นโครงข่ายประสาทเทียมแบบวนกลับชนิดพิเศษที่สามารถเรียนรู้อนุกรมเวลาระยะยาว หน่วยความจำระยะสั้นแบบยาวได้รับการออกแบบมาเพื่อหลีกเลี่ยงปัญหาการจดจำในระยะยาว โดยออกแบบให้มีเกต (gate) ทั้งหมด 3 เกต ได้แก่ อินพุตเกต (input gate) ฟอว์เกตเกต (forget gate) และ เอาต์พุตเกต (output gate) (Apaydin et al., 2020) นอกจากนี้ยังมีสถานะอินพุตของเซลล์ (input cell state: $\tilde{C}_t$) หน่วยความจำระยะสั้นแบบยาวมีโครงสร้างเป็นลูกโซ่ที่คล้ายกับโครงข่ายประสาทเทียมแบบวนกลับอย่างง่าย แต่มีความแตกต่างกันในโครงสร้างของตัวแบบที่ซับซ้อนกว่า เช่น เกต 3 เกต สำหรับทำหน้าที่ควบคุมการไหลของข้อมูลเข้าหรือข้อมูลออก และมีเซลล์สำหรับจดจำช่วงเวลาต่าง ๆ แสดงดังรูปที่ 2 (Cui et al., 2018)



รูปที่ 2 โครงสร้างของหน่วยความจำระยะสั้นแบบยาว (LSTM)

สำหรับการคำนวณค่าอินพุตเกต (i_t) ฟอว์เกตเกต (f_t) เอาต์พุตเกต (o_t) และ สถานะอินพุตของเซลล์ (input cell state: $\tilde{C}_t$) ของเวลา t จะคำนวณจากสมการ (16) ถึง (19)

$$f_t = \sigma_g(W_f x_t + U_f h_{t-1} + b_f) \quad (16)$$

$$i_t = \sigma_g(W_i x_t + U_i h_{t-1} + b_i) \quad (17)$$

$$o_t = \sigma_g(W_o x_t + U_o h_{t-1} + b_o) \quad (18)$$

$$\tilde{C}_t = \tanh(W_C x_t + U_C h_{t-1} + b_C) \quad (19)$$

โดยที่ W_p , W_h , W_o และ W_C เป็นเมทริกซ์น้ำหนักเลเยอร์ซ่อน (hidden layer) ข้อมูลเข้าของแต่ละเกทและสเตท
 U_p , U_h , U_o และ U_C เป็นเมทริกซ์น้ำหนักที่เชื่อมต่อสถานะเอาต์พุตของเซลล์ในอดีตกับแต่ละเกทและสเตท
 b_p , b_h , b_o และ b_C เป็นเวกเตอร์ค่าความเอนเอียงของแต่ละเกทและสเตท

และ σ_g คือ Activation Function ของเกท ซึ่งจะป็นฟังก์ชัน Sigmoid
 จะใช้สมการ (20) ถึง (21) เพื่อคำนวณสถานะเอาต์พุตของเซลล์ (C_t) และเอาต์พุตเลเยอร์ (h_t)

$$C_t = f_t \times C_{t-1} + i_t \times \tilde{C}_t \quad (20)$$

$$h_t = o_t \times \tanh(C_t) \quad (21)$$

โดยที่ $\tanh$ เป็น Activation Function ของเอาต์พุตเลเยอร์

สำหรับ Activation Function ของเอาต์พุตเลเยอร์ที่ใช้ในงานวิจัยนี้มี 3 ชนิด ได้แก่ rectified linear unit (relu), hyperbolic tangent (tanh) และ sigmoid ซึ่งเป็นฟังก์ชันทั้งปวงที่ใช้งานสำหรับ LSTM สามารถคำนวณได้ดังสมการที่ (22) ถึง สมการที่ (24)

- Sigmoid

$$g(z) = \frac{1}{1 + e^{-z}} \quad (22)$$

- Tanh

$$g(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}} \quad (23)$$

- ReLU

$$g(z) = \max(0, z) \quad (24)$$

6. การกำหนดพารามิเตอร์ของตัวแบบ

ในการสร้างตัวแบบการจำแนกเสียงเปิดวาล์วของอุปกรณ์คักไอน้ำ โดยใช้วิธีการทั้งสองวิธีจำเป็นต้องกำหนดและปรับค่าพารามิเตอร์เพื่อให้ได้ตัวแบบที่ให้ประสิทธิภาพในการจำแนกเสียงดีที่สุด การปรับค่าพารามิเตอร์ของตัวแบบพิจารณาจากการศึกษางานวิจัยที่เกี่ยวข้องกับการศึกษาในครั้งนี้ (Albon, 2017; Navlani, 2019)

โดยวิธี SVM จะทำการปรับค่าพารามิเตอร์ของตัวแบบทั้งหมด 3 ค่า ได้แก่ kernel function, C และ gamma แสดงค่าพารามิเตอร์ที่ปรับดังต่อไปนี้

-

- kernel function คือ การแปลงข้อมูลอินพุตของชุดข้อมูลที่กำหนดให้อยู่ในรูปแบบที่ต้องการ ในงานวิจัยนี้ใช้ kernel function 4 ฟังก์ชันได้แก่ linear, polynomial of 3rd degree, sigmoid และ radial basis function
- C คือ ค่าที่กำหนดให้มีจุดข้อมูลละเมิดเส้นขอบ โดยส่งผลต่อความกว้างของเส้นขอบ ในงานวิจัยนี้กำหนด 4 ค่า ได้แก่ 10, 100, 500 และ 1,000
- gamma คือ การกำหนดการแพร่กระจายของ kernel จึงเป็นขอบเขตการตัดสินใจ หากมีค่าต่ำ เส้นโค้งของขอบเขตการตัดสินใจจะต่ำ ในทางตรงข้ามหากมีค่าสูงขอบเขตการตัดสินใจจะกว้าง ดังนั้นเส้นโค้งของขอบเขตการตัดสินใจจะสูง ซึ่งจะสร้างเป็นขอบเขตการตัดสินใจรอบ ๆ จุดข้อมูล ในงานวิจัยนี้กำหนด 7 ค่า ได้แก่ 0.000001 0.00001 0.0001 0.001 0.01 0.1 และ 1.0 ยกเว้น gamma ของ linear function จะมีค่าเท่ากับ 1

วิธี LSTM จะทำการสร้างตัวแบบจำลองตามลำดับ (sequential model) เป็นการเรียงซ้อนของชั้นเชิงเส้น โดยมีชั้นของข้อมูลเข้า (input layer) และชั้นของข้อมูลผลการวิเคราะห์ (output layer) นอกจากนี้ยังมีฟังก์ชันสำหรับการเรียนรู้ และการทำนาย ซึ่งจะมีการปรับค่าพารามิเตอร์ (Zhu & Chollet, 2020; Brownlee, 2017; Lim et al., 2019; TensorFlow, 2021) ดังนี้

ชั้นของข้อมูลเข้า (input layer) ซึ่งเป็นการกำหนดลักษณะต่าง ๆ ของข้อมูลเข้า จะกำหนดคุณสมบัติของชั้นข้อมูลเข้า 4 ค่า ได้แก่ units, input shape, dropout และ return sequences มีการกำหนดค่าดังต่อไปนี้

- units กำหนดค่าเท่ากับ 64
- input shape กำหนดค่าเป็น (1, 536) โดยค่าแรกคือ จำนวนแถวของข้อมูลเข้าในแต่ละรอบ ซึ่งในแต่ละรอบจะนำเข้า 1 แถว และค่าที่สอง คือจำนวนของตัวแปรหรือจำนวนคุณลักษณะ (feature) ของข้อมูล ซึ่งมี 536 ค่า
- dropout กำหนดให้มีค่าตั้งแต่ 0.1 จนถึง 0.5 เพิ่มทีละ 0.1
- return sequences กำหนดให้มีค่าเป็นจริง (true)

ชั้นของข้อมูลผลการวิเคราะห์ (output layer) จะต้องกำหนดคุณสมบัติของชั้นผลการวิเคราะห์ 2 ค่า ได้แก่ units และ activation function โดย units กำหนดค่าเท่ากับ 1 เพราะเป็นการสร้างตัวแบบเพื่อจำแนกกลุ่ม 2 กลุ่มจะแสดงค่าความน่าจะเป็นที่จะเป็นกลุ่ม 1 หรือกลุ่มที่สนใจ ส่วน activation function ของ output layer ใช้ 3 ฟังก์ชัน ได้แก่ tanh, relu และ sigmoid ดังแสดงในสมการที่ (22) ถึง สมการที่ (24)

ฟังก์ชันสำหรับการเรียนรู้จะใช้ฟังก์ชัน compile เป็นการกำหนดค่ากระบวนการเรียนรู้ของตัวแบบ มีการกำหนดค่า 3 ค่า ได้แก่ loss, optimizer และ metrics โดยกำหนดค่าดังต่อไปนี้

- loss เป็นการกำหนดฟังก์ชันการสูญเสีย (loss function) ให้เป็น binary_crossentropy เนื่องจากเป็นปัญหาการจำแนกกลุ่ม 2 กลุ่ม ผลลัพธ์การทำนายของตัวแบบจะบอกค่าความน่าจะเป็นว่ามีโอกาสที่จะเป็นกลุ่ม 1 โดยปกติจะกำหนดให้ 1 แทนกลุ่มที่สนใจและ 0 แทนกลุ่มที่ไม่สนใจ
- optimizer ใช้อัลกอริทึม ADAM ซึ่งอัลกอริทึมนี้สามารถปรับค่า learning rate สำหรับพารามิเตอร์ในแต่ละครั้งได้ ดังนั้นจึงปรับค่า learning rate ทั้งหมด 6 ค่า ได้แก่ 0.1 0.01 0.001 0.0001 และ 0.00001

- metrics เป็นการกำหนดเกณฑ์การประเมินผลตัวแบบ งานวิจัยนี้ใช้เกณฑ์ accuracy คือให้แสดงในรูปแบบของค่าความถูกต้อง

ฟังก์ชันสำหรับการทำนายจะใช้ฟังก์ชัน fit เป็นฟังก์ชันในการเรียนรู้โดยใช้ข้อมูลเรียนรู้ (training set) ซึ่งทำการกำหนดค่าอาร์กิวเมนต์ 2 ค่า ได้แก่ epochs และ batch size โดยกำหนดค่าดังต่อไปนี้

- epochs เท่ากับ 100 เป็นการวนซ้ำข้อมูล x และ y ในการเรียนรู้ตัวแบบโดยแต่ละ Epoch จะทำให้ Loss ลดลง ในขณะที่ Accuracy เพิ่มขึ้น
- batch size เท่ากับ 64 เป็นการกำหนดจำนวนตัวอย่างต่อการอัปเดตค่า gradients

ในการแยกคุณลักษณะของเสียงจากโดยใช้วิธีการ 5 วิธีและการสร้างตัวแบบทั้งสองวิธี เขียนโปรแกรมโดยใช้ภาษา Python บน Google Colab (<https://colab.research.google.com/>)

7. การวัดประสิทธิภาพของตัวแบบ

งานวิจัยนี้ใช้เกณฑ์การวัดประสิทธิภาพของตัวแบบ ได้แก่ ค่าความเที่ยง (precision) ค่าการเรียกคืน (recall) และค่า F1 (F1-score) โดยมีทางเลือกในการสรุปผลการทำนายดังตารางที่ 2

ตารางที่ 2 Confusion Matrix

ค่าทำนาย	ค่าจริง	
	เสียงเปิดวาล์ว	ไม่ใช่เสียงเปิดวาล์ว
เสียงเปิดวาล์ว	TP	FP
ไม่ใช่เสียงเปิดวาล์ว	FN	TN

โดยที่ true positive (TP) คือ ข้อมูลที่ค่าจริงเป็นเสียงเปิดวาล์ว ทำนายถูกว่าเป็นเสียงเปิดวาล์ว
 true negative (TN) คือ ข้อมูลที่ค่าจริงไม่ใช่เสียงเปิดวาล์ว ทำนายถูกว่าไม่ใช่เสียงเปิดวาล์ว
 false positive (FP) คือ ข้อมูลที่ค่าจริงไม่ใช่เสียงเปิดวาล์ว ทำนายผิดว่าเป็นเสียงเปิดวาล์ว
 false negative (FN) คือ ข้อมูลที่ค่าจริงเป็นเสียงเปิดวาล์ว ทำนายผิดว่าไม่ใช่เสียงเปิดวาล์ว

7.1 ค่าความเที่ยง (precision)

คำนวณหาการทำนายที่ถูกต้องเพื่อวัดความแม่นยำของตัวแบบ โดยคิดจากสัดส่วนของจำนวนการทำนายเสียงเปิดวาล์วถูกต้อง (true positive) เทียบกับผลรวมของจำนวนการทำนายว่าเป็นเสียงเปิดวาล์วทั้งหมด

$$\text{Precision} = \frac{TP}{TP + FP} \quad (25)$$

7.2 ค่าการเรียกคืน (recall)

คำนวณหาความถูกต้องในการทำนายเชิงเปิดวาล์วของตัวแบบโดยคิดจากสัดส่วนของจำนวนการทำนายว่าเป็นเชิงเปิดวาล์วถูกต้อง (true positive) เทียบกับจำนวนเชิงเปิดวาล์วในข้อมูลชุดนั้น

$$\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}} \quad (26)$$

7.3 ค่า F1 (F1-score)

ค่าเฉลี่ยฮาร์โมนิก (harmonic mean) ระหว่างค่าความเที่ยง (precision) และค่าการเรียกคืน (recall) และใช้เพื่อหาตัวแบบที่มีความสมดุลที่ดีระหว่างค่าความเที่ยง และค่าการเรียกคืน คำนวณได้จากสมการที่ (24)

$$\text{F1-Score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (27)$$

8. ผลและวิจารณ์ผลการทดลอง

จากการทดลองสร้างตัวแบบด้วยวิธี SVM และวิธี LSTM โดยทำการปรับค่าพารามิเตอร์ตามที่กล่าวไว้ในหัวข้อ 6 ดำเนินการสร้างตัวแบบโดยใช้ชุดข้อมูลการเรียนรู้ที่ข้อมูลไม่สมดุลกับชุดข้อมูลการเรียนรู้ที่ข้อมูลสมดุลที่ปรับแก้ด้วยการใช้วิธีผสมผสาน ผลการศึกษาเมื่อทดสอบด้วยข้อมูลชุดทดสอบแสดงดังตารางที่ 3 และตารางที่ 4

ตารางที่ 3 ตารางผลการวิเคราะห์ของข้อมูลชุดทดสอบที่สร้างตัวแบบด้วยวิธี SVM

พารามิเตอร์			ชุดการเรียนรู้ข้อมูลที่ไม่สมดุล			ชุดการเรียนรู้ข้อมูลสมดุล		
kernel	gamma	C	Precision (%)	Recall (%)	F1-score (%)	Precision (%)	Recall (%)	F1-score (%)
linear	1	100	62.50	50.00	55.56	50.00	60.00	54.55
linear	1	500	62.50	50.00	55.56	50.00	60.00	54.55
linear	1	1000	62.50	50.00	55.56	50.00	60.00	54.55
poly	0.000001	10	0.00	0.00	0.00	0.00	0.00	0.00
poly	0.000001	100	0.00	0.00	0.00	0.00	0.00	0.00
poly	0.000001	500	0.00	0.00	0.00	0.00	0.00	0.00
poly	0.000001	1000	0.00	0.00	0.00	0.00	0.00	0.00
poly	0.00001	10	0.00	0.00	0.00	0.00	0.00	0.00
poly	0.00001	100	0.00	0.00	0.00	0.00	0.00	0.00
poly	0.00001	500	0.00	0.00	0.00	50.00	10.00	16.67
poly	0.00001	1000	0.00	0.00	0.00	28.57	20.00	23.53

poly	0.0001	10	0.00	0.00	0.00	27.78	50.00	35.71
poly	0.0001	100	0.00	0.00	0.00	50.00	70.00	58.33
poly	0.0001	500	100.00	10.00	18.18	53.85	70.00	60.87
poly	0.0001	1000	50.00	10.00	16.67	58.33	70.00	63.64
poly	0.001	10	100.00	30.00	46.15	58.33	70.00	63.64
poly	0.001	100	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.001	500	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.001	1000	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.01	10	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.01	100	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.01	500	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.01	1000	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.1	10	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.1	100	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.1	500	80.00	40.00	53.33	63.64	70.00	66.67
poly	0.1	1000	80.00	40.00	53.33	63.64	70.00	66.67
poly	1	10	80.00	40.00	53.33	63.64	70.00	66.67
poly	1	100	80.00	40.00	53.33	63.64	70.00	66.67
poly	1	500	80.00	40.00	53.33	63.64	70.00	66.67
poly	1	1000	80.00	40.00	53.33	63.64	70.00	66.67
sigmoid	0.000001	10	0.00	0.00	0.00	13.16	100.00	23.26
sigmoid	0.000001	100	0.00	0.00	0.00	40.00	80.00	53.33
sigmoid	0.000001	500	0.00	0.00	0.00	44.44	80.00	57.14
sigmoid	0.000001	1000	33.33	10.00	15.38	44.44	80.00	57.14
sigmoid	0.00001	10	0.00	0.00	0.00	40.00	80.00	53.33
sigmoid	0.00001	100	33.33	10.00	15.38	44.44	80.00	57.14
sigmoid	0.00001	500	50.00	30.00	37.50	50.00	70.00	58.33
sigmoid	0.00001	1000	57.14	40.00	47.06	54.55	60.00	57.14
sigmoid	0.0001	10	33.33	10.00	15.38	47.06	80.00	59.26

sigmoid	0.0001	100	55.56	50.00	52.63	53.85	70.00	60.87
sigmoid	0.0001	500	41.18	70.00	51.85	42.11	80.00	55.17
sigmoid	0.0001	1000	36.84	70.00	48.28	33.33	80.00	47.06
sigmoid	0.001	10	27.27	30.00	28.57	9.28	90.00	16.82
sigmoid	0.001	100	25.00	30.00	27.27	9.18	90.00	16.67
sigmoid	0.001	500	23.08	30.00	26.09	9.18	90.00	16.67
sigmoid	0.001	1000	23.08	30.00	26.09	9.18	90.00	16.67
sigmoid	0.01	10	0.00	0.00	0.00	0.47	70.00	0.94
sigmoid	0.01	100	0.00	0.00	0.00	0.47	70.00	0.94
sigmoid	0.01	500	0.00	0.00	0.00	0.47	70.00	0.94
sigmoid	0.01	1000	0.00	0.00	0.00	0.47	70.00	0.94
sigmoid	0.1	10	50.00	10.00	16.67	0.27	60.00	0.55
sigmoid	0.1	100	50.00	10.00	16.67	0.27	60.00	0.55
sigmoid	0.1	500	50.00	10.00	16.67	0.27	60.00	0.55
sigmoid	0.1	1000	50.00	10.00	16.67	0.27	60.00	0.55
sigmoid	1	10	33.33	10.00	15.38	0.40	80.00	0.79
sigmoid	1	100	25.00	10.00	14.29	0.40	80.00	0.79
sigmoid	1	500	25.00	10.00	14.29	0.40	80.00	0.79
sigmoid	1	1000	25.00	10.00	14.29	0.40	80.00	0.79
rbf	0.000001	10	0.00	0.00	0.00	22.50	90.00	36.00
rbf	0.000001	100	0.00	0.00	0.00	42.11	80.00	55.17
rbf	0.000001	500	33.33	10.00	15.38	44.44	80.00	57.14
rbf	0.000001	1000	33.33	10.00	15.38	50.00	80.00	61.54
rbf	0.00001	10	0.00	0.00	0.00	42.11	80.00	55.17
rbf	0.00001	100	33.33	10.00	15.38	47.06	80.00	59.26
rbf	0.00001	500	57.14	40.00	47.06	54.55	60.00	57.14
rbf	0.00001	1000	57.14	40.00	47.06	50.00	60.00	54.55
rbf	0.0001	10	33.33	10.00	15.38	43.75	70.00	53.85
rbf	0.0001	100	57.14	40.00	47.06	53.85	70.00	60.87

rbf	0.0001	500	62.50	50.00	55.56	53.85	70.00	60.87
rbf	0.0001	1000	62.50	50.00	55.56	53.85	70.00	60.87
rbf	0.001	10	50.00	40.00	44.44	50.00	70.00	58.33
rbf	0.001	100	50.00	40.00	44.44	50.00	70.00	58.33
rbf	0.001	500	50.00	40.00	44.44	50.00	70.00	58.33
rbf	0.001	1000	50.00	40.00	44.44	50.00	70.00	58.33
rbf	0.01	10	0.00	0.00	0.00	0.00	0.00	0.00
rbf	0.01	100	0.00	0.00	0.00	0.00	0.00	0.00
rbf	0.01	500	0.00	0.00	0.00	0.00	0.00	0.00
rbf	0.01	1000	0.00	0.00	0.00	0.00	0.00	0.00
rbf	0.1	10	0.00	0.00	0.00	0.00	0.00	0.00
rbf	0.1	100	0.00	0.00	0.00	0.00	0.00	0.00
rbf	0.1	500	0.00	0.00	0.00	0.00	0.00	0.00
rbf	0.1	1000	0.00	0.00	0.00	0.00	0.00	0.00
rbf	1	10	0.00	0.00	0.00	0.00	0.00	0.00
rbf	1	100	0.00	0.00	0.00	0.00	0.00	0.00
rbf	1	500	0.00	0.00	0.00	0.00	0.00	0.00
rbf	1	1000	0.00	0.00	0.00	0.00	0.00	0.00

ตารางที่ 4 ตารางผลการวิเคราะห์ของข้อมูลชุดทดสอบที่สร้างตัวแบบโดยใช้วิธี LSTM

พารามิเตอร์			ชุดการเรียนรู้ข้อมูลที่ไม่สมดุล			ชุดการเรียนรู้ข้อมูลสมดุล		
Learning rate	Activation func.	Dropout	Precision (%)	Recall (%)	F1-score (%)	Precision (%)	Recall (%)	F1-score (%)
0.1	tanh	0.1	0.00	0.00	0.00	20.59	70.00	31.82
0.1	tanh	0.2	0.00	0.00	0.00	2.01	80	3.91
0.1	tanh	0.3	0.00	0.00	0.00	23.33	70.00	35.00
0.1	tanh	0.4	0.00	0.00	0.00	21.05	80.00	33.33
0.1	tanh	0.5	0.00	0.00	0.00	17.50	70.00	28.00
0.1	relu	0.1	0.00	0.00	0.00	0.39	80	0.78

0.1	relu	0.2	0.00	0.00	0.00	0.29	60	0.59
0.1	relu	0.3	0.00	0.00	0.00	0.47	100	0.93
0.1	relu	0.4	0.00	0.00	0.00	2.12	50	4.07
0.1	relu	0.5	0.00	0.00	0.00	2.34	90	4.57
0.1	sigmoid	0.1	33.33	20.00	25.00	25.93	70	37.84
0.1	sigmoid	0.2	50.00	30.00	37.50	31.58	60	41.38
0.1	sigmoid	0.3	23.08	30.00	26.09	29.17	70	41.18
0.1	sigmoid	0.4	33.33	20.00	25.00	29.17	70	41.18
0.1	sigmoid	0.5	33.33	10.00	15.38	25	60	35.29
0.01	tanh	0.1	0.00	0.00	0.00	33.33	60.00	42.86
0.01	tanh	0.2	0.00	0.00	0.00	33.33	60.00	42.86
0.01	tanh	0.3	10.00	10.00	10.00	25.64	100.00	40.82
0.01	tanh	0.4	50.00	20.00	28.57	38.89	70.00	50.00
0.01	tanh	0.5	0.00	0.00	0.00	42.11	80.00	55.17
0.01	relu	0.1	0.00	0.00	0.00	33.33	70.00	45.16
0.01	relu	0.2	0.00	0.00	0.00	23.08	90.00	36.73
0.01	relu	0.3	0.00	0.00	0.00	40.91	90.00	56.25
0.01	relu	0.4	0.00	0.00	0.00	38.89	70.00	50.00
0.01	relu	0.5	33.33	20.00	25.00	40.00	80.00	53.33
0.01	sigmoid	0.1	57.14	40.00	47.06	33.33	60.00	42.86
0.01	sigmoid	0.2	50.00	40.00	44.44	44.44	80.00	57.14
0.01	sigmoid	0.3	50.00	40.00	44.44	36.84	70.00	48.28
0.01	sigmoid	0.4	50.00	50.00	50.00	42.11	80.00	55.17
0.01	sigmoid	0.5	50.00	50.00	50.00	38.89	70.00	50.00
0.001	tanh	0.1	40.00	20.00	26.67	35.29	60.00	44.44
0.001	tanh	0.2	40.00	20.00	26.67	38.89	70.00	50.00
0.001	tanh	0.3	50.00	40.00	44.44	52.94	90.00	66.67
0.001	tanh	0.4	50.00	50.00	50.00	41.18	70.00	51.85
0.001	tanh	0.5	50.00	50.00	50.00	42.11	80.00	55.17

0.001	relu	0.1	33.33	10.00	15.38	38.89	70.00	50.00
0.001	relu	0.2	0.00	0.00	0.00	38.89	70.00	50.00
0.001	relu	0.3	60.00	60.00	60.00	37.50	60.00	46.15
0.001	relu	0.4	60.00	30.00	40.00	42.11	80.00	55.17
0.001	relu	0.5	62.50	50.00	55.56	38.89	70.00	50.00
0.001	sigmoid	0.1	50.00	60.00	54.55	35.29	60.00	44.44
0.001	sigmoid	0.2	50.00	60.00	54.55	38.89	70.00	50.00
0.001	sigmoid	0.3	60.00	60.00	60.00	40.00	60.00	48.00
0.001	sigmoid	0.4	50.00	60.00	54.55	41.18	70.00	51.85
0.001	sigmoid	0.5	50.00	60.00	54.55	41.18	70.00	51.85
0.0001	tanh	0.1	50.00	40.00	44.44	37.50	60.00	46.15
0.0001	tanh	0.2	45.45	50.00	47.62	37.50	60.00	46.15
0.0001	tanh	0.3	50.00	60.00	54.55	38.89	70.00	50.00
0.0001	tanh	0.4	45.45	50.00	47.62	35.29	60.00	44.44
0.0001	tanh	0.5	38.46	50.00	43.48	38.89	70.00	50.00
0.0001	relu	0.1	55.56	50.00	52.63	37.50	60.00	46.15
0.0001	relu	0.2	46.15	60.00	52.17	41.18	70.00	51.85
0.0001	relu	0.3	54.55	60.00	57.14	41.18	70.00	51.85
0.0001	relu	0.4	33.33	10.00	15.38	38.89	70.00	50.00
0.0001	relu	0.5	54.55	60.00	57.14	38.89	70.00	50.00
0.0001	sigmoid	0.1	54.55	60.00	57.14	37.50	60.00	46.15
0.0001	sigmoid	0.2	60.00	60.00	60.00	35.29	60.00	44.44
0.0001	sigmoid	0.3	54.55	60.00	57.14	38.89	70.00	50.00
0.0001	sigmoid	0.4	60.00	60.00	60.00	41.18	70.00	51.85
0.0001	sigmoid	0.5	55.56	50.00	52.63	42.11	80.00	55.17
0.00001	tanh	0.1	50.00	50.00	50.00	38.89	70.00	50.00
0.00001	tanh	0.2	44.44	40.00	42.11	40.91	90.00	56.25
0.00001	tanh	0.3	40.00	20.00	26.67	40.00	80.00	53.33
0.00001	tanh	0.4	0.00	0.00	0.00	33.33	90.00	48.65

0.00001	tanh	0.5	0.00	0.00	0.00	11.11	100.00	20.00
0.00001	relu	0.1	100.00	10.00	18.18	8.79	80.00	15.84
0.00001	relu	0.2	57.14	40.00	47.06	7.30	100.00	13.61
0.00001	relu	0.3	45.45	50.00	47.62	1.44	100	2.84
0.00001	relu	0.4	25.00	10.00	14.29	1.22	100	2.41
0.00001	relu	0.5	0.00	0.00	0.00	1.15	100	2.27
0.00001	sigmoid	0.1	40.00	20.00	26.67	42.11	80.00	55.17
0.00001	sigmoid	0.2	60.00	30.00	40.00	42.11	80.00	55.17
0.00001	sigmoid	0.3	50.00	10.00	16.67	40.00	80.00	53.33
0.00001	sigmoid	0.4	50.00	10.00	16.67	36.00	90.00	51.43
0.00001	sigmoid	0.5	0.00	0.00	0.00	24.32	90.00	38.30

จากตารางที่ 3 และตารางที่ 4 สามารถสรุปผลที่ได้ค่า F1 สูงที่สุดของแต่ละตัวแบบที่สร้างจากข้อมูลการเรียนรู้ทั้ง 2 ชุดข้อมูลได้ดังตารางที่ 5

ตารางที่ 5 เปรียบเทียบประสิทธิภาพของตัวแบบที่สร้างจากชุดข้อมูลการเรียนรู้ทั้ง 2 ชุด โดยแสดงค่า F1 สูงที่สุด

ตัวแบบ	ข้อมูลที่ไม่สมดุล			ข้อมูลที่สมดุล		
	Precision (%)	Recall (%)	F1-score (%)	Precision (%)	Recall (%)	F1-score (%)
SVM	62.50	50.00	55.56	63.64	70.00	66.67
LSTM	60.00	60.00	60.00	52.94	90.00	66.67

จากตารางที่ 5 จะพบว่าตัวแบบที่เรียนรู้จากชุดข้อมูลที่สมดุลให้ค่า F1 สูงกว่าตัวแบบที่เรียนรู้จากชุดข้อมูลที่ไม่สมดุล สำหรับการสร้างตัวแบบจากชุดข้อมูลการเรียนรู้ที่ไม่สมดุล ตัวแบบ LSTM จะมีค่า F1 และการเรียกคืนมากกว่าตัวแบบ SVM โดยมีค่าเท่ากับ 60.00% แต่ค่าความเที่ยงต่ำกว่าตัวแบบ SVM เล็กน้อยโดยมีค่าเท่ากับ 60.00% สำหรับชุดข้อมูลเรียนรู้ที่สมดุล ผลการศึกษาพบว่า ตัวแบบ SVM ให้ค่า F1 เท่ากับตัวแบบ LSTM โดยมีค่าเท่ากับ 66.67% และมีความเที่ยงสูงกว่าตัวแบบ LSTM ซึ่งมีค่าเท่ากับ 63.64% แต่มีการเรียกคืนต่ำกว่าตัวแบบ LSTM มีค่าเท่ากับ 70.00% โดยฟังก์ชันเคอร์เนลที่ให้ผลลัพธ์ดีที่สุดคือ polynomial 3rd degree เมื่อกำหนดค่า gamma เป็น 0.001, 0.01, 0.1 และ 1 กำหนดค่า C เป็น 10, 100, 500 และ 1,000 ส่วนตัวแบบ LSTM ที่สร้างจากชุดข้อมูลเรียนรู้ที่สมดุลมีประสิทธิภาพดีที่สุดเมื่อกำหนดค่า learning rate ของ optimizer ADAM เท่ากับ 0.001 ใช้ activation function ของ output layer เป็นฟังก์ชัน tanh และกำหนดค่า dropout เท่ากับ 0.3 โดยมีค่าการเรียกคืนสูงที่สุดเท่ากับ 90.00%

ค่าความเที่ยงต่ำที่สุดเท่ากับ 52.94% และค่า F1 เท่ากับ 66.67% แม้ว่าค่า F1 ของทั้งสองตัวแบบจะมีค่าเท่ากันแต่พบว่าตัวแบบ LSTM สามารถทำนายเสียงเปิดวาล์วได้ถูกต้องมากกว่าตัวแบบ SVM แต่ให้ค่าความเที่ยงน้อยกว่าตัวแบบ SVM เมื่อพิจารณาค่าความเที่ยงและค่าการเรียกคืนพร้อมกัน จึงได้ค่า F1 เท่ากัน

อย่างไรก็ตามจากตารางที่ 3 และตารางที่ 4 จะพบว่าตัวแบบที่มีประสิทธิภาพดีที่สุด สามารถกำหนดค่าพารามิเตอร์ที่แตกต่างกันได้ โดยจะให้ค่า F1 ค่าความเที่ยง และค่าการเรียกคืนเท่ากัน

สรุปผลการทดลอง

การวิจัยนี้สร้างตัวแบบการจำแนกเสียงเปิดวาล์วของอุปกรณ์คักไอน้ำ โดยใช้วิธีการจำแนกกลุ่ม 2 วิธี ได้แก่ วิธี SVM และวิธี LSTM และกำหนด activation function ของ output layer สามฟังก์ชัน ได้แก่ tanh relu และ sigmoid และใช้การแยกคุณลักษณะของเสียง 5 วิธี ได้แก่ อัตราค่าตัดศูนย์ spectral centroid spectral rolloff สัมประสิทธิ์เซปสตรัมบนความถี่แบบเมล และการแปลงฟูเรียร์ระยะสั้น นอกจากนี้ได้ทำการทดลองสร้างตัวแบบจำแนกกลุ่มโดยใช้ชุดข้อมูลเรียนรู้ 2 ชุด ได้แก่ ชุดข้อมูลเรียนรู้ที่ไม่สมดุลกับชุดข้อมูลเรียนรู้ที่สมดุลที่ได้ปรับแก้ด้วยวิธีผสมผสาน ผลการวิเคราะห์จำแนกกลุ่มโดยใช้ข้อมูลชุดทดสอบพบว่าตัวแบบ SVM และ LSTM ที่สร้างจากชุดข้อมูลเรียนรู้ที่สมดุลมีค่า F1 สูงกว่าการฝึกสร้างจากชุดข้อมูลเรียนรู้ที่ไม่สมดุล โดยตัวแบบทั้งสองให้ค่า F1 เท่ากันเท่ากับ 66.67% อย่างไรก็ตามพบว่าตัวแบบ SVM มีค่าความเที่ยงมากกว่า ในขณะที่ตัวแบบ LSTM มีค่าการเรียกคืนมากกว่า ดังนั้นตัวแบบ LSTM ทำนายผลเสียงเปิดวาล์วถูกต้องมากกว่าตัวแบบ SVM แต่ทำนายเสียงที่ไม่ใช่เสียงเปิดวาล์วผิดเป็นเสียงเปิดวาล์วมากกว่าตัวแบบ SVM เนื่องจากมีค่าความเที่ยงต่ำกว่า เมื่อนำค่าความเที่ยงและค่าการเรียกคืนมาคำนวณค่า F1 แล้วจึงได้ค่าเท่ากัน ทั้งนี้ตัวแบบ SVM ที่เหมาะสมคือตัวแบบที่ใช้ kernel function คือ polynimail 3rd degree และสามารถกำหนดค่า C ไม่น้อยกว่า 100 และ gamma ไม่น้อยกว่า 0.001 สำหรับตัวแบบ LSTM ที่เหมาะสมมีจำนวนโหนดใน Input Layer เท่ากับ 64 unit กำหนดค่า dropout เท่ากับ 0.3 และ output layer มีโหนด 1 unit กำหนด activation function ของ output layer เป็น tanh และกำหนดค่า learning rate เท่ากับ 0.001 อย่างไรก็ตามพบปัญหาในการวิจัยคือขณะบันทึกเสียงอุปกรณ์คักไอน้ำจะมีเสียงรบกวนจากเครื่องจักรในโรงงาน ควรศึกษาวิธีการบันทึกเสียงที่ลดเสียงรบกวนได้ หรือสร้างพื้นที่ที่ลดเสียงรบกวนจากเครื่องจักรรอบ ๆ อุปกรณ์คักไอน้ำ สำหรับการวิจัยในอนาคตอาจพิจารณาวิธีการเรียนรู้แบบรวมกลุ่ม (Ensemble Learning) สำหรับการจำแนกข้อมูลอาจทำให้ความแม่นยำของตัวแบบจำแนกเสียงสูงขึ้นและอาจใช้วิธีการปรับค่าพารามิเตอร์โดยวิธีการ Grid Search

กิตติกรรมประกาศ

ขอขอบพระคุณบัณฑิตวิทยาลัย มหาวิทยาลัยเชียงใหม่สำหรับทุนผู้ช่วยสอน/ผู้ช่วยวิจัย และขอขอบพระคุณคณะวิทยาศาสตร์และคณะวิศวกรรมศาสตร์ มหาวิทยาลัยเชียงใหม่ที่ให้การสนับสนุนการทำวิจัย

เอกสารอ้างอิง

- Ahmadizadeh, M. (2014). *An Introduction to Short-Time Fourier Transform (STFT)*.
<https://sharif.edu/~ahmadizadeh/courses/advstdyn/Short-Time%20Fourier%20Transform.pdf>
- Albon, C. (2017). *SVC Parameters When Using RBF Kernel*. https://chrisalbon.com/code/machine_learning/support_vector_machines/svc_parameters_using_rbf_kernel/
- Apaydin, H., Feizi, H., Sattari, M. T., Colak, M. S., Shamshirband, S., & Chau, K.W. (2020). Comparative Analysis of Recurrent Neural Network Architectures for Reservoir Inflow Forecasting. *Water*, 12(5), 18.
<https://doi.org/10.3390/w12051500>
- Brownlee, J. (2016). *Sequence Classification with LSTM Recurrent Neural Networks in Python with Keras*.
<https://machinelearningmastery.com/sequence-classification-lstm-recurrent-neural-networks-python-keras/>
- Brownlee, J. (2017). *Dropout with LSTM Networks for Time Series Forecasting*.
<https://machinelearningmastery.com/use-dropout-lstm-networks-time-series-forecasting/>
- Brownlee, J. (2017). *How to Reshape Input Data for Long Short-Term Memory Networks in Keras*.
<https://machinelearningmastery.com/reshape-input-data-long-short-term-memory-networks-keras/>
- Burred, J. J., & Lerch, A. (2004). Hierarchical Automatic Audio Signal Classification. *Journal of The Audio Engineering Society*, 52(7/8), 724-739.
- Cohen, L. (1995). *Time-frequency analysis*. Englewood Cliffs, New Jersey: Prentice hall.
- Cui, Z., Ke, R., & Wang, Y. (2018). Deep Bidirectional and Unidirectional LSTM Recurrent Neural Network for Network-wide Traffic Speed Prediction. *arXiv preprint arXiv: 1801.02143*.
- Department of Industrial Works. (2010). *Safety Technology Bureau, Manual and maintenance of the boiler*.
<https://php.diw.go.th/safety/wp-content/uploads/2015/01/boiler2.pdf> (in Thai)
- Galván-Tejada, C. E., Galván-Tejada, J. I., Celaya-Padilla, J. M., Delgado-Contreras, J. R., Magallanes-Quintanar, R., Martinez-Fierro, M. L., Garza-Veloz, M. F., López-Hernández, Y., & Gamboa-Rosales, H. (2016). An analysis of audio features to develop a human activity recognition model using genetic algorithms, random forests, and neural networks. *Mobile Information Systems*, 2016.
- Giannakopoulos, T. (2015). pyAudioAnalysis: An Open-Source Python Library for Audio Signal Analysis. *PLoS One*, 10(12), e0144610. <https://doi.org/10.1371/journal.pone.0144610>

- Hassan, S., Rafi, M., & Shaikh, M. S. (2011). Comparing SVM and naïve Bayes classifiers for text categorization with Wikitology as knowledge enrichment. In *2011 IEEE 14th International Multitopic Conference* (pp. 31-34). Karachi, Pakistan. <https://doi.org/10.1109/INMIC.2011.6151495>
- Junto, A. (2015). *Sound-based road traffic detection using artificial intelligent approach* [Unpublished master's thesis]. Suranaree University of Technology. (in Thai)
- Kaewtai, O. (2009). *Thai Classical Music Classification* [Unpublished Master's Thesis]. Silpakorn University. (in Thai)
- Kesornsit, W., Lorchirachoonkul, V., & Jitthavech, J. (2018). Imbalanced Data Problem Solving in Classification of Diabetes Patients. *KKU Research journal*, 18(3), 11. (in Thai)
- Kittinaradorn, C. (2020). *Neural Network Programming*. <https://guopai.github.io/ml-blog15.html> (in Thai)
- Kopanyapipat, R. (2015). *Noise robust speech recognition for thai language using N-based limabeam algorithm*. [Unpublished Master's Thesis]. Thammasat University. (in Thai)
- Li, T., Chen, X. R., Tang, H., & Xu, X. K. (2018). Identification of the Normal/Abnormal Heart Sounds Based on Energy Features and Xgboost. In *Biometric Recognition: 13th Chinese Conference, CCBR 2018, Urumqi, China, 11-12 August, 2018, Proceeding 13* (pp. 536-544). https://doi.org/10.1007/978-3-319-97909-0_57
- Lim, S., Jang S., Lim, J., & Ko, J. (2019). Classification of Snoring Sound Based on a Recurrent Neural Network. *Expert Systems with Applications*, 123, 237-245. <https://doi.org/10.1016/j.eswa.2019.01.020>
- Ministry of Energy. (2010). *Industrial Steam System*. https://www2.dede.go.th/bhrd/old/Download/file_handbook/Pre_Build/Build_12.pdf (in Thai)
- Ministry of Energy. (2011). *Steam Trap Management*. https://www2.dede.go.th/bhrd/old/Download/file_handbook/Pre_T_H/Heat_4.pdf (in Thai)
- Navlani, A. (2019). *Support Vector Machines with Scikit-learn*. <https://www.datacamp.com/community/tutorials/svm-classification-scikit-learn-python#tuning>
- Pinmuang, N., & Thongkam, J. (2017). Classifying Thai opinions on online media using text mining. *J Sci Technol MSU*, 37(3), 372-379. (in Thai)
- TensorFlow. (2021). *tf.keras.metrics.binary_crossentropy*. https://www.tensorflow.org/api_docs/python/tf/keras/metrics/binary_crossentropy
- Thamsiri, D., & Meesad, P. (2011). Ensemble Data Classification Based on Decision Tree, Artificial Neuron Network and Support Vector Machine Optimized by Genetic Algorithm. *The Journal of KMUTNB*, 21(2), 293-303. (in Thai)

Tzanetakis, G., & Cook, P. (2002). Musical Genre Classification of Audio Signals. *IEEE Transactions on Speech and Audio Processing*, 10(5), 293-302. <https://doi.org/10.1109/TSA.2002.800560>

Zhu, S., & Chollet, F. (2020). *Working with RNNs*. https://keras.io/guides/working_with_rnn/

Research Article

การบำบัดน้ำเสียโรงงานอุตสาหกรรมฟอกย้อมโดยการตกตะกอนด้วยไฟฟ้า

Wastewater Treatments for Textile Industry by Electrocoagulation Process

ทิพย์สุรีย์ กรบุญรักษา^{1*} และ ธิดารัตน์ ชมภูราช²

Thipsuree Kornboonraksa^{1*} and Tidarad Chompurach²

ภาควิชาวิศวกรรมเคมี คณะวิศวกรรมศาสตร์ มหาวิทยาลัยบูรพา อำเภอเมือง ชลบุรี 20131 ประเทศไทย

^{1,2} Department of Chemical Engineering, Faculty of Engineering, Burapha University, Muang District, Chonburi 20131, Thailand

*E-mail: thipsuree@eng.buu.ac.th

Received: 20/04/2021; Revised: 24/07/2021; Accepted: 17/08/2021

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อศึกษาการบำบัดน้ำเสียโรงงานอุตสาหกรรมฟอกย้อมโดยการตกตะกอนด้วยไฟฟ้ากับน้ำเสียสังเคราะห์ โดยมีสภาวะที่ทำการศึกษา ได้แก่ ค่า pH ที่ 7 และ 10 ความหนาแน่นกระแสไฟฟ้าที่ 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. และระยะเวลาในการทำปฏิกิริยาตั้งแต่ 10 ถึง 60 นาที จากผลการทดลองพบว่า เมื่อน้ำเสียที่ค่า pH 10 ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. เวลา 20 นาที มีประสิทธิภาพในการกำจัดสี ความขุ่นและซีโอดีได้ดีกว่าที่ค่า pH 7 โดยมีค่าเป็น 92.6%, 69.1% และ 89.1% ตามลำดับ ในขณะที่ค่า pH 7 ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. เวลา 40 นาที มีประสิทธิภาพในการกำจัดสี ความขุ่นและซีโอดีเป็น 90.0 %, 10.7 % และ 92.6% ตามลำดับ ค่าใช้จ่ายในการบำบัดจากสภาวะที่เหมาะสมของการตกตะกอนด้วยไฟฟ้าที่ค่า pH 10 เท่ากับ 58.0 บาท/ลบ.ม. ในขณะที่ค่า pH 7 เท่ากับ 94.1 บาท อัตราการเกิดปฏิกิริยาการกำจัดสีของการตกตะกอนด้วยไฟฟ้าที่ค่า pH 7 และ 10 เป็นปฏิกิริยาอันดับสอง เมื่อเปรียบเทียบประสิทธิภาพการบำบัดน้ำเสียอุตสาหกรรมฟอกย้อมของระบบบำบัดแบบตะกอนเร่ง (Activated Sludge) กับการตกตะกอนด้วยไฟฟ้าพบว่า ระบบบำบัดแบบตะกอนเร่งมีประสิทธิภาพดีกว่าในการกำจัดสี ความขุ่นและซีโอดี แต่หากนำระบบบำบัดแบบตะกอนเร่งมาใช้ร่วมกับการตกตะกอนด้วยไฟฟ้าจะทำให้สามารถกำจัดสีของน้ำเสียได้เพิ่มขึ้น ซึ่งสามารถใช้เป็นทางเลือกหนึ่งในการเพิ่มประสิทธิภาพของระบบบำบัดน้ำเสียโรงงานอุตสาหกรรมฟอกย้อมได้

คำสำคัญ: การบำบัดน้ำเสีย โรงงานอุตสาหกรรมฟอกย้อม การตกตะกอนด้วยไฟฟ้า

Abstract

The objective of this research was to study the wastewater treatment for the textile industry by electrocoagulation process with synthetic wastewater. The experimental conditions including the initial pH of 7 and 10, the current densities of 0.01, 0.02 and 0.03 A/cm² with the reaction time from 10 to 60 min. The experimental results revealed that the initial pH of 10, the current density of 0.02 A/cm² and the reaction time of 20 min were the optimum conditions, resulting in the color, turbidity, and chemical oxygen demand (COD) removal efficiencies of 92.6%, 69.1%, and 89.1%, respectively. While the initial pH of 7, the current density of 0.01 A/cm² and the reaction time of 40 min were the optimum conditions, resulting in the color, turbidity, and chemical oxygen demand (COD) removal efficiencies of 90.0%, 10.7%, and 92.6%, respectively. The operational costs of initial pH of 10 and pH of 7 were 58.0 bath/m³ and 94.1 bath/m³, respectively. The study indicated that electrocoagulation process at an initial pH of 7 and 10 followed the second-order kinetics. It was found that the activated sludge (AS) system was capable to treat the real textile wastewater better than the electrocoagulation in terms of higher turbidity and COD removal efficiency. Nevertheless, the results suggested that the activated sludge system followed by the electrocoagulation increases the color removal efficiency of the textile industry wastewater. It can be the alternative wastewater treatment for textile industry.

Keywords: Wastewater treatment, textile industry, electrocoagulation

บทนำ

อุตสาหกรรมฟอกย้อมสิ่งทอเป็นอุตสาหกรรมหนึ่งที่มีการใช้น้ำและสารเคมีในกระบวนการผลิตในปริมาณมากและมีปริมาณน้ำเสียเกิดขึ้นสูงถึง 20,000-40,000 ลบม./เดือน รวมไปถึงมลสารต่าง ๆ ที่ปนเปื้อนอยู่ในน้ำเสียนั้นสามารถก่อให้เกิดผลกระทบต่อสิ่งแวดล้อมได้ โดยปัญหาหลักที่พบคือมีค่าสีปนเปื้อนอยู่ในปริมาณที่สูงถึง 600-4,000 ADMI (Tuba, 2017) โดยสีที่เกิดขึ้นส่วนใหญ่มาจากกระบวนการย้อมสี การพิมพ์ผ้าและการฟอกขาว สีบางชนิดสามารถบำบัดได้ด้วยวิธีการกายภาพและทางเคมีทั่วไป แต่สีบางชนิดก็ไม่สามารถบำบัดได้ด้วยวิธีการดังกล่าว ซึ่งสีย้อมที่นิยมใช้ทั่วไปนั้นมีอยู่หลายชนิด เช่น สีย้อมชนิดรีแอคทีฟ (Reactive dyes) เป็นสีย้อมที่ติดอยู่บนเส้นใยได้ด้วยพันธะโควาเลนต์ (Covalent bond) ซึ่งเป็นพันธะที่แข็งแรง ทำให้ทนทานต่อการซักล้าง สามารถย้อมได้เกือบทุกชนิด จึงทำให้สีย้อมชนิดนี้เป็นที่นิยมอย่างมาก โดยทั่วไปอุตสาหกรรมฟอกย้อมนิยมใช้ระบบบำบัดแบบตะกอนเร่ง (Activated sludge) ในการบำบัดน้ำเสียเนื่องจากมีราคาไม่แพง แต่ก็ยังคงมีปัญหาเรื่องสีในน้ำทิ้ง

ภายหลังผ่านการบำบัดทำให้น้ำทิ้งที่ผ่านออกจากระบบบำบัดไม่ได้ตามมาตรฐานที่กำหนดไว้ ทั้งนี้อาจพบปริมาณสารอินทรีย์ในรูปของซีโอดีและบีโอดีในน้ำทิ้งมีค่าที่สูงและมีปริมาณสีที่ตกค้างอยู่มาก (จริยา, 2557) ปัจจุบันเทคโนโลยีที่สามารถนำมาใช้ในการแก้ไขปัญหานี้ ได้แก่ การตกตะกอนด้วยไฟฟ้า (Electrocoagulation) ซึ่งเป็นวิธีที่ไม่มีการใช้สารเคมีในการทำปฏิกิริยาโดยตรง แต่อาศัยขั้วแอโนดและแคโทดในการเกิดปฏิกิริยาออกซิเดชัน -รีดักชันและปลดปล่อยประจุบวกอิสระที่มีความสามารถในการกำจัดมลสารในน้ำเสียได้ แต่ก็ต้องมีการควบคุมปริมาณกระแสไฟฟ้า ค่า pH เวลาในการทำปฏิกิริยาและระยะเวลาการตกตะกอนให้เหมาะสม จากผลการศึกษาการบำบัดน้ำเสียฟอกย้อมด้วยกระบวนการตกตะกอนด้วยไฟฟ้าพบว่าวิธีนี้สามารถกำจัดสีย้อมประเภท รีแอคทีฟได้มากกว่า 90% (Sengil & Özacar, 2009; Daneshvar et al., 2004) และสามารถกำจัดซีโอดีได้ในช่วง 37-76% (Arslan-Alaton et al., 2008) ดังนั้นในงานวิจัยครั้งนี้จึงทำการศึกษาเปรียบเทียบการบำบัดน้ำเสียอุตสาหกรรมฟอกย้อมสังเคราะห์ซึ่งเตรียมจากสีย้อมชนิดรีแอคทีฟโดยใช้การตกตะกอนด้วยไฟฟ้าเพื่อหาสภาวะในการบำบัดที่เหมาะสม การศึกษาจลนพลศาสตร์ของปฏิกิริยาการกำจัดสีโดยอาศัยแบบจำลองสมการปฏิกิริยาอันดับที่หนึ่งและปฏิกิริยาอันดับที่สอง รวมถึงการประเมินค่าใช้จ่ายในการบำบัดน้ำเสียเมื่อใช้กระบวนการตกตะกอนด้วยไฟฟ้า

วัสดุอุปกรณ์และวิธีการทดลอง

การเตรียมน้ำเสียสังเคราะห์

น้ำเสียสังเคราะห์ที่ใช้ในการศึกษาครั้งนี้ถูกเตรียมขึ้นให้มีคุณลักษณะใกล้เคียงกับน้ำเสียจริงจากโรงงานอุตสาหกรรมฟอกย้อม แสดงดังรูปที่ 1 โดยมีค่าซีโอดีอยู่ในช่วง 500-600 มก./ล. ทำการปรับค่า pH ของน้ำเสียสังเคราะห์ก่อนนำไปทดลองที่ค่า pH 7 และ 10 แสดงสารเคมีที่ใช้ในการเตรียมน้ำเสียสังเคราะห์เตรียมดังตารางที่ 1 และคุณลักษณะของน้ำเสียฟอกย้อมสังเคราะห์ที่เตรียมได้ดังตารางที่ 2



รูปที่ 1 น้ำเสียจากโรงงานอุตสาหกรรมฟอกย้อมแห่งหนึ่งในนิคมอุตสาหกรรมแหลมฉบัง จังหวัดชลบุรี

ตารางที่ 1 สารเคมีที่ใช้ในการเตรียมน้ำเสียฟอกย้อมสังเคราะห์ (Şahinkaya, 2013)

สารเคมี	ความเข้มข้น
สีย้อมชนิด Reactive blue dye	100 มก./ล.
แป้ง (Starch)	500 มก./ล.
โซเดียมคาร์บอเนต (Na_2CO_3)	200 มก./ล.
โซเดียมไบคาร์บอเนต (NaHCO_3)	200 มก./ล.
โซเดียมคลอไรด์ (NaCl)	300 มก./ล.
โซเดียมไฮดรอกไซด์ (NaOH)	100 มก./ล.
กรดซัลฟิวริก (H_2SO_4)	60 มก./ล.

การเลือกสีย้อมชนิด Reactive blue dye เพื่อใช้ในการเตรียมน้ำเสียฟอกย้อมสังเคราะห์ครั้งนี้เนื่องจากเป็นสีย้อมที่นิยมใช้ในอุตสาหกรรมสิ่งทอ โดยโครงสร้างทางเคมีสีย้อมชนิด Reactive blue dye มีประจุเป็นลบที่ประกอบด้วย กลุ่มพื้นฐาน 4 กลุ่ม ซึ่งสามารถแสดงเป็นโครงสร้างโดยทั่วไปได้ดังนี้ (วรรณวรรณ เทียงวรรณ กานต์, 2546)

S-D-T-X

- โดย S คือ กลุ่มที่มีความสามารถในการละลายน้ำได้สูง โดยทั่วไปจะเป็นพวกซัลโฟเนต ($-\text{SO}_2\text{Na}$) ซึ่งติดอยู่กับกลุ่มโครโมฟอร์
- D คือ กลุ่มของเคมีที่ทำให้เกิดสี เรียกว่า กลุ่มโครโมฟอร์ (Chromophore)
- T คือ กลุ่มอะตอมที่ทำหน้าที่เป็นตัวเชื่อมระหว่างกลุ่มรีแอคทีฟกับกลุ่มโครโมฟอร์
- X คือ กลุ่มรีแอคทีฟ (Reactive group) ซึ่งจะเป็นกลุ่มที่ทำให้สีทำปฏิกิริยากับกลุ่มไฮดรอกซิลในเส้นใย

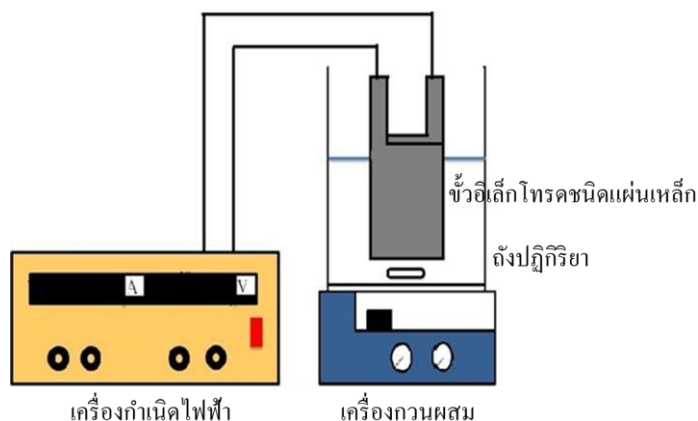
ในปัจจุบันนี้สีรีแอคทีฟ (Reactive dyes) เป็นสีย้อมที่ได้รับความนิยมในการนำมาใช้ในการย้อมเส้นใยอย่างกว้างขวาง เพราะสีรีแอคทีฟนั้นจะติดอยู่บนเส้นใยด้วยพันธะโควาเลนต์ ซึ่งเป็นพันธะที่แข็งแรง ทำให้ทนทานต่อการซักล้าง สามารถย้อมได้เกือบทุกเฉด สีที่ได้มีความสดใส ด้วยจุดเด่นนี้จึงทำให้สีรีแอคทีฟเป็นที่นิยมเพิ่มมากขึ้นเรื่อย ๆ

ตารางที่ 2 คุณลักษณะของน้ำเสียฟอกย้อมสังเคราะห์ที่เตรียมได้

คุณลักษณะ	ความเข้มข้น	
	น้ำเสียน้ำ pH เริ่มต้น 7	ค่า pH เริ่มต้น 10
ค่า pH	7.03 ± 0.02	10.03 ± 0.02
ซีโอดี (COD)	564.78 ± 70.26	602.76 ± 29.14
สี (Color)	42.80 ± 19.64	39.99 ± 10.98
ความขุ่น (Turbidity)	70.89 ± 7.00	80.23 ± 6.21
ของแข็งละลายน้ำ (TDS)	1052.20 ± 24.92	1010.41 ± 63.54
ค่าการนำไฟฟ้า (Conductivity)	2238.39 ± 52.98	2149.31 ± 135.20

การทดลองของกระบวนการตกตะกอนด้วยไฟฟ้า (Electrocoagulation)

การทดลองโดยติดตั้งขั้วอิเล็กโทรดชนิดแผ่นเหล็ก จำนวน 4 แผ่น ขนาดแผ่นละ 8.8x24 ตร.ซม. วางห่างกันแผ่นละ 1.5 ซม. ลงในถังปฏิกรณ์ขนาด 17x32x28 ลบ.ซม. ปริมาตร 9 ลิตร อัตราการกวนผสม 2000 รอบต่อนาที แสดงดังรูปที่ 2 โดยมีค่า pH ที่ใช้ในการทดลอง คือ 7 และ 10 ความหนาแน่นกระแสไฟฟ้าที่ใช้ในการทดลอง คือ 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. ระยะเวลาในการทำปฏิกิริยา คือ 0, 10, 20, 30, 40, 50 และ 60 นาที ตามลำดับ ทำการกวนน้ำเสียด้วยเครื่องกวนแม่เหล็ก (Magnetic stirrer) เมื่อสิ้นสุดการทดลองนำตัวอย่างน้ำที่บำบัดได้ไปตรวจวัดค่า pH ความขุ่น สี และซีโอดี ตามวิธีการวิเคราะห์คุณภาพน้ำของ APHA, AWWA & WEF (2017)



รูปที่ 2 การบำบัดน้ำเสียอุตสาหกรรมฟอกย้อมโดยการตกตะกอนด้วยไฟฟ้า

การวิเคราะห์จลนศาสตร์ของปฏิกิริยา (Kinetic analysis)

การวิเคราะห์จลนศาสตร์ของปฏิกิริยาเป็นการศึกษาอัตราเร็วในการเกิดขึ้นของกระบวนการ รวมไปถึงการศึกษาปัจจัยต่าง ๆ ที่มีผลต่อการเกิดปฏิกิริยา สามารถนำข้อมูลที่ได้มาใช้ในการกำหนดสภาวะที่เหมาะสมในการบำบัดน้ำเสียฟอกย้อม โดยปฏิกิริยาอันดับหนึ่งและอันดับสองนั้นเป็นปฏิกิริยาที่พบมากที่สุด ซึ่งปฏิกิริยาอันดับหนึ่ง (First-order reaction) คือ ปฏิกิริยาที่อัตราการเกิดปฏิกิริยาขึ้นอยู่กับความเข้มข้นของสารตั้งต้นยกกำลังหนึ่งดังสมการที่ 1 (Khan et al., 2010)

$$\frac{-dC}{dt} = kC \quad (1)$$

โดยที่ dC/dt คือ อัตราการเปลี่ยนแปลงความเข้มข้นของมลสาร

– คือ เครื่องหมายลบที่บ่งบอกถึงการลดลงของความเข้มข้นของมลสารที่เวลา t

k คือ ค่าคงที่อัตราการเกิดปฏิกิริยา (นาที่⁻¹)

C คือ ความเข้มข้นของมลสารที่เวลา t

ถ้าความเข้มข้นเริ่มต้นของมลสารที่เวลา $t = 0$ เป็น C_0 และที่เวลา t ต่อมา ความเข้มข้นของมลสารจะเป็น C_t โดยสามารถหาอัตราการเกิดปฏิกิริยาอันดับหนึ่ง (First order) ได้จากสมการที่ 2

$$\ln C_t = \ln C_0 - k_1 t \quad (2)$$

สำหรับปฏิกิริยาอันดับสอง (Second-order reaction) นั้นเป็นปฏิกิริยาที่อัตราการเกิดปฏิกิริยาขึ้นอยู่กับความเข้มข้นของสารตั้งต้นยกกำลังสอง ดังสมการที่ 3 ซึ่งถ้าความเข้มข้นเริ่มต้นของมลสารที่เวลา $t = 0$ เป็น C_0 และที่เวลา t ต่อมาความเข้มข้นของมลสารจะเป็น C_t โดยสามารถหาอัตราการเกิดปฏิกิริยาอันดับสอง (Second order) ได้ดังสมการที่ 4 (Al-Shannag et al., 2015)

$$\frac{-dC}{dt} = kC^2 \quad (3)$$

$$\frac{1}{C_t} = \frac{1}{C_0} + k_2 t \quad (4)$$

โดยที่ C_0 คือ ความเข้มข้นเริ่มต้นของมลสารที่เวลา $t = 0$

C_t คือ ความเข้มข้นของมลสารที่เวลา t

k_1 คือ ค่าคงที่อัตราการเกิดปฏิกิริยาอันดับหนึ่ง (นาที่⁻¹)

k_2 คือ ค่าคงที่อัตราการเกิดปฏิกิริยาอันดับสอง ((มก/ล)⁻¹.นาที่⁻¹)

t คือ เวลาที่ใช้ของปฏิกิริยา (นาที่)

การประเมินค่าใช้จ่ายจากการบำบัดน้ำเสียโดยการตกตะกอนด้วยไฟฟ้า

การบำบัดน้ำเสียโดยการตกตะกอนด้วยไฟฟ้ามีปัจจัยนำเข้าที่ถูกลำเอียงใช้ในการประเมินค่าใช้จ่าย ได้แก่ ปริมาณแผ่นอิเล็กโทรดที่ใช้ไปในการทำปฏิกิริยา ราคาของขั้วอิเล็กโทรด พลังงานที่ใช้ ค่าใช้จ่ายในการปรับค่า pH ดังแสดงในสมการที่ 4 (Koby & Demirbas, 2015)

$$OC = (\alpha \times C_{energy}) + (\beta \times C_{electrode}) + Cost_{pH} \quad (4)$$

โดยที่ OC คือ ค่าใช้จ่ายในการบำบัดน้ำเสีย (บาท/ลบ.ม.)

α คือ ค่าไฟฟ้าต่อหน่วย (บาท/กิโลวัตต์ชั่วโมง) คือ 2.35 บาท/กิโลวัตต์ชั่วโมง

C_{energy} คือ ปริมาณของพลังงานที่ถูกใช้ไป (กิโลวัตต์ชั่วโมง/ลบ.ม.)

β คือ ราคาขั้วอิเล็กโทรดที่ใช้ (แผ่นเหล็ก) (บาท/กก.) คือ 23 บาท/กก.

$C_{electrode}$ คือ ปริมาณอิเล็กโทรดที่ถูกใช้ไป (กก./ลบ.ม.)

$Cost_{pH}$ คือ ค่าใช้จ่ายในการปรับค่า pH ของน้ำเสีย (บาท/ลบ.ม.)

ปริมาณพลังงาน (Energy consumption) หาได้จากสมการที่ 5 และปริมาณอิเล็กโทรดที่ใช้ไป (Electrode consumptions) หาได้จากสมการที่ 6

$$C_{energy} = \frac{U \times I \times t_{EC}}{v} \quad (5)$$

$$C_{electrode} = \frac{I \times t_{EC} \times MW}{z \times F \times v} \quad (6)$$

โดยที่ C_{energy} คือ ปริมาณของพลังงานที่ถูกใช้ไป (กิโลวัตต์ชั่วโมง/ลบ.ม.)

U คือ ศักย์ไฟฟ้า (โวลต์)

I คือ กระแสไฟฟ้า (แอมแปร์)

t_{EC} คือ เวลาในการทำปฏิกิริยา (วินาที)

v คือ ปริมาตรของน้ำเสียในถังปฏิกิริยา (ลบ.ม.)

$C_{electrode}$ คือ ปริมาณอิเล็กโทรดที่ถูกใช้ไป (กก./ลบ.ม.)

MW คือ มวลโมเลกุลของเหล็ก (Fe) คือ 55.84 กรัม/โมล

z คือ จำนวนอิเล็กตรอนในปฏิกิริยาออกซิเดชัน/รีดักชัน ($z = 2$)

F คือ ค่าคงที่ฟาราเดย์ (Faraday's constant) คือ 96,487 คูลอมป์/โมล

ผลและวิจารณ์ผลการทดลอง

ค่า pH

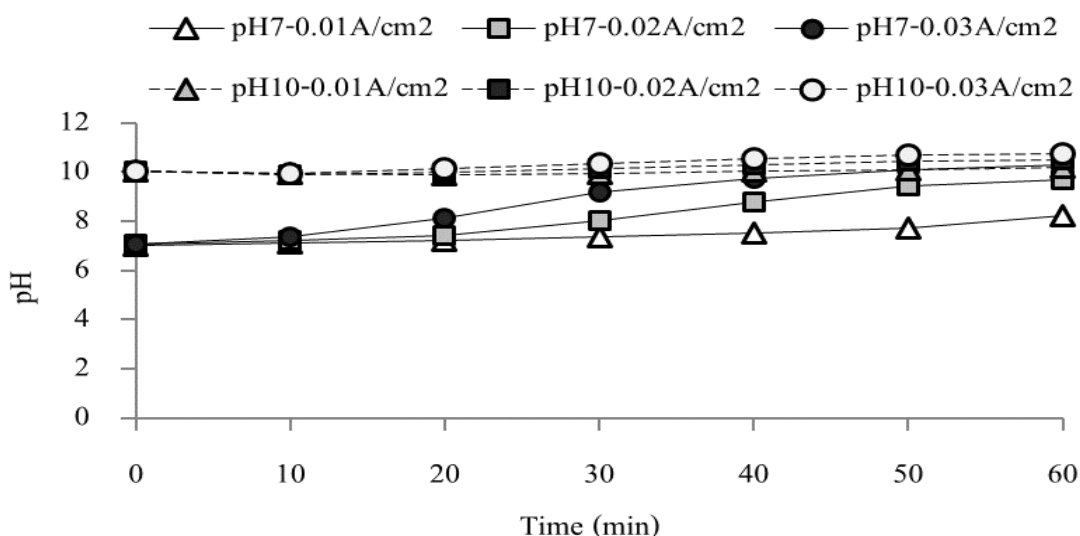
ในการทำงานของกระบวนการตกตะกอนด้วยไฟฟ้านั้นจะทำให้เกิดการเปลี่ยนแปลงของค่า pH ขึ้น โดยขึ้นอยู่กับชนิดของขั้วอิเล็กโทรดที่ใช้ และค่า pH เริ่มต้นของน้ำเสีย ในกรณีที่ใช้แผ่นเหล็กเป็นขั้วอิเล็กโทรดที่ขั้วแอโนดจะเกิดปฏิกิริยาออกซิเดชันทำให้มีไอออนของเหล็ก (Fe^{2+}) ละลายออกมาในน้ำ และที่ขั้วแคโทดเกิดปฏิกิริยารีดักชันได้เป็นไอออนของไฮดรอกไซด์ (OH^-) และก๊าซไฮโดรเจน (H_2) ปริมาณของ OH^- ที่เกิดขึ้นที่ขั้วแคโทดเกิดแสดงดังตารางที่ 3 โดยจะมีค่าเพิ่มขึ้นเมื่อความหนาแน่นกระแสไฟฟ้าและเวลาในการทำปฏิกิริยาเพิ่มขึ้น ดังนั้นจึงทำให้ผลลัพธ์ของปฏิกิริยามีค่า pH ที่สูงในน้ำทิ้งที่ผ่านบำบัดแล้ว

ตารางที่ 3 ความเข้มข้นของ OH^- ที่เกิดขึ้นเมื่อใช้การตกตะกอนด้วยไฟฟ้า

เวลา (นาที)	ความเข้มข้นของ OH^- (โมล/ลิตร)					
	ค่า pH เริ่มต้น 7			ค่า pH เริ่มต้น 10		
	0.01 แอมแปร์/ตร. ซม.	0.02 แอมแปร์/ตร. ซม.	0.03 แอมแปร์/ตร. ซม.	0.01 แอมแปร์/ตร. ซม.	0.02 แอมแปร์/ตร. ซม.	0.03 แอมแปร์/ตร. ซม.
0	1.0×10^{-7}	1.2×10^{-7}	1.1×10^{-7}	1.2×10^{-4}	1.1×10^{-4}	1.1×10^{-4}
10	1.4×10^{-7}	1.7×10^{-7}	2.4×10^{-7}	9.1×10^{-5}	8.3×10^{-5}	9.3×10^{-5}
20	1.6×10^{-7}	2.7×10^{-7}	1.3×10^{-6}	8.0×10^{-5}	9.5×10^{-5}	1.4×10^{-4}
30	2.4×10^{-7}	1.1×10^{-6}	1.6×10^{-5}	8.9×10^{-5}	1.4×10^{-4}	2.3×10^{-4}
40	3.5×10^{-7}	5.8×10^{-6}	5.7×10^{-5}	1.1×10^{-4}	1.9×10^{-4}	3.5×10^{-4}
50	5.3×10^{-7}	2.9×10^{-5}	1.3×10^{-4}	1.2×10^{-4}	2.7×10^{-4}	5.0×10^{-4}
60	1.7×10^{-6}	5.2×10^{-5}	1.9×10^{-4}	1.7×10^{-4}	3.3×10^{-4}	5.7×10^{-4}

จากปฏิกิริยารีดักชันที่ทำให้มี OH^- เกิดขึ้นในระบบแล้ว OH^- บางส่วนยังสามารถรวมตัวกับ Fe^{2+} ในน้ำเกิดเป็น $\text{Fe}(\text{OH})_2$, $\text{Fe}(\text{OH})_3$ ซึ่งเป็นตะกอนที่ไม่ละลายน้ำ และเกิดไอออนไฮดรอกไซด์ของเหล็กในรูปต่าง ๆ ทั้งนี้ขึ้นกับค่า pH ของน้ำ (สุกมาส, 2557) เช่น หากค่า pH อยู่ระหว่าง 7 ถึง 10 จะเกิดเป็น $\text{Fe}(\text{OH})^{2+}$, $\text{Fe}(\text{OH})_2^+$, $\text{Fe}(\text{OH})_4^-$ และ $\text{Fe}(\text{OH})_3$ เป็นต้น และไอออนของเหล็กในเหล่านี้อาจสามารถดูดซับ และจับเกาะสิ่งที่เป็นพิษปนเปื้อนในน้ำได้ทำให้เกิดการดูดซับและเกาะรวมตัวเป็นอนุภาคที่มีขนาดใหญ่ขึ้นและแยกตัวออกจากน้ำได้ง่าย กรณีที่ค่า pH เป็น 10 เหล็กที่เกิดขึ้นจะอยู่ในรูป $\text{Fe}(\text{OH})_3$ เกือบทั้งหมด ซึ่งไฮดรอกไซด์ในรูปดังกล่าวนี้จะมีประสิทธิภาพในการดูดซับมลสารได้ดีกว่าไฮดรอกไซด์ในรูปอื่น ๆ

จากรูปที่ 3 พบว่าเมื่อค่า pH เริ่มต้นเป็น 7 ความหนาแน่นกระแสไฟฟ้า 0.01, 0.02, และ 0.03 แอมแปร์/ตร.ซม. ค่า pH ของน้ำเสียจะมีแนวโน้มเพิ่มขึ้นเมื่อระยะเวลาในการทำปฏิกิริยาเพิ่มขึ้น เมื่อปล่อยกระแสไฟฟ้าจนครบ 60 นาที ที่ความหนาแน่นกระแสไฟฟ้า 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. จะมีค่า pH เป็น 8.2 9.7 และ 10.3 ตามลำดับ และเมื่อค่า pH เริ่มต้นเป็น 10 ความหนาแน่นกระแสไฟฟ้า 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. เมื่อปล่อยกระแสไฟฟ้าจนครบ 60 นาที จะมีค่า pH เป็น 10.2, 10.5 และ 10.8 ตามลำดับ โดยค่า pH ที่เปลี่ยนแปลงไปเมื่อมีค่า pH เริ่มต้นเป็น 7 จะมีค่า pH เพิ่มขึ้นมากอย่างเห็นได้ชัด ในขณะที่ค่า pH เริ่มต้นเป็น 10 ค่า pH จะมีแนวโน้มเพิ่มขึ้นน้อยกว่า เนื่องจากที่ค่า pH ดังกล่าวได้เข้าถึงขั้นสมดุลแล้ว



รูปที่ 3 ความสัมพันธ์ระหว่างค่า pH กับความหนาแน่นกระแสไฟฟ้าที่ระยะเวลาต่าง ๆ

ความขุ่น (Turbidity)

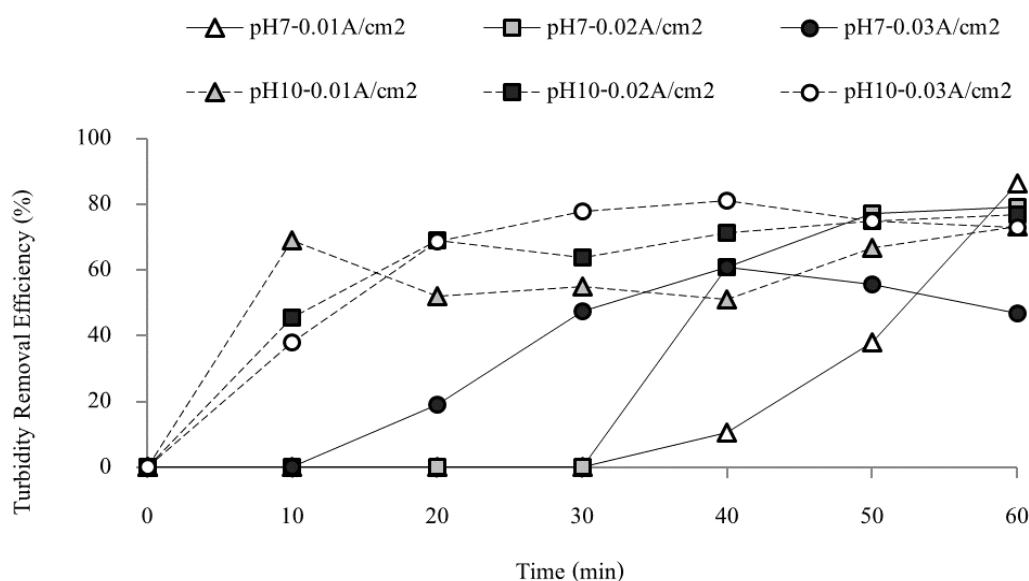
ความขุ่นในน้ำนั้นเกิดจากการที่มีสารที่ไม่ละลายน้ำแขวนลอยอยู่ เมื่อทำการบำบัดด้วยกระบวนการตกตะกอนด้วยไฟฟ้าจะมีตะกอนที่เกิดขึ้นทั้งในรูปของตะกอนที่จมตัวได้และตะกอนที่ลอยอยู่บนผิวน้ำที่เกิดจากการนำพาของก๊าซไฮโดรเจน (H_2) ทั้งนี้ความหนาแน่นกระแสไฟฟ้ามีความสำคัญต่อการควบคุมอัตราการเกิดปฏิกิริยาในกระบวนการตกตะกอนด้วยไฟฟ้า โดยความหนาแน่นกระแสไฟฟ้าที่ใช้นั้นจะเป็นตัวกำหนดปริมาณไอออนเหล็กที่ละลายออกมา เมื่อความหนาแน่นกระแสไฟฟ้าเพิ่มขึ้นจะทำให้ไอออนเหล็กละลายออกมาได้มากขึ้น และสามารถจับตัวกับที่อนุภาคปนเปื้อนในน้ำเสียเกิดเป็นตะกอนได้ดียิ่งขึ้นด้วย (Joseph & Obi, 2017) อีกทั้งความหนาแน่นกระแสไฟฟ้ามีผลต่อสถานะจลนศาสตร์ในการกำจัดความขุ่น หากใช้ความหนาแน่นกระแสไฟฟ้าสูง เวลาที่ใช้ในการบำบัดน้ำเสียก็มีแนวโน้มลดลง และยังส่งผลให้มีอัตราการสร้างฟองก๊าซ H_2 ที่เพิ่มขึ้น และขนาดฟองก๊าซ H_2 จะลดลงตามความหนาแน่นกระแสไฟฟ้าที่เพิ่มขึ้น ซึ่งนับว่าเป็นประโยชน์ต่อการกำจัดความขุ่นโดยอาศัยหลักการลอยของก๊าซ H_2 และนำพาอนุภาคความขุ่นขึ้นสู่ผิวน้ำก่อนกำจัดด้วยวิธีการตักกวาดออกไป (Abeer & Nemr, 2012)

จากรูปที่ 4 เมื่อใช้การตกตะกอนด้วยไฟฟ้าบำบัดน้ำเสียฟอกย้อมสังเคราะห์ พบว่าหากมีค่า pH เริ่มต้นเป็น 7 ในช่วง 10-20 นาทีแรก ระบบไม่สามารถกำจัดความขุ่นลงไปได้ เนื่องจากอยู่ในช่วงของการรวมและสร้างตะกอน (Coagulation and flocculation) ระหว่างอนุภาคคอลลอยด์ในน้ำกับประจุบวกของ Fe^{2+} ที่ถูกกระตุ้นออกมาจากขั้วอิเล็กโทรดด้วยกระแสไฟฟ้า ทำให้ตะกอนที่ถูกทำลายเสถียรภาพแล้วสามารถจับตัวกันได้จนมีขนาดใหญ่ขึ้นแต่ยังแขวนลอยอยู่ในถังปฏิกิริยา อย่างไรก็ตามความขุ่นจะมีค่าลดลงเนื่องจากตะกอนที่รวมตัวกันจนมีขนาดใหญ่สามารถตกตะกอนลงไปที่ก้นถังปฏิกิริยาได้ และเมื่อมีค่า pH เริ่มต้นเป็น 10 พบว่าที่กระแสไฟฟ้าต่าง ๆ จะกำจัดความขุ่นได้เร็วตั้งแต่ที่เวลา 10 นาที เมื่อเวลาเพิ่มขึ้น การกำจัดความขุ่นจะเพิ่มขึ้นและคงที่ โดยที่เวลา 20 นาที จะเห็นว่าความขุ่นนั้นจะลดลงไปได้มากและที่เวลาเพิ่มขึ้นความขุ่นจะเริ่มคงที่ ซึ่งการเพิ่มขึ้นของเวลาจะทำให้ค่า pH ของน้ำเสียมียค่าเพิ่มขึ้นจนเข้าใกล้ขั้นสมดุลของการเกิดตะกอนเฟอร์ริกไฮดรอกไซด์ ($Fe(OH)_3$) โดยตะกอนดังกล่าวนี้จะเกิดมากขึ้นที่ค่า pH สูง ๆ เนื่องจากตะกอนเฟอร์ริกไฮดรอกไซด์ ($Fe(OH)_3$) จะอยู่ในรูปของแข็ง (Solid form) และมีคุณสมบัติในการดูดซับมลสารอนุภาคต่าง ๆ ในน้ำเสียได้ดีกว่าไฮดรอกไซด์ในรูปอื่นที่อยู่ในรูปไอออน จึงทำให้ค่าความขุ่นถูกดูดซับและตกตะกอนลงได้ดีกว่า กรณีที่น้ำเสียที่มีค่า pH เริ่มต้นเป็น 10 นั้นเป็นค่า pH ที่สามารถทำให้เกิด $Fe(OH)_3$ มากขึ้น จึงสามารถกำจัดความขุ่นไปได้มากและรวดเร็วตั้งแต่ช่วงแรกของการเดินระบบ

เมื่อมีค่า pH เริ่มต้นเป็น 7 ความหนาแน่นกระแสไฟฟ้า 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. ที่เวลา 60 นาที สามารถกำจัดค่าความขุ่นไปได้ 46.3% 79.2% และ 46.8% ตามลำดับ และเมื่อค่า pH เริ่มต้นเป็น 10 พบว่าความหนาแน่นกระแสไฟฟ้า 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. เมื่อใช้เวลาเพิ่มขึ้น ประสิทธิภาพการกำจัดความขุ่นจะเพิ่มขึ้นด้วย โดยที่ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. จะกำจัดความขุ่นได้สูงสุด 73.2% ที่เวลา 60 นาที ที่ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. จะกำจัดความขุ่นได้สูงสุด 76.7% ที่เวลา 60 นาที และที่

ความหนาแน่นกระแสไฟฟ้า 0.03 แอมแปร์/ตร.ซม. จะกำจัดความขุ่นได้สูงสุด 81.0% ที่เวลา 40 นาที โดยที่ค่า pH 10 สามารถกำจัดความขุ่นได้ดีกว่าค่า pH 7 ทั้งนี้เมื่อพิจารณาข้อมูลจากตารางที่ 3 พบว่าปริมาณของ OH^- ที่ pH 10 มีค่ามากกว่าที่ pH 7 ที่ความหนาแน่นกระแสไฟฟ้าและเวลาในการทำปฏิกิริยาเดียวกัน จึงส่งผลทำให้เกิดตะกอนเฟอร์ริกไฮดรอกไซด์ ($\text{Fe}(\text{OH})_3$) ที่อยู่ในรูปของแข็ง (Solid form) ที่ pH 10 มากกว่า pH 7 และมีประสิทธิภาพในการกำจัดความขุ่นที่ดีกว่านั่นเอง

จากการสังเกตปฏิกิริยาทางเคมีที่เกิดขึ้นระหว่างการบำบัดพบว่าเมื่อปล่อยกระแสไฟฟ้าและมีไอออนของเหล็กละลายออกมานั้นจะเริ่มสังเกตเห็นตะกอนขนาดเล็กเกิดขึ้นในน้ำเสีย เมื่อเวลาเพิ่มขึ้นตะกอนเริ่มมีขนาดใหญ่ขึ้น ซึ่งตะกอนบางส่วนจะจมอยู่บริเวณด้านล่างและบางส่วนจะลอยขึ้นบริเวณผิวน้ำโดยแก๊สไฮโดรเจน ทำให้สามารถแยกตะกอนออกจากน้ำเสียได้ง่าย ความขุ่นในน้ำจึงมีค่าลดลง นอกจากนี้ยังพบอีกว่าในการปล่อยน้ำเสียให้ตกตะกอนที่เวลา 30 นาทีนั้น ในน้ำเสียจะยังคงมีตะกอนขนาดเล็กอีกส่วนหนึ่งยังคงแขวนลอยอยู่ แต่หากมีการเพิ่มเวลาในการตกตะกอนมากขึ้นก็สามารถที่จะลดค่าความขุ่นในน้ำลงไปอีกได้



รูปที่ 4 ประสิทธิภาพการกำจัดความขุ่น ที่ค่า pH 7 และ 10 กับความหนาแน่นกระแสไฟฟ้าและระยะเวลาต่างกัน

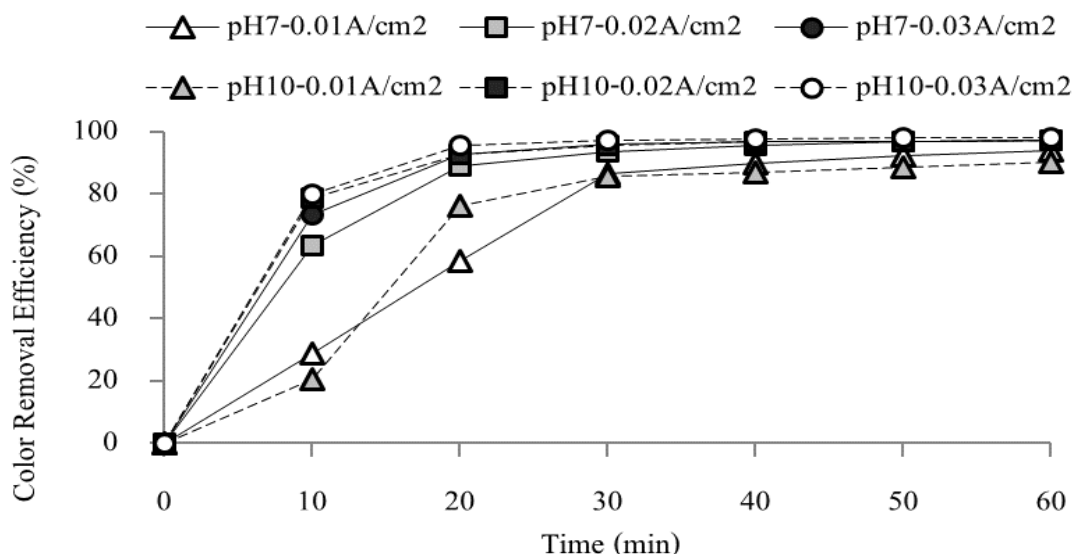
ในกรณีของการเติมแป้ง (Starch) ที่เป็นสารเคมีที่ใช้ในการเตรียมน้ำเสียฟอกข้อมสังเคราะห์ พบว่าแป้งสามารถสร้างกลไกในการกำจัดอนุภาคคอลลอยด์ได้ด้วยการเป็นสะพานเชื่อมต่ออนุภาคคอลลอยด์ (Polymer Bridging) ทั้งนี้การเกาะติดอาจเป็นผลมาจากประจุที่แตกต่างกันของแป้งและอนุภาคคอลลอยด์ หรือเป็นแรงทางปฏิกิริยาเคมีที่เกิดขึ้นระหว่างประจุที่เหมือนกันของแป้งและคอลลอยด์ (รัชฎาวรรณ, 2540) ทั้งนี้แป้งสามารถทำ

หน้าที่เป็นสารสร้างตะกอนหรือโคแอกกูแลนต์ร่วมในการกำจัดคอลลอยด์ได้ (มันสิน, 2542) ซึ่งทำให้สามารถกำจัดความขุ่นได้ดีขึ้น

สี (Color)

ในการกำจัดสีที่ปนเปื้อนในน้ำเสียโดยการตกตะกอนด้วยไฟฟ้าพบว่ากลไกการกำจัดสีนั้นเกิดขึ้นได้ 2 แบบ คือ การตกตะกอนและการดูดซับบนผิวของสารประกอบไฮดรอกไซด์ของโลหะที่เกิดขึ้น การจะเกิดกลไกในการกำจัดสีแบบใดนั้นขึ้นอยู่กับค่า pH ของน้ำเสีย หากน้ำเสียมีค่า pH มากกว่า 6.5 ส่วนใหญ่จะเกิดการดูดซับโดยการตะกอนสารประกอบไฮดรอกไซด์โลหะ ในระหว่างกระบวนการตกตะกอนด้วยไฟฟ้าโมเลกุลสีที่มีประจุลบจะถูกทำให้เป็นกลางด้วยประจุบวกของเหล็กออกไซด์ที่ถูกปล่อยออกมาจากขั้วแอโนดโทรมเหล็ก ทำให้โมเลกุลสีสามารถรวมตัวกันเป็นฟล็อกที่มีขนาดใหญ่ขึ้นและตกตะกอนลงมาได้ ส่วนกลุ่มอนุภาคที่เกาะรวมตัวกันและมีประจุบวกจะเกิดจากการดูดซับของสีในรูปฟุ้งที่เป็นโมเนอ์เมอร์และโพลิเมอร์ (ศุภมาส, 2557) นอกจากนี้การเพิ่มขึ้นของกระแสไฟฟ้านั้นจะส่งผลต่อการเกิดสารรวมตะกอน $\text{Fe}(\text{OH})_3$ มีผลให้ประสิทธิภาพการบำบัดเพิ่มขึ้นและสามารถบำบัดได้เร็วขึ้น กรณีที่ค่า pH มากกว่า 7 พบว่าเป็นค่า pH ที่เหมาะสมต่อการเกิดสารรวมตะกอน $\text{Fe}(\text{OH})_3$ เนื่องจากไฮดรอกไซด์ไอออนที่เกิดจากขั้วแคโทดจะทำปฏิกิริยากับ $\text{Fe}^{2+}(\text{aq})$ หรือ $\text{Fe}^{3+}(\text{aq})$ เกิดเป็น $\text{Fe}(\text{OH})_2(\text{s})$ หรือ $\text{Fe}(\text{OH})_3(\text{s})$ (พลกฤษณ์ และคณะ, 2559) ที่ค่า pH 10 จึงบำบัดได้ดีกว่าที่ค่า pH 7 โดยใช้เวลาในการบำบัดน้อยกว่าที่ความหนาแน่นกระแสไฟฟ้าเดียวกัน ทั้งนี้ผลการทดลองสอดคล้องกับการศึกษาของ Phalakornkule et al. (2010) และ Sengil & Özacar (2009) ที่พบว่าน้ำเสียในช่วงค่า pH สูงนั้นจะมีประสิทธิภาพการกำจัดสีได้ดีกว่า โดยจะใช้เวลาในการกำจัดสีน้อยกว่าเมื่อน้ำเสียมีค่า pH ต่ำ และเมื่อใช้ความหนาแน่นกระแสไฟฟ้าเพิ่มขึ้นก็จะใช้เวลาในการกำจัดสีน้อยลงด้วย

จากผลการทดลองพบว่าประสิทธิภาพการกำจัดสีจะเพิ่มขึ้นเมื่อระยะเวลาในการทำปฏิกิริยาเพิ่มขึ้น ดังรูปที่ 5 เมื่อน้ำเสียมีค่า pH เริ่มต้นเป็น 7 ที่ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. พบว่าการกำจัดสีจะเกิดขึ้นอย่างช้า ๆ และเมื่อเพิ่มความหนาแน่นกระแสไฟฟ้าเป็น 0.02 และ 0.03 แอมแปร์/ตร.ซม. การกำจัดสีจะเกิดขึ้นอย่างรวดเร็วในช่วงแรกของปฏิกิริยา หลังจากนั้นถึงแม้จะเพิ่มระยะเวลาในการทำปฏิกิริยาแต่การกำจัดสีจะเริ่มเข้าสู่สภาวะคงที่ ในขณะที่เมื่อค่า pH เริ่มต้นเป็น 10 ที่ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. พบว่าการกำจัดสีจะเกิดขึ้นในช่วงแรกของปฏิกิริยาและการกำจัดสีจะเพิ่มขึ้นเมื่อระยะเวลาการทำปฏิกิริยาเพิ่มมากขึ้น และเมื่อความหนาแน่นกระแสไฟฟ้าเป็น 0.02 และ 0.03 แอมแปร์/ตร.ซม. การกำจัดสีจะเกิดขึ้นอย่างรวดเร็วในช่วงแรกของปฏิกิริยา หลังจากนั้นประสิทธิภาพในการกำจัดสีจะเริ่มคงที่เช่นเดียวกันกับที่ค่า pH 7



รูปที่ 5 ประสิทธิภาพการกำจัดสี ที่ค่า pH 7 และ 10 กับความหนาแน่นกระแสไฟฟ้าและระยะเวลาต่างกัน

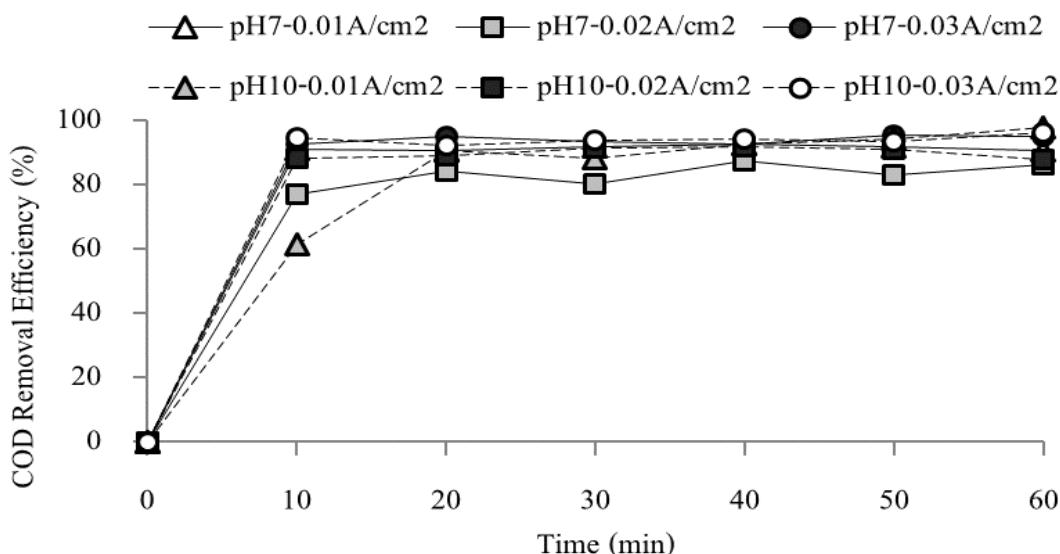
จากรูปที่ 5 พบว่าเมื่อระยะเวลาในการทำปฏิกิริยาเพิ่มขึ้น ประสิทธิภาพในการกำจัดสีจะเพิ่มขึ้น และการกำจัดสีจะเริ่มคงที่เมื่อเวลาผ่านไปประมาณ 30 นาที โดยที่ค่า pH เริ่มต้นเป็น 7 ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. เวลา 10-60 นาที จะกำจัดสีได้ 28.8%-94.1% ที่ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. เวลา 10-60 นาที กำจัดสีได้ 63.5%-97.1% และที่ความหนาแน่นกระแสไฟฟ้า 0.03 แอมแปร์/ตร.ซม. เวลา 10-60 นาที กำจัดสีได้ 73.2%-97.3% โดยน้ำเสียจากที่มีสีน้ำเงินเข้มจะเปลี่ยนเป็นสีเขียวเข้มอมน้ำตาล และมีสีอ่อนลงเรื่อยๆ ในกรณีที่ค่า pH เริ่มต้นเป็น 10 ที่ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. เวลา 10-60 นาที กำจัดสีได้ 20.7%-90.4% ที่ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. กำจัดสีได้ 78.5%-97.1% และที่ความหนาแน่นกระแสไฟฟ้า 0.03 แอมแปร์/ตร.ซม. กำจัดสีได้ 79.8%-98.2% ตามลำดับ เมื่อผ่านการบำบัด น้ำเสียจากที่มีสีน้ำเงินเข้มจะเริ่มเปลี่ยนเป็นสีเขียวอมน้ำตาลและใสขึ้น จากลักษณะสีของน้ำเสียที่มีการเปลี่ยนแปลงไปนั้น อาจเป็นไปได้ว่าเมื่อใช้เหล็กเป็นขั้วอิเล็กโทรดในกระบวนการตกตะกอนด้วยไฟฟ้า หากปล่อยกระแสไฟฟ้าไปที่ขั้วเหล็ก แผ่นเหล็กจะถูกกัดกร่อนได้เป็นไอออนของเหล็กละลายออกมา ซึ่งในน้ำเสียจะมีไอออนของเหล็กในรูปแบบต่างๆ ปะปนกันอยู่ รวมถึง $\text{Fe}(\text{OH})_2$ โดย $\text{Fe}(\text{OH})_2$ นี้จะมีสีเขียวหรือสีเขียวมืดเมื่อถูกออกซิเดชันโดยอากาศ และเมื่อได้รับออกซิเจนมากเกินไปจะได้เป็น Hydrated ferric oxide หรือ Ferric hydroxide ซึ่งมีสีส้มหรือแดงอมน้ำตาล (Moreno et al., 2009) เมื่อสารประกอบไฮดรอกไซด์ที่ได้เกิดการดูดซับกับโมเลกุลของสีย้อมซึ่งมีสีน้ำเงิน จึงทำให้สีของน้ำเสียมีลักษณะเปลี่ยนไปเป็นสีเขียวอมน้ำตาลดังที่ได้จากการทดลอง โมเลกุลของสีที่ถูกดูดซับจะเกิดเป็นตะกอนแยกออกจากน้ำเสียไป เมื่อเวลาเพิ่มขึ้นสีที่ปนเปื้อนจึงลดน้อยลง น้ำเสียที่ได้จึงมีลักษณะใสขึ้น

กรณีที่น้ำเสียสังเคราะห์ที่ pH 10 บำบัดได้ดีกว่าที่ pH 7 นั้น เนื่องจากน้ำเสียที่ค่า pH มากกว่า 7 เป็น pH ที่เหมาะสมต่อการเกิดสารรวมตะกอน ($\text{Fe}(\text{OH})_3$) เนื่องจากไฮดรอกไซด์ไอออนที่เกิดจากขี้แควโทจะทำปฏิกิริยากับ $\text{Fe}^{2+}(\text{aq})$ หรือ $\text{Fe}^{3+}(\text{aq})$ เกิดเป็น $\text{Fe}(\text{OH})_2(\text{s})$ หรือ $\text{Fe}(\text{OH})_3(\text{s})$ (พลกฤษณ์ และคณะ, 2559) นอกจากนี้ผลการทดลองยังสอดคล้องกับการศึกษาของ Phalakornkule et al. (2010) และ Şengil and Özacar (2009) ที่พบว่า น้ำเสียในช่วงค่า pH สูงนั้นจะมีประสิทธิภาพการกำจัดได้ดีกว่า โดยจะใช้เวลาในการกำจัดสั้นน้อยกว่าเมื่อน้ำเสียมีค่า pH ต่ำ และเมื่อใช้ความหนาแน่นกระแสไฟฟ้าเพิ่มขึ้น ก็จะใช้เวลาในการกำจัดสั้นลงด้วย

ในการเลือกสภาวะเดินระบบบำบัดที่เหมาะสมจะพิจารณาจากการกำจัดค่าสีในน้ำทิ้งให้ได้ต่ำกว่า 300 ADMI ซึ่งเป็นไปตามค่ามาตรฐานน้ำทิ้งที่ระบายออกจากโรงงานได้กำหนดไว้ (กระทรวงอุตสาหกรรม, 2560) สำหรับสภาวะที่เหมาะสมในการบำบัดหากน้ำเสียมีค่า pH เริ่มต้น 7 คือ ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. ที่เวลา 30 นาที มีค่าสีเท่ากับ 184.4 ADMI ที่ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. ใช้เวลา 20 นาที มีค่าสีเท่ากับ 192.4 ADMI และที่ความหนาแน่นกระแสไฟฟ้า 0.03 แอมแปร์/ตร.ซม. ใช้เวลา 20 นาที มีค่าสีเท่ากับ 139.3 ADMI ในกรณีที่น้ำเสียมีค่า pH เริ่มต้น 10 พบว่าสภาวะที่เหมาะสม คือ ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. ที่เวลา 20 นาที มีค่าสีเท่ากับ 209.4 ADMI ที่ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. ใช้เวลา 10 นาที มีค่าสีเท่ากับ 254.4 ADMI และที่ความหนาแน่นกระแสไฟฟ้า 0.03 แอมแปร์/ตร.ซม. ใช้เวลา 10 นาที มีค่าสีเท่ากับ 251.2 ADMI

ซีโอดี (COD)

ค่าซีโอดีเป็นค่าใช้ในการบ่งชี้ถึงความสกปรกของน้ำเสียที่มีการปนเปื้อนของสารเคมีต่าง ๆ รวมถึงอนุภาคคอลลอยด์ที่ก่อให้เกิดความขุ่นและโมเลกุลจากสีย้อม หากสามารถกำจัดค่าความขุ่นและสีออกจากน้ำเสียได้โดยกลไกการสร้างตะกอน การรวมตะกอนและการตกตะกอน ก็จะทำให้การปนเปื้อนต่าง ๆ ในน้ำทิ้งมีน้อยลง เมื่อพิจารณาในเชิงประสิทธิภาพของการกำจัดค่าซีโอดีก็จะพบว่ามีความเพิ่มมากขึ้น นอกจากนี้ในปฏิกิริยามีตะกอนไฮดรอกไซด์ของเหล็กเกิดขึ้นซึ่งตะกอนที่เกิดขึ้นนี้จะมีประสิทธิภาพในการดูดซับอนุภาคต่าง ๆ ในน้ำเสียได้เช่นกัน โดยประสิทธิภาพในการกำจัดอนุภาคปนเปื้อนต่าง ๆ ในน้ำเสียนั้นขึ้นอยู่กับปริมาณเฟอร์ริกไฮดรอกไซด์ที่เกิดขึ้นจากปฏิกิริยาตามกฎฟาราเดย์ ซึ่งการเกิดปฏิกิริยาจะแปรผันตรงกับกระแสไฟฟ้าหรือความหนาแน่นกระแสไฟฟ้าที่ใช้ เมื่อความหนาแน่นกระแสไฟฟ้าเพิ่มมากขึ้นทำให้เกิดปริมาณเฟอร์ริกไฮดรอกไซด์มากขึ้น มลสารต่าง ๆ ก็สามารถถูกดูดซับและกำจัดออกได้มากขึ้น เมื่อมลสารในน้ำถูกกำจัดไปทั้งในรูปความขุ่นและสีของน้ำ จึงส่งผลให้ค่าซีโอดีลดลงด้วยเช่นกัน ดังแสดงผลการทดลองในรูปที่ 6



รูปที่ 6 ประสิทธิภาพในการกำจัดซีโอดี ค่า pH 7 และ 10 กับความหนาแน่นกระแสไฟฟ้าที่ระยะเวลาต่าง ๆ

จากรูปที่ 6 พบว่าที่ค่า pH 7 และ 10 สามารถกำจัดค่าซีโอดีได้ค่อนข้างเร็วตั้งแต่ช่วงแรกของการทำปฏิกิริยา และจะเริ่มคงที่เมื่อเวลาผ่านไป 20 นาที ทั้งนี้การเพิ่มขึ้นของความหนาแน่นกระแสไฟฟ้าจะทำให้ประสิทธิภาพการกำจัดค่าซีโอดีเพิ่มขึ้น จากผลการทดลองดังรูปที่ 8 พบว่าการตกตะกอนด้วยไฟฟ้าสามารถกำจัดค่าซีโอดีได้ดีและมีค่าไม่เกินมาตรฐานคุณภาพน้ำที่กำหนด คือ 400 มก./ล. กรณีที่น้ำเสียมีค่า pH 7 ที่ความหนาแน่นกระแสไฟฟ้า 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. ค่าซีโอดีลดลงอย่างรวดเร็วในช่วง 10 นาทีแรก และที่เวลา 20 นาที ประสิทธิภาพในการกำจัดซีโอดีมีแนวโน้มค่อนข้างคงที่ในทุก ๆ ค่าความหนาแน่นกระแสไฟฟ้า คิดเป็น 90.4%, 84.5% และ 95.2% ตามลำดับ สำหรับกรณีที่น้ำเสียมีค่า pH 10 ที่ความหนาแน่นกระแสไฟฟ้า 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. สามารถกำจัดค่าซีโอดีได้ในช่วง 10 นาทีแรกเช่นเดียวกัน และที่เวลา 20 นาที ประสิทธิภาพในการกำจัดซีโอดีมีแนวโน้มคงที่ในทุก ๆ ค่าความหนาแน่นกระแสไฟฟ้า คิดเป็น 90.7%, 89.1% และ 92.1% ตามลำดับ โดยจะเห็นว่าทั้งน้ำเสียที่ค่า pH 7 และ 10 สามารถกำจัดค่าซีโอดีได้ดี ซึ่งการเพิ่มขึ้นของความหนาแน่นกระแสไฟฟ้าที่ 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. จะทำให้ค่า pH ในน้ำเสียมีแนวโน้มสูงขึ้น การที่น้ำเสียมีค่า pH สูงขึ้นนั้นจะมีผลต่อการเกิดเฟอริกไฮดรอกไซด์ที่สามารถดูดซับมลสารต่างๆ ในน้ำเสีย และการเพิ่มขึ้นของความหนาแน่นกระแสไฟฟ้าจะทำให้เกิดปริมาณเฟอริกไฮดรอกไซด์มากขึ้น มลสารต่างๆ ก็สามารถถูกดูดซับและกำจัดออกได้มากขึ้น จึงทำให้ค่าซีโอดีลดลง และเนื่องจากเมื่อน้ำเสียมีค่า pH สูงขึ้นจนใกล้หรือเข้าสู่จุดสมมูล การกำจัดซีโอดีจึงมีแนวโน้มที่จะคงที่

อัตราการเกิดปฏิกิริยาในการกำจัดสี (Kinetic)

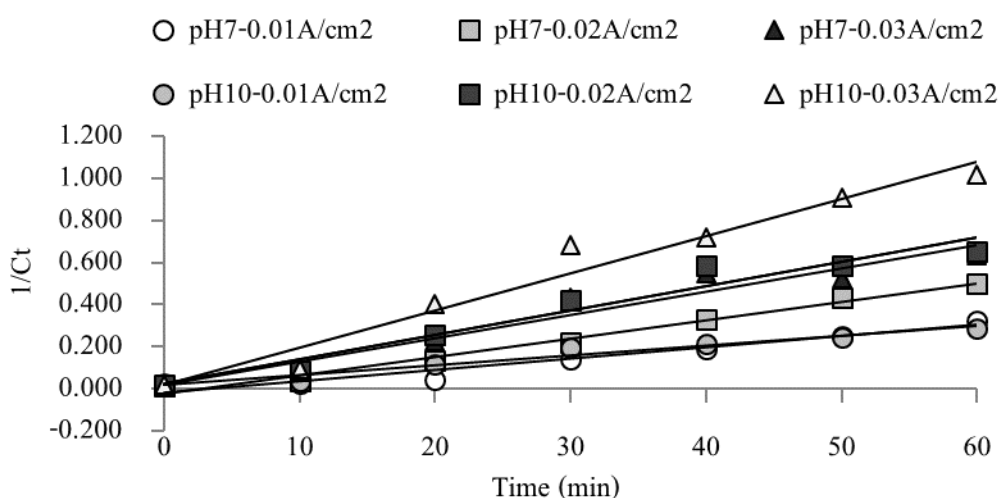
การศึกษ้อัตราการเกิดปฏิกิริยา (Kinetic) ในการกำจัดสี พิจารณาจากความเข้มข้นสีในน้ำเสีย ค่า pH เริ่มต้น ความหนาแน่นกระแสไฟฟ้า และระยะเวลาในการทำปฏิกิริยา ทั้งนี้แบบจำลองสมการอัตราการเกิดปฏิกิริยาของ ปฏิกิริยาอันดับหนึ่ง (First order) และปฏิกิริยาอันดับสอง (Second order) ใช้ในการอธิบายอัตราการเกิดปฏิกิริยาการ กำจัดสี สำหรับอัตราการเกิดปฏิกิริยาอันดับหนึ่ง ทำการพล็อตกราฟเส้นตรงระหว่างค่า $\ln C_t$ กับ t และปฏิกิริยา อันดับสอง ทำการพล็อตกราฟเส้นตรงระหว่างค่า $\frac{1}{C_t}$ กับ t จะได้ค่าคงที่ k_1 และ k_2 จากความชัน ตามลำดับ แสดงดัง ตารางที่ 4

ตารางที่ 4 ค่าคงที่อัตราและค่า R-Square ของอัตราการเกิดปฏิกิริยาอันดับหนึ่งและอันดับสองในการกำจัดสี

pH	ความหนาแน่นกระแสไฟฟ้า (แอมแปร์/ตร.ซม.)	First-order Reaction		Second-order Reaction	
		k_1 (นาที ⁻¹)	R^2	k_2 (มก./ล.·นาที ⁻¹)	R^2
7	0.01	0.0513	0.9505	0.0054	0.9558
	0.02	0.0583	0.9003	0.0087	0.9851
	0.03	0.0570	0.8081	0.0111	0.9373
10	0.01	0.0412	0.8603	0.0046	0.9591
	0.02	0.0544	0.8005	0.0115	0.9514
	0.03	0.0618	0.7890	0.0177	0.9642

กรณีที่มีน้ำเสียมีค่า pH 7 และ 10 ที่ความหนาแน่นกระแสไฟฟ้า 0.01, 0.02 และ 0.03 แอมแปร์/ตร.ซม. สามารถใช้สมการอัตราการเกิดปฏิกิริยาอันดับสอง (Second order) ในการอธิบายได้ทุกสภาวะ โดยดูจากค่า R-Square ที่มีค่ามากกว่า 0.9 และจากค่า k_2 จะเห็นว่า ค่า k_2 จะเพิ่มขึ้นเมื่อค่ากระแสไฟฟ้าเพิ่มขึ้น ดังรูปที่ 7 ที่แสดงความสัมพันธ์ระหว่างค่า $\frac{1}{C_t}$ กับ t ที่มีแนวโน้มมีความเป็นเส้นตรง หมายถึงหากความหนาแน่นกระแสไฟฟ้าเพิ่มขึ้น อัตราการกำจัดสีก็จะเกิดได้เร็วขึ้นด้วย ซึ่งอัตราการเกิดปฏิกิริยาอันดับสองนั้นหมายถึงสารตั้งต้นสองตัวมีผลต่ออัตราการเกิดปฏิกิริยา โดยในการศึกษาครั้งนี้ปัจจัยที่มีผลต่อการเกิดปฏิกิริยา ได้แก่ ค่า pH เริ่มต้นของน้ำเสีย และความหนาแน่นกระแสไฟฟ้า ให้ความเข้มข้นของน้ำเสียเป็นค่าคงที่ ซึ่งค่า pH เริ่มต้นของน้ำเสียกับความหนาแน่นกระแสไฟฟ้ามีผลต่อความเข้มข้นของไอออนเหล็กหรือปริมาณไอออนเหล็กที่เกิดขึ้นเมื่อปล่อยกระแสไฟฟ้า หากค่า pH มีค่ามากจะทำให้ไอออนเหล็กละลายออกมาในปริมาณมาก และเมื่อเพิ่มความหนาแน่นกระแสไฟฟ้าจะทำให้ไอออนเหล็กละลายออกมาได้มากเช่นกัน เมื่อในปฏิกิริยามีปริมาณไอออนเหล็กมาก ไอออนเหล่านี้นี้จะทำปฏิกิริยา

กับโมเลกุลของสีได้มาก ทำให้กำจัดสีออกจากน้ำเสียได้มากขึ้น อัตราการกำจัดสีที่ได้จึงเพิ่มขึ้น ในขณะที่การศึกษาของ Sengil & Özacar (2009) ได้ศึกษาการกำจัดสีย้อมชนิด Reactive blue dye โดยการตกตะกอนด้วยไฟฟ้าและใช้ขั้วอิเล็กโทรดชนิดแผ่นเหล็ก พบว่าเป็นจลนพลศาสตร์ของปฏิกิริยาอันดับหนึ่ง โดยมีค่าสัมประสิทธิ์สหสัมพันธ์ R^2 ที่ได้จากการพล็อตกราฟใกล้เคียง 1 ซึ่งเมื่อเวลาในการทำปฏิกิริยาเพิ่มขึ้นจะทำให้ความเข้มข้นของสีลดลงเมื่อเทียบกับเวลา ทั้งนี้ผลการวิเคราะห์ข้อมูลที่ได้มีความแตกต่างกันเนื่องด้วยความแตกต่างของสภาวะต่าง ๆ ที่ใช้ในการทดลอง เช่น ความเข้มข้นของสี ค่า pH ความหนาแน่นกระแสไฟฟ้า เป็นต้น



รูปที่ 7 ความสัมพันธ์ระหว่าง $1/C_t$ กับ t ของการเกิดปฏิกิริยาอันดับสอง

การประเมินค่าใช้จ่ายในการบำบัด (Operational Cost)

สำหรับการประเมินค่าใช้จ่ายในการบำบัดน้ำเสียโดยกระบวนการตกตะกอนด้วยไฟฟ้านั้นจะขึ้นอยู่กับตัวแปรของปริมาณกระแสไฟฟ้าที่ใช้ ระยะเวลาในการทำปฏิกิริยา ค่าใช้จ่ายในการปรับค่า pH ของน้ำเสีย และปริมาณของแผ่นอิเล็กโทรดเหล็กที่ใช้ หากใช้ความหนาแน่นกระแสไฟฟ้า และระยะเวลาในการทำปฏิกิริยาเพิ่มขึ้น ค่าใช้จ่ายในการบำบัดก็จะเพิ่มขึ้นด้วย จากสภาวะที่เหมาะสมที่ได้จากประสิทธิภาพการกำจัดสีพบว่าหากค่า pH เริ่มต้นเป็น 7 ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร.ซม. ที่เวลา 40 นาที มีค่าใช้จ่ายในการบำบัดเท่ากับ 94.1 บาท/ลบ.ม. หากค่า pH เริ่มต้นเป็น 10 ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. ที่เวลา 20 นาที มีค่าใช้จ่ายในการบำบัดเท่ากับ 58.0 บาท/ลบ.ม. เมื่อเปรียบเทียบประสิทธิภาพและค่าใช้จ่ายในการบำบัด พบว่าที่ค่า pH เริ่มต้นเป็น 10 ใช้ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. นาน 20 นาที เป็นสภาวะที่เหมาะสมที่สุด ซึ่งที่สภาวะดังกล่าวให้ประสิทธิภาพในการบำบัดที่ดี ใช้เวลาน้อยกว่า และมีค่าใช้จ่ายในการบำบัดที่ต่ำกว่าค่า pH 7 เนื่องจากที่ค่า pH เริ่มต้น เป็น 7 ใช้เวลาในการบำบัดมากกว่า อีกทั้งยังมีค่าใช้จ่ายในการปรับค่า pH ของน้ำเสียด้วย

ในสภาวะการณ้จริงน้ำเสียจากโรงงานอุตสาหกรรมฟอกย้อมมีค่า pH ใกล้เคียง 10 ซึ่งนับเป็นสภาวะที่เหมาะสมกับการเดินระบบด้วยการตกตะกอนด้วยไฟฟ้าโดยไม่ต้องปรับค่า pH ในน้ำเสียก่อนการบำบัด อันจะทำให้ลดค่าใช้จ่ายในการปรับค่า pH ก่อนการบำบัดลงได้

อย่างไรก็ตามจากการศึกษาเปรียบเทียบประสิทธิภาพการบำบัดน้ำเสียฟอกย้อมจริงของการตกตะกอนด้วยไฟฟ้าในสภาวะที่เหมาะสมที่ค่า pH 10 กับระบบบำบัดน้ำเสียแบบตะกอนเร่งซึ่งเป็นระบบบำบัดที่ใช้กับน้ำเสียโรงงานอุตสาหกรรมฟอกย้อมในรูปที่ 2 พบว่าระบบบำบัดน้ำเสียแบบตะกอนเร่ง (Activated Sludge, AS) มีประสิทธิภาพมากกว่าการตกตะกอนด้วยไฟฟ้า (Electrocoagulation, EC) โดยระบบบำบัดน้ำเสียแบบตะกอนเร่งมีประสิทธิภาพในการกำจัดความขุ่น สี และซีโอดีอยู่ที่ 70.4%, 93.4% และ 62.1% ตามลำดับ แต่หากมีการนำน้ำทิ้งจากระบบบำบัดน้ำเสียแบบตะกอนเร่งมาใช้ร่วมกับการตกตะกอนด้วยไฟฟ้าจะทำให้สีจริงของน้ำทิ้งภายหลังการกรองมีคุณภาพที่ดีขึ้น โดยมีประสิทธิภาพในการกำจัดสีเพิ่มขึ้นเป็น 98.1% ดังแสดงในรูปที่ 8 (ข)



(ก) การบำบัดด้วยระบบ AS



(ข) การบำบัดด้วยระบบ AS ตามด้วย EC

รูปที่ 8 เปรียบเทียบสีของน้ำเสียภายหลังการบำบัดด้วยระบบ AS และการบำบัดด้วยระบบ AS ตามด้วย EC

สรุปผลการทดลอง

1. สภาวะที่เหมาะสมของการบำบัดน้ำเสียที่ค่า pH 10 คือ ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร. ซม. เวลา 20 นาที มีประสิทธิภาพในการกำจัดสี ความขุ่น และซีโอดีเป็น 92.6%, 69.1% และ 89.1% ตามลำดับ โดยค่าใช้จ่ายในการบำบัดจากสภาวะที่เหมาะสมของการตกตะกอนด้วยไฟฟ้าที่ค่า pH 10 เท่ากับ 58.0 บาท/ลบ.ม.
2. สภาวะที่เหมาะสมของการบำบัดน้ำเสียที่ค่า pH 7 คือ ความหนาแน่นกระแสไฟฟ้า 0.01 แอมแปร์/ตร. ซม. เวลา 40 นาที มีประสิทธิภาพในการกำจัดสี ความขุ่น และซีโอดีเป็น 90.0%, 10.7% และ 92.6% ตามลำดับ โดยค่าใช้จ่ายในการบำบัดจากสภาวะที่เหมาะสมของการตกตะกอนด้วยไฟฟ้าที่ค่า pH 10 เท่ากับ 58.0 บาท/ลบ.ม. ในขณะที่ค่า pH 7 เท่ากับ 94.1 บาท/ลบ.ม.

3. น้ำเสียค่า pH เริ่มต้น 10 ใช้ความหนาแน่นกระแสไฟฟ้า 0.02 แอมแปร์/ตร.ซม. นาน 20 นาที เป็นสภาวะที่เหมาะสมที่สุด ซึ่งที่สภาวะดังกล่าวให้ประสิทธิภาพในการบำบัดที่ดี ใช้เวลาน้อยกว่า และมีค่าใช้จ่ายในการบำบัดที่ต่ำกว่าที่ค่า pH 7

4. จากการศึกษาอัตราการเกิดปฏิกิริยา (Kinetic) การกำจัดสีโดยใช้กระบวนการตกตะกอนด้วยไฟฟ้าทั้งที่ค่า pH 7 และ 10 พบว่า อัตราการเกิดปฏิกิริยาการกำจัดสีเป็นปฏิกิริยาอันดับสอง และจากค่า k นั้นแสดงให้เห็นว่าเมื่อเพิ่มความหนาแน่นกระแสไฟฟ้าจะทำให้อัตราการกำจัดสีเกิดได้เร็วขึ้น

5. ระบบบำบัดแบบตะกอนเร่งมีประสิทธิภาพมากกว่าการตกตะกอนด้วยไฟฟ้า โดยมีประสิทธิภาพในการกำจัดความขุ่น สี และซีโอดีอยู่ที่ 70.4%, 93.4% และ 62.1% ตามลำดับ ทั้งนี้เมื่อใช้ระบบบำบัดแบบตะกอนเร่งร่วมกับการตกตะกอนด้วยไฟฟ้าจะทำให้สีของน้ำทิ้งหลังการบำบัดดีขึ้น ซึ่งสามารถนำระบบบำบัดน้ำเสียแบบตะกอนเร่งมาใช้ร่วมกับการตกตะกอนด้วยไฟฟ้าจะทำให้สีจริงของน้ำทิ้งภายหลังการกรองมีคุณภาพที่ดีขึ้น

เอกสารอ้างอิง

Ministry of Industry. (2017). *Royal Gazette*, Volume 134, Special Issue 153. pp. 11-15.

Yimrattanabovorn, C. (2014). *Lifespan of shale subsurface flow constructed wetland for the treatment of textile wastewater* [Unpublished master's thesis]. Suranaree University of Technology. (in thai)

Jitto, P., Chaikrathang, T., & Nakbanpote, W. (2016). Silk Textile Wastewater Treatment by Electrocoagulation Process. *Journal of Research Unit on Science, Technology and Environment for Learning*, 7(2), 228-239. (in thai)

Tanthulwet, M. (1999). *Water supply engineering*. (3rd ed.). Bangkok: Chulalongkorn university. (in thai)

Panathampon, R. (1997). *Feasibility study on high turbidity and poorly settleable raw water for water supply production by coagulation* [Unpublished master's thesis]. Kasetsart University. (in thai)

Danwittayakul, S. (2014). Electrocoagulation in Wastewater Treatment. *Technology material journal*, 90, 35-41. (in thai)

Abeer, A. M., & Nemr, A. E. (2012). Electro-Coagulation for Textile Dyes Removal. In Ahmed, E.N. (Eds.), *Non-Conventional Textile Waste Water Treatment*, 161-204, Nova Science Publishers.

- Al-Shannag, M., Al-Qodah, Z., Bani-Melhem, K., Qtaishat, M. R., & Alkasrawi, M. (2015). Heavy metal ions removal from metal plating wastewater using electrocoagulation: Kinetic study and process performance. *Chemical Engineering Journal*, 260, 749-756.
- APHA, AWWA, & WEF. (2017). *Standard methods for the examination of water and wastewater* (23rd ed.). Washington, DC: American Public Health Association.
- Arslan-Alaton, I., Kabdaşlı, I., Hanbaba, D., & Kuybu, E. (2008). Electrocoagulation of a real reactive dyebath effluent using aluminum and stainless steel electrodes. *Journal of Hazardous Materials*, 150(1), 166-173.
- Daneshvar, N., Ashassi Sorkhabi, H., & Kasiri, M. B. (2004). Decolorization of dye solution containing Acid Red 14 by electrocoagulation with a comparative investigation of different electrode connections. *Journal of Hazardous Materials*, 112(1), 55-62.
- Joseph, T. N., & Obi, C. C. (2017). Abattoir Wastewater Treatment by Electrocoagulation Using Iron Electrodes. *Der Chemica Sinica*, 8(2), 254-260.
- Khan, H., Ahmad, N., Yasar, A., & Shahid, R. (2010). Advanced Oxidative Decolorization of Red CI-5B: Effects of Dye Concentration, Process Optimization and Reaction Kinetics. *Polish journal of environmental studies*, 19(1), 83-92.
- Kobyas, M., & Demirbas, E. (2015). Evaluations of operating parameters on treatment of can manufacturing wastewater by electrocoagulation. *Journal of Water Process Engineering*, 8, 64-74.
- Moreno C. H. A., Cocke, D. L., Gomes, J. A. G., Morkovsky, P., Parga, J. R., Peterson, E., & Garcia, C. (2009). Electrochemical reactions for electrocoagulation using iron electrodes. *Industrial and Engineering Chemistry Research*, 48(4), 2275-2282.
- Phalakornkule, C., Polgumhang, S., Tongdaung, W., Karakat, B., & Nuyut, T. (2010). Electrocoagulation of blue reactive, red disperse and mixed dyes, and application in treating textile effluent. *Journal of Environmental Management*, 91(4), 918-926.
- Şahinkaya, S. (2013). COD and color removal from synthetic textile wastewater by ultrasound assisted electro-Fenton oxidation process. *Journal of Industrial and Engineering Chemistry*, 19(2), 601-605.
- Şengil, İ. A., & Özacar, M. (2009). The decolorization of C.I. Reactive Black 5 in aqueous solution by electrocoagulation using sacrificial iron electrodes. *Journal of Hazardous Materials*, 161(2), 1369-1376.
- Toprak, T., & Anis, P. (2017). Textile industry's environmental effects and approaching cleaner production and sustainability, an overview. *Journal of Textile Engineering & Fashion Technology*, 2(4), 429-442.

Research Article

สถานะธาตุอาหาร ในดิน ใบยางพารา และปริมาณธาตุอาหารในน้ำยางพาราภายใต้การปลูกพืชร่วมยางพาราที่แตกต่างกัน

Nutrient status in rubber growing soil, rubber leaf and nutrient concentration in latex under different rubber-based intercropping plantation

ปรานชญ์ แซ่เต็ง¹ จำปิ่น อ่อนทอง¹ และขวัญตา ขาวมี^{1*}

Pranchalee Saeteng¹, Jumpen Onthong¹ and Khwunta Khawmee^{1*}

¹สาขาวิชานวัตกรรมเกษตรและการจัดการ คณะทรัพยากรธรรมชาติ มหาวิทยาลัยสงขลานครินทร์ วิทยาเขตหาดใหญ่ จังหวัดสงขลา 90110 ประเทศไทย

¹Agricultural Innovation and Management Division, Faculty of Natural Resources Prince of Songkla University, Hat Yai Campus, Songkhla, 90110, Thailand

*E-mail: khwunta.k@psu.ac.th

Received: 23/04/2021; Revised: 04/08/2021; Accepted: 22/09/2021

บทคัดย่อ

ปัจจุบันมีการปลูกพืชร่วมยางพาราเพิ่มขึ้น แต่การศึกษาเกี่ยวกับสถานะธาตุอาหารในสวนยางพาราที่ปลูกพืชร่วมยังมีน้อย จึงได้ทำการเก็บตัวอย่างดิน ใบยางพารา และน้ำยางพาราในแปลงที่ปลูกไม้ร่วมยางพารา ผักเหลียงร่วมยางพารา สละร่วมยางพารา และยางพาราเชิงเดี่ยวบริเวณใกล้เคียงแปลงปลูกพืชร่วมยางพารามาศึกษาสถานะธาตุอาหาร เพื่อใช้เป็นแนวทางในการจัดการสวนยางพาราที่ปลูกพืชร่วมให้เหมาะสมสำหรับยางพาราและพืชร่วมพบว่า ปริมาณธาตุอาหารในดิน ใบยางพารา และน้ำยางพาราของแปลงปลูกยางพาราเชิงเดี่ยวและปลูกพืชร่วมยางพาราทั้งสามชนิดไม่มีความแตกต่างกันทางสถิติ โดยในดินและใบยางพารามีสถานะธาตุไนโตรเจน โพแทสเซียม และแมกนีเซียมอยู่ในระดับต่ำ ในขณะที่ฟอสฟอรัสที่เป็นประโยชน์และแคลเซียมที่สกัดได้อยู่ในระดับปานกลางถึงสูง สำหรับดินปลูกยางพารา อย่างไรก็ตาม แปลงปลูกพืชร่วมยางพารามีปริมาณฟอสฟอรัสที่เป็นประโยชน์ แคลเซียมและแมกนีเซียมที่สกัดได้มีแนวโน้มสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว ซึ่งสอดคล้องกับปริมาณฟอสฟอรัสและแมกนีเซียมในใบยางพารา และแมกนีเซียมในน้ำยางพารา และยังพบว่าฟอสฟอรัสที่เป็นประโยชน์ในดินมีความสัมพันธ์กับความเข้มข้นของฟอสฟอรัสในใบยางพารา ($r = 0.74^{**}$) จากผลการทดลอง

ชี้ให้เห็นว่าการปลูกพืชร่วมยางพาราช่วยเพิ่มฟอสฟอรัสที่เป็นประโยชน์ แคลเซียมและแมกนีเซียมที่สกัดได้ในดิน แต่เนื่องจากไนโตรเจน โพแทสเซียม แมกนีเซียมในดิน และใบยางพาราในส่วนที่ปลูกพืชร่วมยางพารายังมีปริมาณต่ำ ดังนั้น ควรเพิ่มอัตราปุ๋ยให้แก่ยางพาราและใส่ปุ๋ยตามค่าวิเคราะห์ดินหรือใบยางพารา โดยเฉพาะ ธาตุไนโตรเจน โพแทสเซียม และแมกนีเซียม

คำสำคัญ : สถานะธาตุอาหาร พืชร่วมยางพารา ดินปลูกยางพารา ใบยางพารา น้ำยางพารา

Abstract

At present, there is increased rubber-based intercropping plantation. However, the studies of nutrient status in rubber-based intercropping plantation are limited. The rubber growing soil, rubber leaves and latex in different rubber-based intercrops; rubber-bamboo, rubber-Phak-liang, rubber-Sala and rubber monocropping close to rubber-based intercropping were collected to study nutrient status for optimal fertilizer use management. The results showed that nutrient concentrations of rubber growing soil, rubber leaves and latex in rubber monocropping was not significantly different with rubber-based intercropping. Nitrogen, potassium and magnesium status of rubber growing soil and rubber leaves were low. While available phosphorus and extractable calcium concentrations of rubber growing soil and rubber leaves were medium to high. However, available phosphorus, extractable calcium and magnesium of rubber-based intercropping tended to be higher than rubber monocropping, which corresponded to phosphorus and magnesium concentrations in rubber leaves and magnesium in latex. Available phosphorus of rubber growing soil was significantly correlated with phosphorus concentrations of rubber leaves ($r = 0.74^{**}$). The results indicates that rubber-based intercropping increases available phosphorus, extractable calcium and magnesium concentrations in rubber growing soil. However, nitrogen, potassium and magnesium concentrations in rubber growing soil and rubber leaves remain low. Therefore, fertilizer application based on soil or leaf analysis should be practiced, especially nitrogen, potassium and magnesium based fertilizers.

Keywords: Nutrient status, Rubber-base intercropping, Rubber growing soil, Rubber leaf, latex

บทนำ

ประเทศไทยผลิตและส่งออกผลิตภัณฑ์ยางพารารายใหญ่ของโลก โดยใน พ.ศ. 2562 ประเทศไทยสามารถผลิตยางพาราได้จำนวน 4.77 ล้านตัน และสามารถส่งออกได้จำนวน 4.03 ล้านตัน ซึ่งสามารถทำรายได้เข้าประเทศได้ปีละกว่า 300,000 ล้านบาท (Rubber Intelligence Unit, 2020) ยางพาราเป็นพืชที่สามารถเจริญเติบโตได้ดีในดินที่มีความอุดมสมบูรณ์สูงหรือ ปานกลาง โดยดินปลูกยางพาราแต่ละแหล่งจะมีปริมาณธาตุอาหารแตกต่างกัน จึงส่งผลให้สถานะธาตุอาหารในใบยางพาราแตกต่างกัน (Kungpisadan, 2009) การประเมินสถานะธาตุอาหารเป็นการบอกถึงระดับความเข้มข้นของธาตุอาหารในดินหรือพืชว่ามีปริมาณต่ำ ปานกลาง หรือสูงเมื่อเปรียบเทียบกับระดับวิกฤต (Rubber Research Institute of Thailand, 2007) เพื่อเป็นแนวทางในการจัดการธาตุอาหารในดินให้เหมาะสมต่อการเจริญเติบโตของยางพารา มีรายงานว่า ดินปลูกยางพาราของประเทศไทยมีปริมาณไนโตรเจนทั้งหมด (น้อยกว่า 1.1 กรัม/กิโลกรัม) ฟอสฟอรัสที่เป็นประโยชน์ (น้อยกว่า 11 มิลลิกรัม/กิโลกรัม) และโพแทสเซียมที่สกัดได้ (น้อยกว่า 40 มิลลิกรัม/กิโลกรัม) อยู่ในระดับต่ำถึงต่ำมาก (Kungpisadan et al., 2013) และมีการศึกษาสถานะธาตุอาหารในใบยางพาราในพื้นที่ลุ่มและที่ดอน พบว่า ไนโตรเจน (น้อยกว่า 33 กรัม/กิโลกรัม) และโพแทสเซียม (น้อยกว่า 13 กรัม/กิโลกรัม) ต่ำกว่าระดับที่เหมาะสม (Onthong et al., 2016) อีกทั้ง มีผลการศึกษาของ Suchartkul (2011) ที่ศึกษาสถานะธาตุอาหารในดินและใบยางพาราของสวนยางพาราในจังหวัดชุมพร สุราษฎร์ธานี และนครศรีธรรมราช พบว่า ไนโตรเจน ฟอสฟอรัส และโพแทสเซียมอยู่ในระดับต่ำ จากผลการศึกษาดังกล่าวจะเห็นได้ว่า ส่วนใหญ่ธาตุไนโตรเจน ฟอสฟอรัส และโพแทสเซียมทั้งในดินและใบยางพาราอยู่ในระดับต่ำ และธาตุอาหารบางส่วนสูญเสียไปกับผลผลิตนี้ยางพารา โดยมีรายงานว่า น้ำยางพารา 1 ตัน จะสูญเสียไนโตรเจน 20 กิโลกรัม ฟอสฟอรัส 5 กิโลกรัม โพแทสเซียม 25 กิโลกรัม แคลเซียม 4 กิโลกรัม และแมกนีเซียม 5 กิโลกรัม (Kungpisadan, 2009) ดังนั้น จึงต้องมีการเพิ่มธาตุอาหารให้แกดิน โดยการใส่ปุ๋ยเคมีและปุ๋ยอินทรีย์

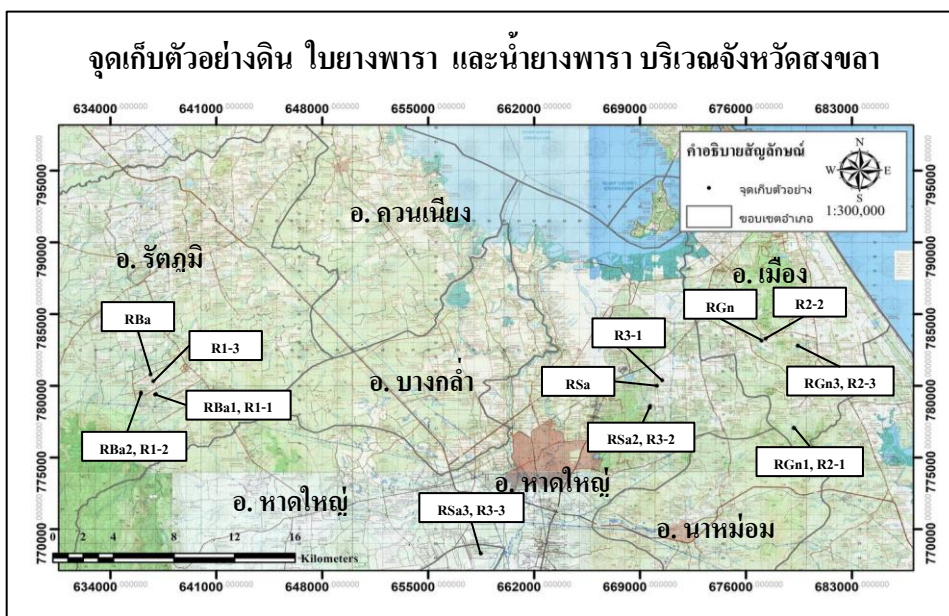
การปลูกพืชร่วมยางพาราเป็นวิธีการหนึ่งที่เพิ่มธาตุอาหารให้แกดิน ผลการศึกษของการปลูกถั่วมูน่ากล้วย และมันสำปะหลังร่วมกับยางพาราส่งผลให้ปริมาณไนโตรเจนและอินทรีย์วัตถุในดินสูงขึ้น (Kaewmapa et al., 2014) และมีการศึกษาการปลูกฝั ผักเหลียง และตะเคียนร่วมกับยางพารา พบว่า การปลูกผักเหลียงร่วมยางพารามีแนวโน้มเพิ่มธาตุอาหารในดิน (Saeateaw, 2019) อีกทั้งมีการศึกษาผลของการปลูกกาแฟอาราบิก้า โกโก้ มะเดื่อ และพะยูง มีผลทำให้ปริมาณอินทรีย์วัตถุและไนโตรเจนในดินเพิ่มขึ้น (Chen et al., 2017) ปัจจุบันประเทศไทยมีการปลูกพืชร่วมยางพาราเพิ่มขึ้นและมีการศึกษาเกี่ยวกับพืชร่วมยางพาราแต่นักศึกษาในด้านของความเป็นไปได้ที่จะปลูกร่วมกับยางพาราอย่างไรก็ตาม การศึกษาในด้านของปริมาณธาตุอาหารและสถานะธาตุอาหารในสวนยางพาราที่ปลูกพืชร่วมยางพารายังมีน้อย

เพื่อศึกษาสมบัติของดินและสถานะธาตุอาหารในดิน ไบยางพารา และปริมาณธาตุอาหารในน้ำยางพารา ในสวนที่ปลูกพืชร่วมยางพาราแตกต่างกัน เพื่อเป็นแนวทางในการเพิ่มธาตุอาหารในสวนยางพาราที่ปลูกพืชร่วม

วิธีดำเนินการวิจัย

1. การเลือกแปลงศึกษา

เลือกแปลงยางพาราที่เปิดกรีดแล้ว (7-30 ปี) และมีการปลูกพืชร่วมยางพารา ได้แก่ ไม้ร่วมยางพารา (3 ปี) (RBa) พักเหียงร่วมยางพารา (RGn) (13 ปี) สละอินโดร่วมยางพารา (RSa) (7 ปี) และยางพาราเชิงเดี่ยว (R) ที่อยู่บริเวณใกล้เคียงและเป็นดินที่เหมือนกับแปลงปลูกพืชร่วมยางพารา เนื้อดินส่วนใหญ่จัดอยู่ในกลุ่มเนื้อหยาบ ภายในอำเภอเมือง หาดใหญ่ และรัตภูมิ จังหวัดสงขลา (ภาพที่ 1) อย่างละ 3 แปลง รวมแปลงที่ใช้ในการศึกษาทั้งหมด 18 แปลง และทำการเก็บตัวอย่างดิน ไบยางพารา (มีนาคม-เมษายน 2562) และน้ำยางพารา (กรกฎาคม 2562) เพื่อนำมาวิเคราะห์ธาตุอาหาร



รูปที่ 1 พิกัดของแปลงภายในจังหวัดสงขลาที่ใช้สำหรับเก็บตัวอย่างดิน ไบยางพารา และน้ำยางพารา

หมายเหตุ : R คือ แปลงปลูกยางพาราเชิงเดี่ยวที่อยู่ใกล้กับแปลงปลูกพืชร่วมยางพารา, RBa คือ แปลงปลูกไม้ร่วมยางพารา, RGn คือ แปลงปลูกพักเหียงร่วมยางพารา และ RSa คือแปลงปลูกสละร่วมยางพารา

2. การเก็บตัวอย่างดินและวิเคราะห์สมบัติเคมีและธาตุอาหารในดิน

ทำการเก็บตัวอย่างดินที่ความลึก 0-30 และ 30-60 เซนติเมตร แบบทำลายโครงสร้างดินโดยสุ่มเก็บแบบ X-shape ทั้งหมด 9 จุด และนำดินผสมคลุกเคล้ากัน จากนั้นนำตัวอย่างดินผึ่งให้แห้งในที่ร่ม บดและร่อนผ่านตะแกรงขนาดช่องเปิด 2 มิลลิเมตร และนำตัวอย่างดินไปวิเคราะห์สมบัติเคมีและธาตุอาหาร จำนวน 3 ซ้ำ ได้แก่ ค่าพีเอช (ดิน:น้ำ = 1:5 w/v) ค่าสภาพนำไฟฟ้าของดิน (ดิน:น้ำ = 1:5 w/v) อินทรีย์วัตถุ (ใช้ดินที่ร่อนผ่านตะแกรงขนาดช่องเปิด 0.5 มิลลิเมตร และวิเคราะห์โดย Walkley and Black method) ไนโตรเจนทั้งหมด (Kjeldahl method) และฟอสฟอรัสที่เป็นประโยชน์ (Bray II method) ส่วนโพแทสเซียม แคลเซียม และแมกนีเซียมที่สกัดได้ ทำโดยการนำดินมาสกัดด้วยแอมโมเนียมอะซิเตท (NH_4OAc) 1 โมลาร์ พีเอช 7 และวัดด้วยเครื่อง atomic absorption spectrophotometer (AAS) ยี่ห้อ Thermo รุ่น iCE 3000 Series จากประเทศอเมริกา โดยโพแทสเซียมใช้หลักการของ atomic emission ส่วนแคลเซียมและแมกนีเซียมใช้หลักการของ atomic absorption (Onthong & Poonpakdee, 2019)

3. การเก็บตัวอย่างใบยางพาราและวิเคราะห์ธาตุอาหารในใบยางพารา

เก็บใบยางพาราอายุ 100-150 วันหลังจากผลิใบใหม่ และเก็บจากตำแหน่งคู่ล่างหรือใบที่ 1 และ 2 ของฉัตรแรกของกิ่งในร่มระหว่างแถว โดยเลือกต้นยางพาราที่ใกล้จุดเก็บดิน เก็บแปลงละ 9 ต้น ต้นละ 4-6 ใบ ในระยะก่อนใส่ปุ๋ย นำตัวอย่างทำความสะอาดโดยการใช้ผ้าสะอาดชุบน้ำพอลหมาด ๆ เช็ดให้ทั่ว และนำไปอบที่อุณหภูมิ 70 องศาเซลเซียส ประมาณ 72 ชั่วโมง เมื่อตัวอย่างแห้งจึงนำไปบดให้ละเอียดด้วยเครื่องบดตัวอย่างพืช ผ่านตะแกรงขนาด 20 เมช และวิเคราะห์ปริมาณธาตุอาหาร จำนวน 3 ซ้ำ ได้แก่ ไนโตรเจนทั้งหมด (Kjeldahl method) ฟอสฟอรัสทั้งหมด (yellow molybdovanadophosphoric acid method) ส่วนโพแทสเซียม แคลเซียม และแมกนีเซียมทั้งหมด ทำโดยการนำตัวอย่างใบมาย่อยด้วยกรดผสมไนตริกและเพอร์คลอริก (HNO_3 ; HClO_4 ; 3:1 v/v) และวัดด้วยเครื่อง atomic absorption spectro-photometer (AAS) ยี่ห้อ Thermo รุ่น iCE 3000 Series จากประเทศอเมริกา โดยโพแทสเซียมใช้หลักการของ atomic emission ส่วนแคลเซียมและแมกนีเซียมใช้หลักการของ atomic absorption (Onthong & Poonpakdee, 2019)

4. การเก็บตัวอย่างน้ำยางพาราและวิเคราะห์ธาตุอาหารในน้ำยางพารา

เก็บน้ำยางพาราจากต้นยางพาราที่เก็บตัวอย่างใบยางพารา เก็บต้นละ 1 มิลลิลิตร (ประมาณ 15 หยด) จำนวน 9 ต้น โดยใช้หลักปลายแหลมเจาะกิ่งกลางใต้รอยกรีดลงมาประมาณ 5 เซนติเมตร แล้วรองรับน้ำยางพาราด้วยหลอดแก้วที่แช่อยู่ในภาชนะที่บรรจุน้ำแข็ง นำน้ำยางพาราที่เก็บได้จาก 9 ต้น ผสมให้เข้ากัน และนำตัวอย่างน้ำยางพารามาวิเคราะห์ธาตุอาหารจำนวน 3 ซ้ำ ได้แก่

ไนโตรเจนทั้งหมดโดยย่อยด้วยกรดซัลฟิวริก (98% w/w) 15 มิลลิลิตร ที่อุณหภูมิ 380 องศาเซลเซียส จนเนื้อยางสลายตัวหมด (สารละลายตัวอย่างเป็นสีดำ) จากนั้นเติมกรดเพอร์คลอริก (70% w/w) 5 มิลลิลิตร และย่อยต่อที่อุณหภูมิ 180 องศาเซลเซียส จนสารละลายตัวอย่างเปลี่ยนเป็นสีเขียวอมฟ้า จากนั้นกรองสารละลายที่ได้จากการ

ย่อยผ่านกระดาษกรองวัดแมนเบอร์ 1 และนำมาวัดไนโตรเจนทั้งหมด (Kjeldahl method) (Onthong & Poonpakdee, 2019)

ฟอสฟอรัส โพแทสเซียม แคลเซียมและแมกนีเซียมทั้งหมดย่อยด้วยกรดไนตริก (65% w/w) 10 มิลลิลิตร ที่อุณหภูมิ 80 องศาเซลเซียส จนกระทั่งควันสีน้ำตาลหมดไป จากนั้นเติมกรดผสมไนตริก-เพอร์คลอริก (HNO_3 : HClO_4 ; 3:1 v/v) 10 มิลลิลิตร ย่อยต่อที่อุณหภูมิ 200 องศาเซลเซียส จนกระทั่งปรากฏควันสีขาว และเติมกรดผสมไนตริก-เพอร์คลอริก (HNO_3 : HClO_4 ; 3:1 v/v) ลงไปอีก 10 มิลลิลิตร และย่อยต่อจนสารละลายตัวอย่างใสไม่มีสี จากนั้นนำสารละลายที่ได้จากการย่อยโดยไม่ต้องกรองมาวัดธาตุอาหารเช่นเดียวกับการวิเคราะห์ธาตุอาหารไนโบรียม (Onthong & Poonpakdee, 2019)

5. ประเมินสถานะธาตุอาหารในดินและใบยางพารา

นำผลการวิเคราะห์ธาตุอาหารในดิน และใบยางพารามาประเมินสถานะธาตุอาหาร โดยการนำผลวิเคราะห์ธาตุอาหารในดินและใบยางพารามาเปรียบเทียบกับระดับธาตุอาหารในดินปลูกยางพาราและใบยางพาราของสถาบันวิจัยยางในประเทศไทย (Kungpisdan, 2011) และประเมินระดับความเข้มข้นของธาตุอาหารในดินและใบยางพาราว่ามีปริมาณต่ำ ปานกลาง หรือสูง ตามเกณฑ์ธาตุอาหารที่เพียงพอสำหรับยางพารา (Rubber Research Institute of Thailand, 2007)

6. การวิเคราะห์สถิติ

นำข้อมูลสมบัติทางเคมีและปริมาณธาตุอาหารในดิน ใบยางพารา และน้ำยางพาราในแต่ละปีซึ่งรวมยางพาราและแปลงยางพาราเชิงเดี่ยวมาหาค่าเฉลี่ย และเปรียบเทียบค่าเฉลี่ยระหว่างสองกลุ่มด้วยวิธี Independent t-test และวิเคราะห์หาค่าความแปรปรวนทางเดียวด้วยวิธี ANOVA และเปรียบเทียบความแตกต่างระหว่างปีซึ่งรวมยางพาราโดยวิธี DMRT และหาความสัมพันธ์ระหว่างปริมาณธาตุอาหารในดิน ใบยางพารา และน้ำยางพารา โดยวิธี Pearson's correlation ที่ระดับความเชื่อมั่น 95 เปอร์เซ็นต์ โดยใช้โปรแกรมคอมพิวเตอร์สำเร็จรูป SPSS version 23

ผลการทดลอง

1. พืชและค่าสภาพนำไฟฟ้า และอินทรีย์วัตถุในดินปลูกยางพารา

พืชและค่าสภาพนำไฟฟ้าของแปลงปลูกยางพาราเชิงเดี่ยวและแปลงปลูกพืชร่วมยางพาราทั้งสองระดับความลึกดินไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$) โดยพืชเฉลี่ยเท่ากับ 4.5-5.4 และค่าสภาพการนำไฟฟ้าเฉลี่ยเท่ากับ 0.01-0.03 เดซิซีเมนส์/เมตร ส่วนปริมาณอินทรีย์วัตถุของแปลงปลูกยางพาราเชิงเดี่ยวและพืชร่วมยางพาราไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$) และมีปริมาณลดลงตามระดับความลึกดิน

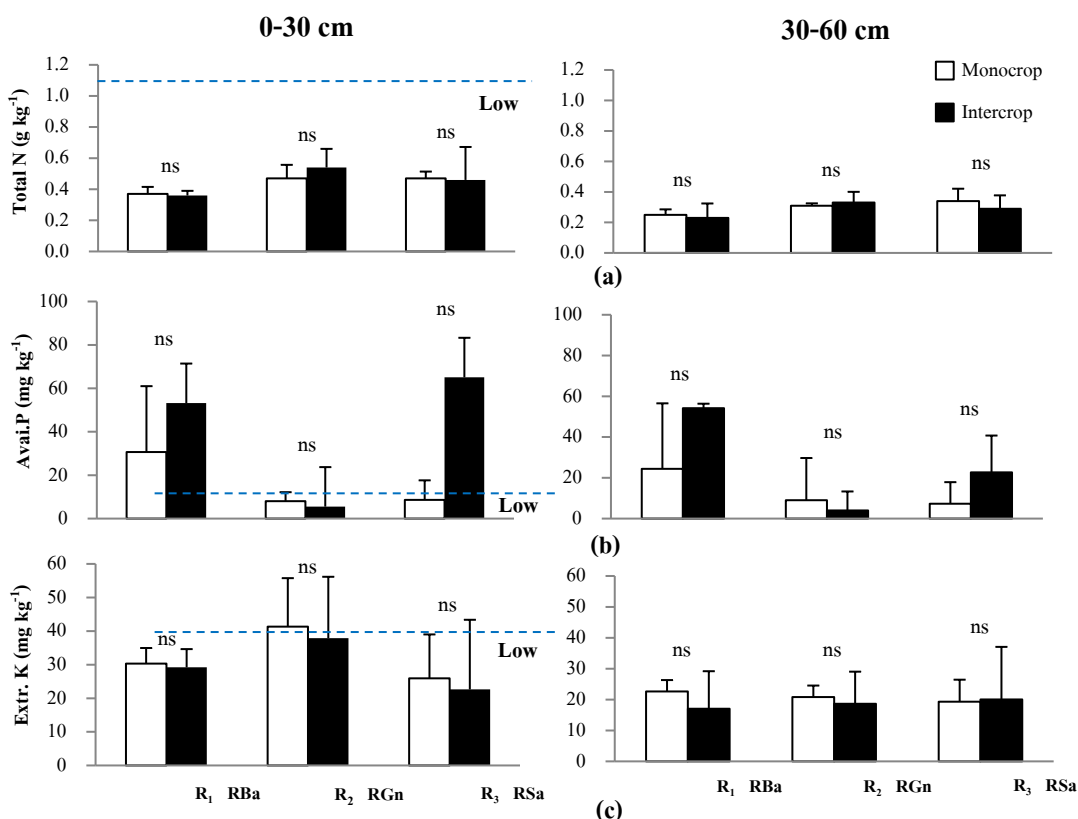
ตารางที่ 1 พีเอช ค่าสภาพน้ำไฟฟ้า และปริมาณอินทรีย์วัตถุในดินที่ความลึก 0-30 และ 30-60 เซนติเมตร

Treatment	pH _(1:5)		EC _{1:5} (dS m ⁻¹)		OM (g kg ⁻¹)	
	0-30 cm	30-60 cm	0-30 cm	30-60 cm	0-30 cm	30-60 cm
R ₁	5.0	5.1	0.02	0.01	8.6	5.9
RBa	5.0	5.4	0.02	0.01	8.6	5.2
T-test	ns	ns	ns	ns	ns	ns
R ₂	5.2	5.1	0.02	0.01	16.2	9.3
R _{Gn}	5.2	5.1	0.02	0.01	15.7	9.4
T-test	ns	ns	ns	ns	ns	ns
R ₃	5.0	4.9	0.01	0.01	13.1	7.4
RSa	5.1	4.5	0.03	0.01	11.2	5.8
T-test	ns	ns	ns	ns	ns	ns

หมายเหตุ : R₁₋₃ คือ แปลงปลูกยางพาราเชิงเดี่ยวที่อยู่ใกล้กับแปลงปลูกพืชร่วมยางพารา, RBa คือ แปลงปลูกไผ่ร่วมยางพารา, R_{Gn} คือ แปลงปลูกผักเหลียงร่วมยางพารา และ RSa คือแปลงปลูกสละร่วมยางพารา และ ns คือ ไม่มีความแตกต่างทางสถิติ

2. สถานะธาตุอาหารในดินปลูกยางพารา

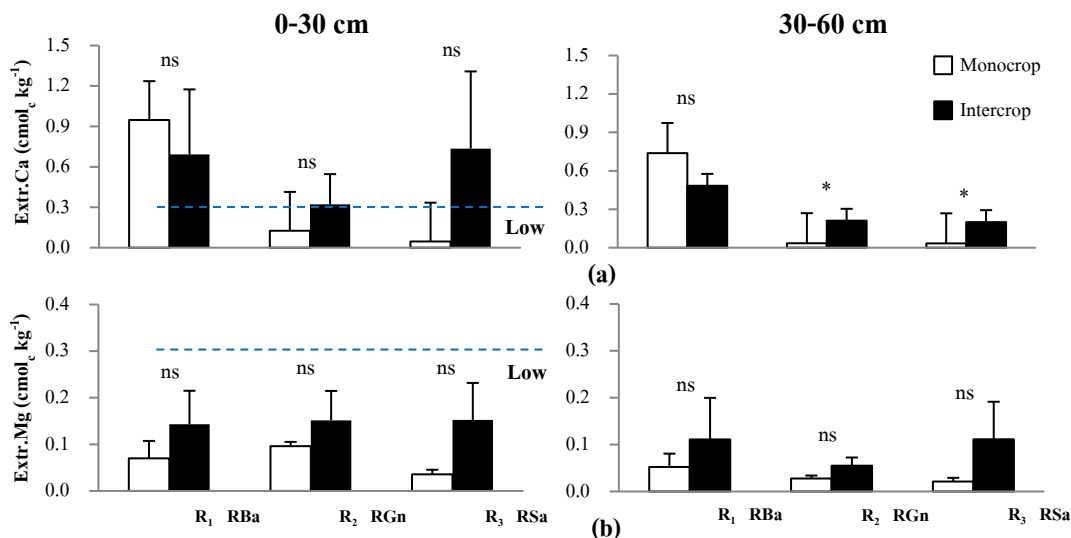
ปริมาณไนโตรเจนทั้งหมด ฟอสฟอรัสที่เป็นประโยชน์ และโพแทสเซียมที่สกัดได้ของแปลงปลูกยางพาราเชิงเดี่ยวและพืชร่วมยางพาราไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$) และมีแนวโน้มลดลงตามระดับความลึกดิน อย่างไรก็ตาม แปลงปลูกผักเหลียงร่วมยางพารามีแนวโน้มของปริมาณไนโตรเจนทั้งหมดและโพแทสเซียมที่สกัดได้สูงที่สุด ส่วนฟอสฟอรัสที่เป็นประโยชน์ของแปลงปลูกพืชร่วมยางพารามีแนวโน้มสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว โดยไนโตรเจนทั้งหมดจัดอยู่ในระดับต่ำทั้งหมด ฟอสฟอรัสที่เป็นประโยชน์ของแปลงปลูกยางพาราเชิงเดี่ยวส่วนใหญ่จัดอยู่ในระดับต่ำ ส่วนแปลงปลูกพืชร่วมยางพาราส่งส่วนใหญ่อยู่ในระดับสูง และ โพแทสเซียมที่สกัดได้ส่วนใหญ่อยู่ในระดับต่ำ เมื่อเปรียบเทียบกับเกณฑ์มาตรฐานของสถาบันวิจัยยางพารา (Kungpisdan, 2011)



รูปที่ 2 ไนโตรเจนทั้งหมด (a) ฟอสฟอรัสที่เป็นประโยชน์ (b) และโพแทสเซียมที่สกัดได้ (c) ของดินที่ความลึก 0-30 และ 30-60 เซนติเมตร เทียบกับเกณฑ์มาตรฐานของสถาบันวิจัยยางพารา (Kungpisdan, 2011)

หมายเหตุ : R₁₋₃ คือ แปลงปลูกยางพาราเชิงเดี่ยวที่อยู่ใกล้กับแปลงปลูกพืชร่วมยางพารา, R_{Ba} คือ แปลงปลูกไม้ร่วมยางพารา, R_{Gn} คือ แปลงปลูกผักเหลียงร่วมยางพารา และ R_{Sa} คือแปลงปลูกสละร่วมยางพารา และ ns คือ ไม่มีความแตกต่างทางสถิติ

ปริมาณแคลเซียมและแมกนีเซียมที่สกัดได้ของแปลงปลูกยางพาราเชิงเดี่ยวและพืชร่วมยางพาราไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$) และมีแนวโน้มลดลงตามระดับความลึกดิน อย่างไรก็ตาม แคลเซียมและแมกนีเซียมที่สกัดได้ของแปลงปลูกพืชร่วมยางพารามีแนวโน้มสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว โดยแคลเซียมที่สกัดได้ของแปลงปลูกยางพาราเชิงเดี่ยวส่วนใหญ่จัดอยู่ในระดับต่ำ ส่วนแปลงปลูกพืชร่วมยางพาราส่งส่วนใหญ่จัดอยู่ในระดับปานกลาง และแมกนีเซียมที่สกัดได้อยู่ในระดับต่ำทั้งหมด เมื่อเปรียบเทียบกับเกณฑ์มาตรฐานของสถาบันวิจัยยางพารา (Kungpisdan, 2011)

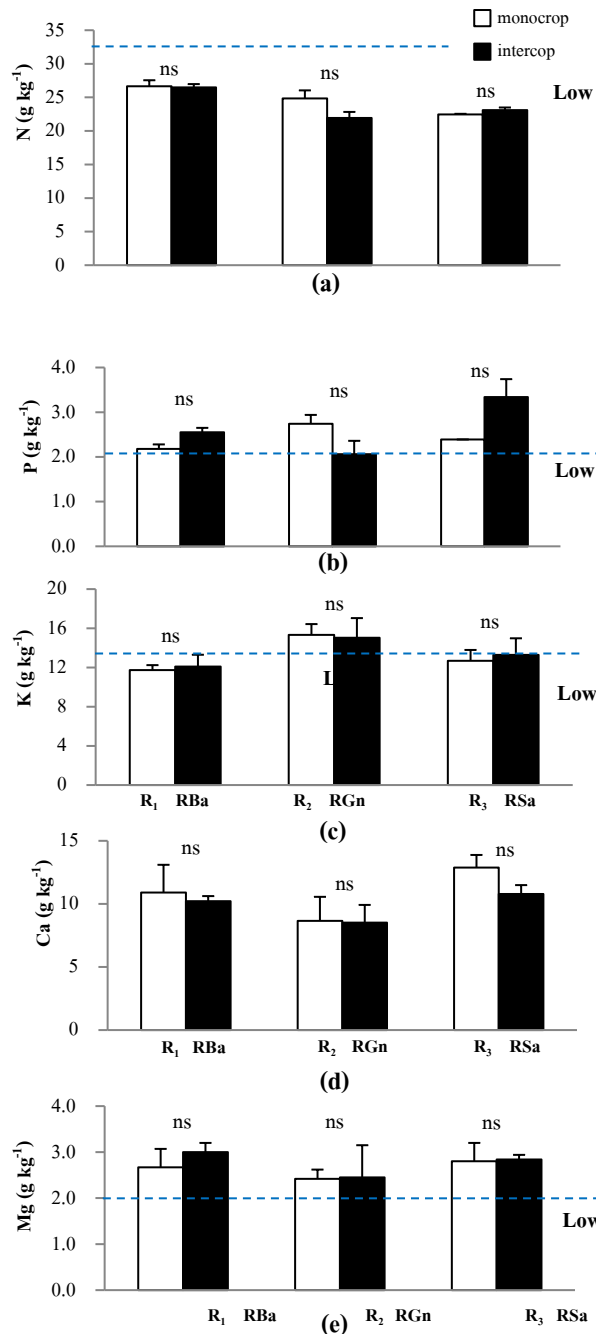


รูปที่ 3 แคลเซียมที่สกัดได้ (a) และแมกนีเซียมที่สกัดได้ (b) ของดินที่ความลึก 0-30 และ 30-60 เซนติเมตร เทียบกับเกณฑ์มาตรฐานของสถาบันวิจัยยางพารา (Kungpisdan, 2011)

หมายเหตุ : R₁₋₃ คือ แปลงปลูกยางพาราเชิงเดี่ยวที่อยู่ใกล้กับแปลงปลูกพืชร่วมยางพารา, RBa คือ แปลงปลูกไผ่ร่วมยางพารา, RGn คือ แปลงปลูกผักเหลียงร่วมยางพารา และ RSa คือแปลงปลูกสละร่วมยางพารา
* แตกต่างทางสถิติอย่างมีนัยสำคัญที่ระดับความเชื่อมั่น 95 เปอร์เซนต์ และ ns คือ ไม่มีความแตกต่างทางสถิติ

3. สถานะธาตุอาหารไนโตรเจนในใบยางพารา

ปริมาณไนโตรเจน ฟอสฟอรัส โพแทสเซียม แคลเซียม และแมกนีเซียมในใบยางพาราของแปลงปลูกยางพาราเชิงเดี่ยวและพืชร่วมยางพาราไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$) โดยไนโตรเจนจัดอยู่ในระดับต่ำทั้งหมด ฟอสฟอรัสจัดอยู่ในระดับปานกลางถึงสูง โพแทสเซียมส่วนใหญ่จัดอยู่ในระดับต่ำ และแมกนีเซียมจัดอยู่ในระดับปานกลางถึงสูงเมื่อเปรียบเทียบกับเกณฑ์มาตรฐานของสถาบันวิจัยยางพารา (Kungpisdan, 2011) (ภาพที่ 4) โดยแปลงปลูกพืชร่วมยางพารามีแนวโน้มของปริมาณฟอสฟอรัสและแมกนีเซียมในใบยางพาราสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว ส่วนปริมาณไนโตรเจน โพแทสเซียม และแคลเซียมของแปลงปลูกพืชร่วมยางพาราและแปลงปลูกยางพาราเชิงเดี่ยวไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$) (ตารางที่ 2)



รูปที่ 4 ไนโตรเจน (a) ฟอสฟอรัส (b) โพแทสเซียม (c) แคลเซียม (d) และแมกนีเซียม (e) ในใบยางพาราภายใต้การปลูกพืชร่วมยางพาราที่แตกต่างกัน เทียบกับเกณฑ์มาตรฐานของสถาบันวิจัยยางพารา (Kungpisdan, 2011)

หมายเหตุ : R₁₋₃ คือ แปลงปลูกยางพาราเชิงเดี่ยวที่อยู่ใกล้กับแปลงปลูกพืชร่วมยางพารา, RBa คือ แปลงปลูกไฟ้ร่วมยางพารา, RGn คือ แปลงปลูกผักเหลียงร่วมยางพารา และ RSa คือแปลงปลูกสละร่วมยางพารา และ ns คือ ไม่มีความแตกต่างทางสถิติ

ตารางที่ 2 ปริมาณธาตุอาหารในใบยางพาราภายใต้การปลูกพืชร่วมยางพาราที่แตกต่างกัน (ค่าเฉลี่ย $\pm$ SD)

Treatment	Concentration (g kg ⁻¹)				
	N	P	K	Ca	Mg
R ₁	26.7 $\pm$ 0.9 a	2.2 $\pm$ 0.1 b	11.7 $\pm$ 0.5	10.9 $\pm$ 2.2	2.7 $\pm$ 0.4
RBa	26.5 $\pm$ 1.2 a	2.6 $\pm$ 0.19 ab	12.1 $\pm$ 1.1	10.2 $\pm$ 1.9	3.0 $\pm$ 0.2
R ₂	24.8 $\pm$ 0.1 ab	2.7 $\pm$ 0.0 ab	15.3 $\pm$ 1.1	8.7 $\pm$ 1.0	2.4 $\pm$ 0.4
RGn	21.9 $\pm$ 0.5 c	2.1 $\pm$ 0.1 b	15.0 $\pm$ 1.2	8.5 $\pm$ 0.4	2.5 $\pm$ 0.2
R ₃	22.4 $\pm$ 0.9 bc	2.4 $\pm$ 0.3 b	12.7 $\pm$ 2.0	12.9 $\pm$ 1.4	2.8 $\pm$ 0.7
RSa	23.1 $\pm$ 0.4 bc	3.3 $\pm$ 0.4 a	13.3 $\pm$ 1.7	10.8 $\pm$ 0.7	2.8 $\pm$ 0.1
F-test	**	*	ns	ns	ns
CV (%)	7.8	17.5	19.2	26.1	21.9

หมายเหตุ : R₁₋₃ คือ แปลงปลูกยางพาราเชิงเดี่ยวที่อยู่ใกล้กับแปลงปลูกพืชร่วมยางพารา, RBa คือ แปลงปลูกไผ่ร่วมยางพารา, RGn คือ แปลงปลูกผักเหลียงร่วมยางพารา และ R₃ คือแปลงปลูกสละร่วมยางพารา

** แตกต่างทางสถิติอย่างมีนัยสำคัญที่ระดับความเชื่อมั่นที่ 99 เปอร์เซนต์, * แตกต่างทางสถิติอย่างมีนัยสำคัญที่ระดับความเชื่อมั่น 95 เปอร์เซนต์ และ ns คือ ไม่มีความแตกต่างทางสถิติ ตัวอักษรที่แตกต่างในแนวคอลัมน์แสดงถึงความแตกต่างเมื่อทดสอบด้วยวิธี DMRT

4. ปริมาณธาตุอาหารในน้ำยางพารา

ปริมาณไนโตรเจน ฟอสฟอรัส โพแทสเซียม แคลเซียม และแมกนีเซียมในน้ำยางพาราของแปลงปลูกยางพาราเชิงเดี่ยวและแปลงปลูกพืชร่วมยางพาราไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$) อย่างไรก็ตาม แมกนีเซียมของแปลงปลูกพืชร่วมยางพารามีแนวโน้มสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว (ตารางที่ 3)

ตารางที่ 3 ปริมาณธาตุอาหารในน้ำยางพาราภายใต้การปลูกพืชร่วมยางพาราที่แตกต่างกัน (ค่าเฉลี่ย $\pm$ SD)

Treatment	Concentrations (g kg ⁻¹)				
	N	P	K	Ca	Mg
R ₁	1.8 $\pm$ 0.2 a	0.9 $\pm$ 0.1	2.0 $\pm$ 0.4	0.06 $\pm$ 0.0	0.4 $\pm$ 0.1
RBa	2.1 $\pm$ 0.5 a	0.7 $\pm$ 0.1	1.9 $\pm$ 0.0	0.05 $\pm$ 0.0	0.5 $\pm$ 0.1
R ₂	1.1 $\pm$ 0.1 b	1.1 $\pm$ 0.1	2.3 $\pm$ 0.3	0.09 $\pm$ 0.0	0.5 $\pm$ 0.0
RGn	1.2 $\pm$ 0.1 b	1.0 $\pm$ 0.2	1.7 $\pm$ 0.2	0.07 $\pm$ 0.0	0.5 $\pm$ 0.2
R ₃	1.9 $\pm$ 0.7 a	0.6 $\pm$ 0.1	1.3 $\pm$ 0.1	0.05 $\pm$ 0.0	0.2 $\pm$ 0.1
RSa	1.9 $\pm$ 0.1 a	0.7 $\pm$ 0.1	1.5 $\pm$ 0.5	0.06 $\pm$ 0.0	0.3 $\pm$ 0.2
F-test	*	ns	ns	ns	ns
CV (%)	26.1	29.3	30.1	45.2	58.6

หมายเหตุ : R_{1-3} คือ แปลงปลูกยางพาราเชิงเดี่ยวที่อยู่ใกล้กับแปลงปลูกพืชร่วมยางพารา, RBa คือ แปลงปลูกไผ่ร่วมยางพารา, R_{Gn} คือ แปลงปลูกผักเหลียงร่วมยางพารา และ RSa คือแปลงปลูกสละร่วมยางพารา

* แตกต่างทางสถิติอย่างมีนัยสำคัญที่ระดับความเชื่อมั่น 95 เปอร์เซนต์ และ ns คือ ไม่มีความแตกต่างทางสถิติตัวอักษรที่แตกต่างในแนวคอลัมน์แสดงถึงความแตกต่างเมื่อทดสอบด้วยวิธี DMRT

5. ความสัมพันธ์ระหว่างธาตุอาหารในดิน ใบยางพารา และน้ำยางพารา

เมื่อนำธาตุอาหารในดิน ใบยางพารา และน้ำยางพารามาหาสัมประสิทธิ์สหสัมพันธ์ พบว่า ฟอสฟอรัสและโพแทสเซียมในดินมีสหสัมพันธ์กับฟอสฟอรัสและโพแทสเซียมในใบยางพารา ($r = 0.74^{**}$ และ 0.66^{**} ตามลำดับ) แต่ปริมาณธาตุอาหารอื่น ๆ ทั้งในดิน ใบยางพารา และในน้ำยางพาราไม่มีความสัมพันธ์กัน (ตารางที่ 4)

ตารางที่ 4 ความสัมพันธ์ระหว่างธาตุอาหารในดิน ใบยางพารา และน้ำยางพาราภายใต้การปลูกพืชร่วมยางพาราที่แตกต่างกัน

Nutrient	Composition	Composition		
		Soil	Leaves	Latex
N	Soil	1		
	Leaves	-0.39	1	
	Latex	0.23	-0.12	1
P	Soil	1		
	Leaves	0.74**	1	
	Latex	-0.26	-0.29	1
K	Soil	1		
	Leaves	0.66**	1	
	Latex	0.12	-0.01	1
Ca	Soil	1		
	Leaves	0.43	1	
	Latex	-0.21	-0.30	1
Mg	Soil	1		
	Leaves	0.26	1	
	Latex	0.25	0.15	1

หมายเหตุ : ** มีความสัมพันธ์กันอย่างมีนัยสำคัญที่ระดับความเชื่อมั่น 99 เปอร์เซนต์

วิจารณ์ผลการทดลอง

1. พีเอช ค่าสภาพนาไฟฟ้า และอินทรีย์วัตถุในดิน

จากผลการศึกษา พบว่า พีเอชของดินปลูกยางพาราเชิงเดี่ยวและปลูกพืชร่วมยางพาราทั้งสองระดับความลึกดิน เฉลี่ยเท่ากับ 4.5-5.4 ซึ่งเหมาะสมสำหรับปลูกยางพารา โดยมีรายงานว่าพีเอชที่เหมาะสมสำหรับปลูกยางพาราอยู่ในช่วง 4.5-5.5 (Kungpisadan, 2011) อย่างไรก็ตาม พีเอชของแปลงที่ศึกษามีค่าต่ำและจัดเป็นกรดปานกลางถึงกรดจัด เนื่องจากประเทศไทยตั้งอยู่ในเขตร้อนชื้นที่มีอุณหภูมิและปริมาณน้ำฝนสูงทำให้ดินมีพัฒนาการสูง (Popan, 2000) ส่งผลให้แคตไอออนชนิดเบสถูกชะละลายง่ายจึงเหลือในดินต่ำ แต่มีแคตไอออนชนิดกรด คือ ไฮโดรเจนและอลูมิเนียมไอออนสูง จึงส่งผลให้ดินมีพีเอชต่ำ (Brady & Weil, 2008) และค่าสภาพนาไฟฟ้า เฉลี่ยเท่ากับ 0.01-0.03 เดซิซิเมนส์/เมตร ซึ่งต่ำกว่า 2 เดซิซิเมนส์/เมตร แสดงให้เห็นว่าปริมาณเกลือที่ละลายในดินมีปริมาณต่ำและไม่มีผลกระทบต่อน้ำที่เป็นประโยชน์ต่อพืชจึงทำให้พืชสามารถดูดน้ำและธาตุอาหารไปใช้เพื่อการเจริญเติบโตได้ (Onthong & Poonpakdee, 2019)

ปริมาณอินทรีย์วัตถุของแปลงปลูกยางพาราเชิงเดี่ยวและปลูกพืชร่วมยางพาราลดลงตามความลึกดินที่เพิ่มขึ้น โดยปริมาณอินทรีย์วัตถุที่ระดับความลึกดิน 0-30 เซนติเมตร เฉลี่ยเท่ากับ 8.6-16.2 กรัม/กิโลกรัม (ตารางที่ 1) ซึ่งจัดอยู่ในระดับต่ำ (< 10 กรัม/กิโลกรัม) ถึงระดับปานกลาง (10-25 กรัม/กิโลกรัม) เนื่องจากภาคใต้มีภูมิอากาศแบบร้อนชื้นส่งผลให้การย่อยสลายเศษซากพืชเป็นไปอย่างรวดเร็ว จึงทำให้ปริมาณอินทรีย์วัตถุต่ำ (Carter, 2002) และอาจเนื่องมาจากอินทรีย์วัตถุมักจะสะสมที่บริเวณผิวดิน แต่จากแปลงที่ศึกษาได้เก็บดินที่ระดับความลึก 0-30 เซนติเมตร และผสมคลุกเคล้ากันจึงส่งผลให้อินทรีย์วัตถุจากผิวดินกระจายตัวทำให้มีปริมาณต่ำ อย่างไรก็ตาม การใส่วัสดุอินทรีย์หรือปุ๋ยอินทรีย์เพื่อเพิ่มอินทรีย์วัตถุและธาตุอาหารให้เพียงพอต่อความต้องการของยางพารานั้นต้องใช้เวลานาน และใช้ในปริมาณที่มาก โดยมีการศึกษาการใช้ปุ๋ยอินทรีย์อัตรา 0, 5, 10 และ 15 เมกกะกรัม/เฮกแตร์ ต่อปริมาณอินทรีย์วัตถุ พบว่า การใช้ปุ๋ยอินทรีย์ที่อัตรา 15 เมกกะกรัม/เฮกแตร์ ถึงจะมีผลให้อินทรีย์วัตถุในดินเพิ่มขึ้น (Barzegar et al., 2002) สอดคล้องกับผลการศึกษาการใส่ฟางข้าวสาลีอัตรา 0, 2, 4, 8 และ 16 เมกกะกรัม/เฮกแตร์/ปี ต่อปริมาณคาร์บอนอินทรีย์ พบว่า การใส่ฟางข้าวสาลีที่อัตรา 16 เมกกะกรัม/เฮกแตร์/ปี ส่งผลให้ปริมาณคาร์บอนอินทรีย์ในดินเพิ่มสูงขึ้นอย่างชัดเจน (Duiker & Lal, 1999) นอกจากนี้ มีรายงานว่าเกษตรกรแปลงปลูกพืชร่วมยางพารามีการใส่เศษซากพืชร่วมเพื่อให้ย่อยสลายเป็นอินทรีย์วัตถุลงในแปลง พบว่า ปริมาณเศษซากใบพืชร่วมยางพาราเฉลี่ยทั้งปีของเศษซากใบไม้ เท่ากับ 483.26 กิโลกรัม/ไร่/ปี และเศษซากใบผักเหลียง เท่ากับ 379.29 กิโลกรัม/ไร่/ปี (Saeteaw, 2019) ซึ่งเป็นปริมาณที่ต่ำเมื่อเปรียบเทียบกับปริมาณเศษซากที่ร่วงหล่นในป่าฝนธรรมชาติ (1,408-1,920 กิโลกรัม/ไร่/ปี) (Paoli & Curran, 2007) ดังนั้น จึงแนะนำให้เกษตรกรแปลงปลูกพืชร่วมยางพาราใส่ปุ๋ยอินทรีย์ร่วมกับใส่เศษซากพืชร่วมให้แก่ยางพารา

2. สถานะธาตุอาหารในดินปลูกยางพารา

ดินปลูกยางพารามีระดับไนโตรเจนทั้งหมด โพแทสเซียมและแมกนีเซียมที่สกัดได้ทั้งในแปลงปลูกยางพาราเชิงเดี่ยวและพืชร่วมยางพาราอยู่ในระดับต่ำ ($N < 1.1$ กรัม/กิโลกรัม, $K < 40.0$ มิลลิกรัม/กิโลกรัม และ $Mg < 0.3$ เซนติโมล/กิโลกรัม) (Kungpisadan, 2011) (ภาพที่ 2) เนื่องจากภาคใต้อยู่ในเขตร้อนชื้นที่มีปริมาณน้ำฝนและอุณหภูมิสูงทำให้อินทรีย์วัตถุซึ่งเป็นแหล่งที่มาของไนโตรเจนในดินเกิดการย่อยสลายได้อย่างรวดเร็ว (Popan, 2000) และไนโตรเจนสูญเสียออกจากพื้นที่โดยติดไปกับผลผลิตน้ำยางพารา ($N = 4.6-7.4$ กรัม/กิโลกรัม) โดยมีรายงานว่าน้ำยางพารา 1 ตัน มีไนโตรเจนติดไปกับผลผลิต 20 กิโลกรัม (Kungpisadan, 2009) รวมทั้งเกิดการชะละลาย โพแทสเซียมและแมกนีเซียมสูง (Popan, 2000) และแปลงที่ศึกษาส่วนใหญ่มีการใส่ปุ๋ยเคมีสูตร 15-15-15 เฉลี่ย 50 กิโลกรัม/ไร่/ปี เฉลี่ยต้นละไม่ถึง 1 กิโลกรัม/ไร่/ปี ซึ่งต่ำกว่าที่สถาบันวิจัยยางพาราแนะนำ และยังเป็นสูตรปุ๋ยที่ต่ำกว่าสูตรที่แนะนำ คือ 29-5-18 อย่างไรก็ตาม แปลงปลูกผักเหลียงร่วมยางพารามีแนวโน้มของปริมาณไนโตรเจนทั้งหมดและโพแทสเซียมที่สกัดได้สูงที่สุด (ภาพที่ 2) เนื่องจากเศษซากใบผักเหลียงที่เกษตรกรตัดแต่งกิ่งและใส่คืนกลับในแปลง จัดอยู่ในชั้นคุณภาพของสารอินทรีย์ ระหว่าง 1 และ 2 (ไนโตรเจน > 2.5 เปอร์เซ็นต์, ดินอิน < 15 เปอร์เซ็นต์ และโพธิ์ฟีนอล < 4 เปอร์เซ็นต์ ของน้ำหนักแห้ง) โดยชั้นคุณภาพที่ 1 จะเป็นแหล่งไนโตรเจนให้กับดินโดยตรง และชั้นคุณภาพที่ 2 คือ ส่วนที่ปลดปล่อยธาตุอาหารออกมาเป็นประโยชน์สู่ดิน และมีการศึกษาปริมาณธาตุอาหารจากเศษซากพืชร่วมยางพารา ได้แก่ ตะเคียน ไม้ และผักเหลียง พบว่า เศษซากใบผักเหลียงมีปริมาณไนโตรเจนและโพแทสเซียมคืนกลับสู่ดินสูงที่สุดเมื่อเปรียบเทียบกับพืชชนิดอื่น (Saeteaw, 2019; Palm et al., 2001) ส่วนแมกนีเซียมที่สกัดได้ในดิน พบว่า แปลงปลูกพืชร่วมยางพาราทั้งสามชนิด (ไม้ ผักเหลียง และสละ) มีปริมาณใกล้เคียงกัน 0.14-0.15 เซนติโมล/กิโลกรัม แต่มีแนวโน้มสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว (ภาพที่ 3) เนื่องจากแมกนีเซียมเป็นองค์ประกอบของคลอโรฟิลล์ทำให้แมกนีเซียมมีการสะสมในใบ (Osotsapar, 2015) และเมื่อเศษซากใบพืชร่วมร่วงหล่นหรือเศษซากใบพืชร่วมที่เกษตรกรใส่กลับคืนในพื้นที่ย่อยสลายและมีการปลดปล่อยแมกนีเซียมจึงส่งผลให้ดินของแปลงปลูกพืชร่วมยางพารามีการสะสมแมกนีเซียม ซึ่งมีการศึกษาปริมาณแมกนีเซียมที่คืนกลับสู่ดินจากเศษซากพืชร่วมยางพารา ได้แก่ ไม้ ผักเหลียง และตะเคียน พบว่า ไม้และผักเหลียงมีปริมาณแมกนีเซียมที่คืนกลับสู่ดิน (0.82 และ 0.58 กิโลกรัม/ไร่/ปี ตามลำดับ) สูงที่สุดเมื่อเปรียบเทียบกับพืชชนิดอื่น (Saeteaw, 2019) อีกทั้งแมกนีเซียมเป็นธาตุอาหารรองที่พืชต้องการในปริมาณน้อยเมื่อเปรียบเทียบกับไนโตรเจนและโพแทสเซียม จึงส่งผลให้แปลงปลูกพืชร่วมยางพารามีการสะสมแมกนีเซียม

ฟอสฟอรัสที่เป็นประโยชน์ของแปลงปลูกยางพาราเชิงเดี่ยวส่วนใหญ่จัดอยู่ในระดับต่ำ (< 11 มิลลิกรัม/กิโลกรัม) แต่แปลงปลูกพืชร่วมยางพาราส่วนใหญ่อยู่ในระดับสูง (> 30 มิลลิกรัม/กิโลกรัม) (ภาพที่ 2) เนื่องจากการปลูกพืชร่วมยางพาราเพิ่มปริมาณเศษซากพืชลงสู่ดิน และเมื่อเศษซากพืชย่อยสลายจึงมีการปลดปล่อยฟอสฟอรัส ซึ่งมีการศึกษาปริมาณฟอสฟอรัสที่คืนกลับสู่ดินจากเศษซากใบพืชร่วมยางพารา ได้แก่ ไม้ ผักเหลียง และตะเคียน พบว่า ไม้และผักเหลียงมีปริมาณฟอสฟอรัสที่คืนกลับสู่ดิน (0.14 และ 0.60 กิโลกรัม/ไร่/ปี ตามลำดับ) สูงที่สุดเมื่อ

เปรียบเทียบกับพืชชนิดอื่น (Saeteaw, 2019) อีกทั้งฟอสฟอรัสเป็นธาตุที่สูญเสียได้น้อย และพืชต้องการฟอสฟอรัสต่ำกว่าไนโตรเจนและโพแทสเซียมมาก จึงส่งผลให้แปลงปลูกพืชร่วมยางพารามีการสะสมฟอสฟอรัสโดยเฉพาะแปลงปลูกไผ่และสละร่วมยางพารา อาจเนื่องจากเศษซากใบไม้จัดอยู่ในชั้นคุณภาพสารอินทรีย์ที่ 3 (Saeteaw, 2019; Palm et al., 2001) และเศษซากใบสละจากผลการทดลองที่มีปริมาณไนโตรเจน 1.4 เปอร์เซ็นต์ โดยน้ำหนักแห้ง (ไม่แสดงข้อมูล) และสละจัดเป็นพืชวงศ์เดียวกับปาล์มน้ำมันซึ่งในทางใบของพืชตระกูลปาล์มมีปริมาณลิกนิน 21 เปอร์เซ็นต์ โดยน้ำหนักแห้งพืช (Loow & Wu, 2018) จึงคาดว่าเศษซากใบสละน่าจะจัดอยู่ในชั้นคุณภาพสารอินทรีย์ที่ 4 โดยชั้นคุณภาพที่ 3 เป็นส่วนที่ค่อยๆ ปลดปล่อยธาตุอาหารอย่างช้าๆ และชั้นคุณภาพที่ 4 เป็นกลุ่มที่ไวคลุมดินเพื่อลดการสูญเสียธาตุอาหารและการกร่อนดิน (Palm et al., 2001) ดังนั้น การปลูกไผ่และสละร่วมยางพาราจึงสามารถลดการสูญเสียและเกิดการสะสมฟอสฟอรัสในดิน อย่างไรก็ตาม แปลงปลูกผักเหลียงร่วมยางพารามีฟอสฟอรัสที่เป็นประโยชน์ในดินต่ำที่สุด อาจเนื่องมาจากเศษซากใบผักเหลียงอยู่ในชั้นคุณภาพสารอินทรีย์ที่ 2 ซึ่งสามารถปลดปล่อยธาตุอาหารออกมาเป็นประโยชน์ (Saeteaw, 2019; Palm et al., 2001) ส่งผลให้ยางพาราและผักเหลียงสามารถดูดฟอสฟอรัสไปใช้ได้ทันที จึงส่งผลให้ดินมีการสะสมฟอสฟอรัสที่เป็นประโยชน์ได้ต่ำ และปริมาณแคลเซียมที่สกัดได้ในดินมีแนวโน้มเช่นเดียวกัน (รูปที่ 3)

3. สถานะธาตุอาหารในใบยางพาราพารา

ไนโตรเจนในใบยางพาราของแปลงปลูกยางพาราเชิงเดี่ยวและปลูกพืชร่วมยางพาราอยู่ในระดับต่ำ (<33 กรัม/กิโลกรัม) (ภาพที่ 4) สอดคล้องกับปริมาณไนโตรเจนทั้งหมดในดินที่พบว่าอยู่ในระดับต่ำ ซึ่งเมื่อดินมีไนโตรเจนต่ำส่งผลให้ต้นยางพาราได้รับไนโตรเจนไม่เพียงพอต่อการเจริญเติบโต มีรายงานว่า การใส่ปุ๋ยไนโตรเจนเพิ่มขึ้นในยางพาราก่อนเปิดกรีดทำให้มีการสะสมของไนโตรเจนในใบสูงขึ้น (Boonmanee et al., 2013) อย่างไรก็ตาม ฟอสฟอรัสในใบยางพาราของแปลงปลูกยางพาราเชิงเดี่ยวส่วนใหญ่อยู่ในระดับสูง (>2.5 กรัม/กิโลกรัม) (ภาพที่ 4) ในขณะที่ดินมีฟอสฟอรัสที่เป็นประโยชน์ต่ำ เนื่องจากยางพาราที่ปลูกในดินกรดเขตร้อนสามารถปรับตัวด้วยกลไกการหลังกรีดอินทรีย์ออกจากราก โดยเฉพาะกรดซิตริก ซึ่งจะละลายอินทรีย์ฟอสฟอรัส โดยในดินเขตร้อนจะมีอะลูมิเนียมฟอสเฟตและเหล็กฟอสเฟตอยู่มาก เมื่อรูปดังกล่าวถูกละลายด้วยกรดอินทรีย์ ทำให้ฟอสฟอรัสละลายออกมาอยู่ในรูปที่เป็นประโยชน์ (Onthong & Osaki, 2006) อย่างไรก็ตาม แปลงปลูกพืชร่วมยางพาราส่วนใหญ่มีฟอสฟอรัสอยู่ในระดับปานกลางถึงสูง โดยแปลงปลูกไผ่และสละร่วมยางพารามีปริมาณฟอสฟอรัสสูงที่สุด (2.55 และ 3.34 กรัม/กิโลกรัม ตามลำดับ) เนื่องจากดินของแปลงปลูกไผ่และสละร่วมยางพารามีปริมาณฟอสฟอรัสที่เป็นประโยชน์สูง (ภาพที่ 2) ซึ่งสอดคล้องกับผลการศึกษาความสัมพันธ์ระหว่างฟอสฟอรัสในใบยางพาราและในดินปลูกยางพารา พบว่า ฟอสฟอรัสในใบยางพารามีสหสัมพันธ์เชิงบวกกับในดิน (0.738**) (ตารางที่ 4)

โพแทสเซียมในใบยางพาราของแปลงปลูกยางพาราเชิงเดี่ยวและแปลงปลูกพืชร่วมยางพาราส่วนใหญ่อยู่ในระดับต่ำ (<13.6 กรัม/กิโลกรัม) (ภาพที่ 4) เนื่องจากดินของแปลงปลูกยางพาราเชิงเดี่ยวและปลูกพืชร่วมยางพารา

มีปริมาณโพแทสเซียมที่สกัดได้ต่ำ (ภาพที่ 1) เช่นเดียวกับยางพาราที่ปลูกในพื้นที่ลุ่มและที่ดอนในจังหวัดสงขลา ที่พบว่า ยางพารามีโพแทสเซียมในใบต่ำ (Onthong et al., 2016) และยังสอดคล้องกับผลการศึกษาความสัมพันธ์ระหว่างโพแทสเซียมในใบยางพาราและในดินปลูกยางพารา พบว่า โพแทสเซียมในใบยางพารามีสหสัมพันธ์เชิงบวกกับในดิน (0.66^{**}) (ตารางที่ 4) ส่วนแมกนีเซียมในใบยางพาราของแปลงปลูกพืชร่วมยางพาราส่วนใหญ่อยู่ในระดับสูง (>2.5 กรัม/กิโลกรัม) (ภาพที่ 4) และมีแนวโน้มสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว (ตารางที่ 2) สอดคล้องกับปริมาณแมกนีเซียมที่สกัดได้ในดินของแปลงปลูกพืชร่วมยางพาราที่สูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว (ภาพที่ 3) แสดงให้เห็นว่าการปลูกพืชร่วมยางพาราทำให้ปริมาณแมกนีเซียมในดินเพิ่มสูงขึ้น จึงส่งผลให้ยางพาราสามารถดูดแมกนีเซียมไปใช้ได้มากขึ้น อย่างไรก็ตาม แมกนีเซียมในใบยางพาราของแปลงปลูกพืชร่วมยางพาราอยู่ในระดับสูง แต่แมกนีเซียมที่สกัดได้ในดินอยู่ในระดับต่ำ ซึ่งสอดคล้องกับการศึกษาสถานะแมกนีเซียมในดินและในใบยางพาราที่ปลูกในที่ลุ่มและที่ดอน พบว่า สถานะของแมกนีเซียมในดินปลูกยางพาราอยู่ในระดับต่ำแต่สถานะในใบยางพาราจัดอยู่ในระดับสูง (Kongmak et al., 2017) และมีรายงานว่าความเข้มข้นของแมกนีเซียมในใบยางพาราสูงกว่าค่าที่เหมาะสม ($2.0-2.5$ กรัม/กิโลกรัม) (Kungpisdan, 2011) แต่ยางพาราก็ยังเจริญเติบโตได้ดีเมื่อมีการใส่แมกนีเซียม (Pongthai, 2017) แสดงให้เห็นว่า ยางพาราอาจต้องการแมกนีเซียมในปริมาณที่สูงกว่าค่าที่เหมาะสมตามที่มียางานไว้

4. ปริมาณธาตุอาหารในน้ำยางพารา

ปริมาณไนโตรเจน ฟอสฟอรัส โพแทสเซียม แคลเซียม และแมกนีเซียมในน้ำยางพาราของแปลงปลูกยางพาราเชิงเดี่ยวและแปลงปลูกพืชร่วมยางพารามีปริมาณใกล้เคียงกัน อย่างไรก็ตาม แมกนีเซียมของแปลงปลูกพืชร่วมยางพารามีแนวโน้มสูงกว่าแปลงปลูกยางพาราเชิงเดี่ยว (ตารางที่ 3) เช่นเดียวกับแมกนีเซียมในดิน (ภาพที่ 3) และในใบยางพารา (ตารางที่ 2) ในขณะที่ธาตุอาหารในน้ำยางพารา พบว่า ปริมาณไนโตรเจนมีแนวโน้มสูงที่สุดรองลงมา คือ โพแทสเซียม ฟอสฟอรัส แมกนีเซียม และแคลเซียมต่ำที่สุด (ตารางที่ 3) ซึ่งสอดคล้องกับผลการศึกษาธาตุอาหารในน้ำยางพาราของ Yingjajaval & Bangjan (2006) ที่พบไนโตรเจนในน้ำยางพารามีปริมาณสูงที่สุดรองลงมา คือ โพแทสเซียม ฟอสฟอรัส แมกนีเซียม และแคลเซียมต่ำที่สุด สำหรับไนโตรเจนที่พบในน้ำยางพาราสูงที่สุด เนื่องจากไนโตรเจนเป็นองค์ประกอบของโปรตีน ซึ่งในน้ำยางพารามีโปรตีนประมาณ 1 เปอร์เซ็นต์ อีกทั้งไนโตรเจนเป็นองค์ประกอบของคลอโรฟิลล์ที่มีผลต่อการสังเคราะห์แสงและปริมาณซูโครสที่ใบยางพารา ซึ่งซูโครสจะเป็นสารตั้งต้นในการสร้างน้ำยางพารา โดยถ้าปริมาณซูโครสในน้ำยางพาราสูง แสดงว่าต้นยางพารามีศักยภาพในการสร้างน้ำยางพาราสูง แสดงให้เห็นว่าปริมาณของไนโตรเจนในน้ำยางพาราสอดคล้องกับปริมาณไนโตรเจนในใบยางพารา ซึ่งจากผลการศึกษา พบว่า แปลงปลูกพืชร่วมยางพารามีไนโตรเจนในน้ำยางพาราสูงที่สุด ซึ่งสอดคล้องกับไนโตรเจนในใบยางพาราของแปลงปลูกพืชร่วมยางพาราที่มีไนโตรเจนสูงที่สุด (26.5 กรัม/กิโลกรัม) (ตารางที่ 2)

โพแทสเซียมในน้ำยางพาราที่พบปริมาณรองจากไนโตรเจน ทั้งนี้เนื่องจากโพแทสเซียมเป็นธาตุที่เคลื่อนย้ายง่ายในท่ออาหาร จึงมีหน้าที่สำคัญในการเคลื่อนย้ายน้ำตาลซูโครสที่ได้จากกระบวนการสังเคราะห์แสงที่ใบ (Osotsapar, 2015) เพื่อนำไปสร้างน้ำยางพารา ส่งผลให้น้ำยางพารามีปริมาณโพแทสเซียมสูง และพบฟอสฟอรัสรองลงมา เนื่องจากฟอสฟอรัสเป็นองค์ประกอบที่สำคัญของสารที่ให้พลังงาน (adenosine triphosphate : ATP) โดยกระบวนการสร้างน้ำยางพาราจากซูโครสต้องใช้พลังงาน ATP เพื่อนำซูโครสที่ได้จากการสังเคราะห์แสงที่ใบเข้าสู่เซลล์ที่น้ำยางพารา และในกระบวนการเปลี่ยนกรดเม-วาโลนิค (mevalonic acid) จนได้เป็นไอโซเพนทีนิลไพโรฟอสเฟต (isopentenyl pyrophosphate : IPP) ก็ต้องใช้ ATP จากนั้น IPP ก็จะเปลี่ยนแปลงต่อไปเป็นโมเลกุลของไอโซพรีนและเชื่อมต่อกันจนได้เป็นน้ำยางพารา (Chanthuma et al., 2003) และพบแมกนีเซียมต่ำ เนื่องจากแมกนีเซียมเป็นธาตุที่เคลื่อนย้ายง่ายในน้ำเลี้ยงไหลเอม และเป็นธาตุที่มีหน้าที่เป็นตัวกระตุ้นเอนไซม์ที่เกี่ยวข้องกับการสังเคราะห์ยางพารา เช่น เอทีพีเอส และทรานสเฟอเรส อีกทั้งแมกนีเซียมยังเกี่ยวข้องกับเสถียรภาพของลูทอยด์ โดยหากลูทอยด์มีเสถียรภาพต่ำจะทำให้ลูทอยด์แตกและมีการปลดปล่อยแมกนีเซียมออกมา มีผลทำให้น้ำยางพาราจับตัวเป็นก้อน เกิดการอุดตันท่อน้ำยางพาราทำให้น้ำยางพาราไหลช้าจึงส่งผลให้ผลผลิตต่ำ (Chanthuma et al., 2003) ดังนั้น น้ำยางพาราที่มีเสถียรภาพสูงจะมีแมกนีเซียมต่ำ ส่วนแคลเซียมที่พบในน้ำยางพาราต่ำที่สุด เนื่องจากแคลเซียมเป็นธาตุที่เคลื่อนย้ายในท่ออาหารได้ยาก จึงพบน้อยมากในท่ออาหารพืช ซึ่งส่วนใหญ่สะสมอยู่ที่ผนังเซลล์ โดยเป็นส่วนของโครงสร้างในพืช (Osotsapar, 2015)

สรุปผลการทดลอง

การปลูกพืชร่วมยางพาราส่งผลให้ดินมีปริมาณฟอสฟอรัสที่เป็นประโยชน์ แคลเซียมและแมกนีเซียมที่สกัดได้สูงเมื่อเปรียบเทียบกับปลูกยางพาราเชิงเดี่ยว ซึ่งสอดคล้องกับปริมาณฟอสฟอรัสและแมกนีเซียมในใบยางพารา และแมกนีเซียมในน้ำยางพารา ในขณะที่การปลูกพืชร่วมยางพาราไม่ส่งผลต่อพีเอช ค่าสภาพการนำไฟฟ้า และปริมาณอินทรีย์วัตถุ อย่างไรก็ตาม สถานะธาตุอาหารในดินและใบยางพาราของไนโตรเจน โพแทสเซียม และแมกนีเซียมอยู่ในระดับต่ำ แต่ฟอสฟอรัสที่เป็นประโยชน์และแคลเซียมที่สกัดได้อยู่ในระดับปานกลางถึงสูงสำหรับดินปลูกยางพารา จากผลการทดลองชี้ให้เห็นว่า การปลูกพืชร่วมยางพาราช่วยเพิ่มฟอสฟอรัสที่เป็นประโยชน์ แคลเซียมและแมกนีเซียมที่สกัดได้ในดิน และส่งผลให้ยางพาราสามารถดูดไปใช้เพื่อการเจริญเติบโตและสร้างผลผลิตได้เพิ่มขึ้น เนื่องจากไนโตรเจน โพแทสเซียม และแมกนีเซียมในดินและใบยางพาราของการปลูกพืชร่วมยางพารายังมีปริมาณต่ำ ดังนั้น ควรเพิ่มอัตราปุ๋ยให้แก่ยางพาราและใส่ปุ๋ยตามค่าวิเคราะห์ดินหรือใบยางพาราเพื่อให้เพียงพอต่อความต้องการของยางพาราและพืชร่วมยางพารา โดยเฉพาะธาตุไนโตรเจน โพแทสเซียม และแมกนีเซียม

กิตติกรรมประกาศ

งานวิจัยนี้ได้รับทุนสนับสนุนการทำวิจัยจากงบประมาณแผ่นดิน (แผนบูรณาการวิจัยและนวัตกรรม) ประจำปีงบประมาณ 2562 รหัสโครงการ NAT620160g

เอกสารอ้างอิง

- Barzegar, A. R., Yousefi, A., & Daryashenas, A. (2002). The effect of addition of different amounts and types of organic materials on soil physical properties and yield of wheat. *Plant and Soil*, 247, 295-301.
- Boonmanee, S., Onthong, J., & Khawmee, K. (2013). Comparison of Fertilizer application based on soil testing and 20-8-20 mixed fertilizer in immature rubber trees. *King Mongkut's Agricultural Journal*, 31(1), 53-62. (in Thai)
- Brady, N. C., & Weil, R. R. (2008). *The Nature and Properties of Soils*. New Jersey: Pearson Prentice Hall.
- Carter, M. R. (2002). Soil quality for sustainable land management: organic matter and aggregation interactions that maintain soil functions. *Agronomy journal*, 94(1), 38-47.
- Chanthuma, P., Chanthuma, A., & Somnak, S. (2003). *Change of bio-chemical compositions in rubber latex on tapping system and yeid of rubber* (Rubber Research Report 2003). Bangkok: Rubber Authority of Thailand. (in Thai)
- Chen, C., Liu, W., Jiang, X., & Wu, J. (2017). Effects of rubber-based agroforestry systems on soil aggregation and associated soil organic carbon: Implications for land use. *Geoderma*, 299, 13-24.
- Duiker, S. W., & Lal, R. (1999). Crop residue and tillage effects on carbon sequestration in a Luvisol in central Ohio. *Soil and Tillage Research*, 52, 73-81.
- Kaewmapa, N., Choosai, C., Isarangkool, Na Ayutthaya S., Gonkhamdee, S., Sungthongwises, K., & Wongcharoen, A. (2014). Effect of different types of intercrops with rubber tree on some soil nutrients and soil fertility. *Khon Kaen Agriculture Journal*, 42, 443-449. (in Thai)
- Kongmak, P., Khawmee, K. & Onthong, J. (2017). Status and K/Mg ratio in soil and leaves of rubber trees grown in lowland and upland areas. *Songklanakarin Journal of Plant Science*, 4, 66-72. (in Thai)
- Kunpisadan, N. (2009). *Sustainable rubber management: Soil, Water and Plant nutrient*. Bangkok: Rubber Research Institute Department of Agriculture. (in Thai)
- Kunpisadan, N. (2011). *Recommend of the use rubber fertilizer*. Bangkok: Rubber Research Institute Department of Agriculture. (in Thai)

- Kungpisdan, N., Rattanachot, M., Permkrasin, P., Kiwrum, T., Chunamporn, L., & Thongphu, A. (2013). *Development of Technology on Nutrition Management of Rubber*. Bangkok: Rubber Research Institute Department of Agriculture. (in Thai)
- Loow, Y. L., & Wu, T. Y. (2018). Transformation of oil palm fronds into pentose sugars using copper (II) sulfate pentahydrate with the assistance of chemical additive. *Journal of environmental management*, 216, 192-203.
- Onthong, J., & Osaki, M. (2006). Adaptation of tropical plants to acid soils. *Tropics*, 15, 337-347.
- Onthong, J., & Poonpakdee, C. (2019). *Soil and Plant analysis*. Faculty of Natural Resources, Prince of Songkla University, SongKhla. (in Thai)
- Onthong, J., Khawmee, K., & Keawmano, C. (2016). Growth of immature rubber trees planted in abandoned paddy field and upland areas in relation to soil properties and leaf nutrients. *Songklanakarin Journal of Science and Technology*, 39, 675-683.
- Osotsapar, Y. (2015). *Plant nutrients*. Kasetsart University press, Bangkok. (in Thai)
- Palm, C. A., Gachengo, C. N., Delve, R. J., Cadisch, G., & Giller, K. E. (2001). Organic inputs for soil fertility management in tropical agroecosystems : application of an organic resource database. *Agriculture, Ecosystems and Environment*, 83, 27-42.
- Paoli, G. D., & Curran, L. M. (2007). Soil nutrients limit fine litter production and tree growth in mature lowland forest of Southwestern Borneo. *Ecosystems*, 10, 503-518.
- Pongthai, T. Onthong, J., & Khawmee, K. (2017). Effect of magnesium on nutrient concentration and growth of rubber tree sapling. *Journal of Agr. Research and Extension*, 34, 1-12. (in Thai)
- Popan, A. (2000). *Tropical soil*. King Mongkut's Institute of Technology Ladkrabang, Department of Soil Science, Bangkok. (in Thai)
- Rubber Intelligence Unit. (2020). *Production and Export Statistics - Natural Rubber*. <http://rubber.oie.go.th/Production.aspx> (in Thai)
- Rubber Research Institute of Thailand. (2007). *Academic information of rubber in 2007*. Rubber Research Institute Department of Agriculture, Bangkok. (in Thai)
- Saeteaw, W. (2019). *Organic Carbon Composition and Properties of Soils under Different Rubber-based Intercrops* [Unpublished master's thesis]. Prince of Songkla University. (in Thai)

Suchartkul, S. (2011). *Establishment of standard values for nutritional diagnosis in soil and leaves of immature rubber tree* [Unpublished master's thesis]. Walailak University. (in Thai)

Yingjajaval, S., & Bangjan, J. (2006). Major plant nutrients contents in para rubber (RRIM 600). *Agricultural Science Journal*, 37(4), 353-364. (in Thai)

Research Article

การพัฒนาผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือด

Development of cracker products from purple yam flour (*Dioscorea alata* L.)

ฤทัยศักดิ์ ชาญศรี^{*1}, รัสรินทร์ ฉัตรทองพิสุทธิ์¹, รัชวรรณ จันทเขต²

Ruethaipak Chansri^{*1}, Rassarin Chatthongpisut¹, Ratchawan Jantakhat²

¹ภาควิชาเกษตรและสิ่งแวดล้อม หลักสูตรวิทยาศาสตรบัณฑิตและเทคโนโลยีอาหาร มหาวิทยาลัยราชภัฏสุรินทร์ ประเทศไทย

²ภาควิชาเกษตรและสิ่งแวดล้อม หลักสูตรครุศาสตรบัณฑิต สาขาวิชาคหกรรมศาสตร์ มหาวิทยาลัยราชภัฏสุรินทร์ ประเทศไทย

¹Department of Agriculture and Environment, Food Science and Technology Program, Surindra Rajabhat University, Thailand

²Department of Agriculture and Environment, Bachelor of Education Program in Home Economics, Surindra Rajabhat University, Thailand

*E-mail: Ruethaipak_d@hotmail.com

Received: 24/05/2021; Revised: 16/04/2022; Accepted: 17/04/2022

บทคัดย่อ

การวิจัยนี้มีวัตถุประสงค์เพื่อพัฒนาผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือดที่ผ่านกรรมวิธีการทำแห้ง หลังจากการลวก นำมาวิเคราะห์ความหนืดของแป้ง และพัฒนาเป็นผลิตภัณฑ์ข้าวเกรียบ ทดสอบการยอมรับทางประสาทสัมผัส วิเคราะห์องค์ประกอบทางเคมี และพลังงานที่ได้รับ ผลการวิจัยพบว่า แป้งมันเลือด มีความหนืดสูงสุด 91.92 ± 2.24 RVU และความหนืดสุดท้าย 113.42 ± 2.48 RVU การพัฒนาเป็นผลิตภัณฑ์ข้าวเกรียบพบว่าสูตรที่ 3 ที่ใช้แป้งมันเลือดทดแทนแป้งมันสำปะหลังร้อยละ 40 ได้คะแนนการยอมรับทางประสาทสัมผัสในด้านลักษณะปรากฏ ความกรอบ สี และความชอบโดยรวมสูงที่สุด ผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือด 100 กรัม มีปริมาณความชื้น คาร์โบไฮเดรต ไขมัน โปรตีน และเส้นใย ร้อยละ 1.73 ± 0.04 , 65.37 ± 1.49 , 12.42 ± 0.62 , 17.87 ± 0.77 , 2.61 ± 0.32 และ 1.41 ± 0.47 ตามลำดับ มีพลังงานที่ได้รับทั้งหมด 466.39 ± 0.96 กิโลแคลอรี

คำสำคัญ: ข้าวเกรียบ แป้งมันเลือด แป้งมันสำปะหลัง

Abstract

The objective of this research was to study the preparation of cracker from purple yam flour that has been dried after blanching it before drying. The viscosity of starch and development into a cracker product, sensory acceptance, proximate composition and energy were determined. The results showed that purple yam

flour had the highest viscosity of 91.92 ± 2.24 RVU and final viscosity of 113.42 ± 2.48 RVU. In the development of cracker product, the third formula with 40% of the purple yam flour as a substitute for cassava starch received the highest sensory acceptance scores for appearance, crispness, color, and overall preference. The percentage of moisture content, carbohydrate, ash, total fats, protein and fiber of the cracker were shown at 1.73 ± 0.04 , 65.37 ± 1.49 , 12.42 ± 0.62 , 17.87 ± 0.77 , 2.61 ± 0.32 and $1.41 \pm 0.47\%$, respectively, while its total energy value derived was 466.39 ± 0.96 kcal.

Keywords: Crackers, purple yam flour, Cassava starch

บทนำ

มันพื้นบ้านจัดเป็นพืชอาหารที่มีให้รับประทานได้ตามช่วงฤดูกาลของทุกปี จัดเป็นพืชหัวที่ให้คุณค่าทางอาหารและมีประโยชน์ต่อร่างกาย ปัจจุบันมีงานวิจัยเกี่ยวกับพืชหัวค่อนข้างมากเนื่องจากเป็นพืชอาหารที่มีปริมาณผลผลิตสูง เป็นพืชหัวที่สามารถผลิตได้ปริมาณมากเป็นอันดับที่สามของโลกรองจากมันเทศ (Lebot, 2020) เป็นแหล่งคาร์โบไฮเดรตที่สำคัญในประเทศเขตร้อนและกึ่งเขตร้อน แต่ในประเทศไทยการบริโภคยังคงมีแค่จำกัด โดยเฉพาะมันพื้นบ้าน (yam) ที่มีอยู่หลายชนิด เช่น มันเลือด เป็นมันพื้นบ้านชนิดหนึ่งที่มีปริมาณผลผลิตต่อต้นสูง มันเลือดเป็นพืชมีหัวในวงศ์ถอย (Dioscoreaceae) มีชื่อวิทยาศาสตร์ว่า *Dioscorea alata* L. รู้จักในชื่อ purple yam ซึ่งเป็นคนละชนิดกับมันเทศสีม่วง (purple sweet potato) ที่อยู่ในวงศ์ผักบุ้ง (Ipomoea) (Cakrawati, 2021) ในมันเลือดมีสารอาหารที่มีคุณค่าทางโภชนาการครบถ้วน ทั้งแป้ง น้ำตาล โปรตีน และธาตุอาหารที่มีประโยชน์ เช่น แคลเซียม แมกนีเซียม โพแทสเซียม กรดโฟลิก และสารสี เช่น สารแคโรทีนอยด์และสารแอนโทไซยานิน (Promdang et al., 2018) ซึ่งเป็นสารต้านอนุมูลอิสระที่มีประโยชน์แก่ร่างกาย (Larief et al., 2018) นอกจากนี้มันเลือดยังเป็นแหล่งของวิตามินซีที่สำคัญ (Chen et al., 2003) และเป็นแหล่งของฮอร์โมน diosgenin ที่มีคุณสมบัติคล้ายกับฮอร์โมนเพศหญิง (Phytoestrogen) diosgenin สามารถถูกเปลี่ยนให้เป็นฮอร์โมน progesterone ได้อีกด้วย (Satour et al., 2007) บางสายพันธุ์มีค่าเส้นใยอาหารสูงสามารถนำมาทำอาหารสำหรับผู้ป่วยเบาหวานได้ด้วย (Okon & Famurewa, 2015) ช่วยลดปัญหาด้านสุขภาพของหัวใจและหลอดเลือด ป้องกันการเกิดโรคเบาหวานชนิดที่ 2 และโรคอ้วน (Cory et al., 2018) สารอาหารและสารประกอบทางชีวภาพต่างๆ เหล่านี้ทำให้เกิดการส่งเสริมให้ปลูกมันเลือดเพื่อเป็นการอนุรักษ์พันธุ์กรรมพืชสำหรับการใช้ประโยชน์ในอนาคตทางด้านอาหารและยา งานวิจัยส่วนใหญ่จึงมุ่งเน้นไปทางการวิเคราะห์องค์ประกอบทางเคมีในอาหาร คุณค่าทางโภชนาการ และคุณสมบัติการเป็นสารออกฤทธิ์ทางชีวภาพต่างๆ อย่างไรก็ตามการนำมันเลือดมาแปรรูปเป็นอาหารในประเทศไทยยังพบน้อยและมักนิยมรับประทานในวงจำกัด โดยส่วนใหญ่จะนำมานึ่งแล้วรับประทานพร้อมกับน้ำตาล หรือเชื่อมทานเป็นขนมหวานแบบต่างๆ ในต่างประเทศการพัฒนาแป้งมันเลือดเพื่อเป็นผลิตภัณฑ์อาหารส่วนใหญ่จะคำนึงถึงประโยชน์ที่ได้รับ

จากแป้งมันเลือดหรือการออกฤทธิ์ของสารที่สำคัญในแป้งมันเลือด เช่น การศึกษาผลของการผลิตขนมเค้กจากมันเลือดที่มีต่อการคงอยู่ของสารแอนโทไซยานินและฤทธิ์การต้านอนุมูลอิสระในก่อนและหลังการผลิต (Larief et al., 2018) ซึ่งพบว่าหลังจากผลิตเค้กปริมาณของแอนโทไซยานินและฤทธิ์ต้านอนุมูลอิสระนั้นลดลง ซึ่งการลดลงของฤทธิ์ต้านอนุมูลอิสระเป็นสัดส่วนโดยตรงกับการลดลงของระดับแอนโทไซยานิน หรือศึกษาการย่อยของขนมปังที่ทดแทนแป้งสาลีด้วยแป้งมันเลือดในหลอดทดลอง ซึ่งพบว่าขนมปังที่มีการทดแทนด้วยแป้งมันเลือดจะย่อยได้ยากกว่าเนื่องจากเม็ดแป้งของมันเลือดไม่ถูกทำลายด้วยการมีสารยับยั้งเอนไซม์ย่อยอาหารในแป้งมันเลือด (Liu et al., 2019) หรือ การศึกษาผลของการต้มต่อโปรไฟล์คาร์โบไฮเดรต การย่อยได้ของแป้ง และดัชนีระดับน้ำตาลในเลือดของผงมันเลือดของ Rosales & Barrion (2019) พบว่าแป้งมันเลือดอยู่ในหมวดอาหารดัชนีน้ำตาลปานกลาง การบริโภคจะเพิ่มระดับการตอบสนองของน้ำตาลในเลือดในปริมาณปานกลางและยังคงเป็นพืชหัวที่แนะนำสำหรับผู้ป่วยโรคเบาหวานได้ เป็นต้น

ผู้วิจัยเล็งเห็นถึงความสำคัญของการใช้ประโยชน์จากมันเลือด เพื่อเพิ่มมูลค่าของมันเลือดให้มีการใช้ประโยชน์มากขึ้นจึงต้องการนำมาทำผลิตภัณฑ์อาหารแปรรูป และพัฒนาเป็นผลิตภัณฑ์อาหารที่คงคุณค่าทางโภชนาการ สามารถผลิตได้ง่ายในครัวเรือน ช่วยให้เกิดอาชีพในท้องถิ่น และสามารถพัฒนาให้เป็นผลิตภัณฑ์ชุมชนต่อไปได้ ผู้วิจัยจึงได้นำมันเลือดที่มีในท้องถิ่นจังหวัดสุรินทร์มาทำการศึกษาในครั้งนี้ โดยมีวัตถุประสงค์เพื่อจะศึกษาวิธีการผลิตแป้ง ศึกษาพฤติกรรมความหนืดของแป้ง ศึกษาองค์ประกอบทางเคมีของแป้ง ศึกษาวิธีการผลิตข้าวเกรียบ ทดสอบการยอมรับทางประสาทสัมผัส วิเคราะห์องค์ประกอบทางเคมีของข้าวเกรียบมันเลือดและพลังงานที่ได้รับ ซึ่งในระยะแรกของงานวิจัยได้ทำการศึกษาการผลิตแป้งมันเลือดโดยเลือกวิธีที่เหมาะสมที่สุดจากปริมาณแอนโทไซยานินที่มีมากที่สุดซึ่งพบว่าวิธีการผลิตแป้งโดยการทำให้มันเลือดหลังการลวกให้ปริมาณแอนโทไซยานินมากกว่ามีค่าเท่ากับ 183.39 ± 1.88 mg/l ในขณะที่เดียวกันวิธีการผลิตแป้งโดยไม่ผ่านการลวกให้ปริมาณแอนโทไซยานินต่ำกว่า มีค่าเท่ากับ 152.56 ± 1.92 mg/l (Chansri et al., 2019) ดังนั้นผู้วิจัยจึงเลือกผลิตแป้งมันเลือดด้วยวิธีการดังกล่าว นำแป้งที่ได้มาวิเคราะห์องค์ประกอบทางเคมี และศึกษาพฤติกรรมความหนืดของแป้ง พบว่าแป้งมันเลือดมีลักษณะการฟูรูปแบบซี (C-Type) ตามที่ Schoch และ Maywald (1968) ได้จำแนกเอาไว้ว่าเป็นแป้งที่มีการพองตัวจำกัด (Restricted-swelling starch) และแป้งที่พองตัวทนทานต่อแรงเฉือนจากการกวน จัดเป็นแป้งที่มีคุณลักษณะที่เหมาะสมต่อการนำมาทำเป็นผลิตภัณฑ์ขนิดพองกรอบ (Sesang, 2020) ผู้วิจัยจึงเลือกแป้งมันเลือดมาผลิตเป็นข้าวเกรียบเนื่องจากทำได้ง่าย และเป็นอาหารว่างที่คนนิยมรับประทาน

วัสดุ อุปกรณ์ และวิธีการทดลอง

1. การเตรียมแป้งมันเลือด

ผลิตแป้งมันเลือดด้วยวิธีการบดแห้งแบบลวกมันก่อนตาก (Suriya et al., 2016) นำมันเลือดที่ขูดได้ในพื้นที่จังหวัดสุรินทร์มาปอกเปลือกออกให้หมด และล้างน้ำกำจัดเศษดินต่างๆ ออก และนำมันเลือดที่ล้างแล้วมาผ่านให้

เป็นแผ่นหนาประมาณ 2 มิลลิเมตร นำมันที่ผ่านแล้วเทลงในน้ำต้มอุณหภูมิ 70 °C นาน 5 นาที ตักขึ้นสะเด็ดน้ำ และล้างด้วยน้ำเปล่าจนหมดเมือก ตากมันจนแห้งสนิท วัดความชื้นได้ 12.95 ± 0.26 % และบดด้วยเครื่องปั่นจนละเอียด นำมาร้อนผ่านรูตะแกรงขนาด 120 เมช นำมาอบที่อุณหภูมิ 50 °C นาน 2 ชั่วโมง ทิ้งให้เย็น (ความชื้น 7.05 ± 0.06 %) และเก็บรักษาในถุงโพลีโพรพิลีนที่ปิดสนิท

2. การวิเคราะห์ความหนืดของแป้งมันเลือด

โดยเครื่อง Rapid Visco Analyzer (RVA; RVA-4 SA, Newport Scientific PYT Ltd., NSW, ประเทศออสเตรเลีย) เตรียมตัวอย่างแป้งให้มีความเข้มข้นร้อยละ 8 (w/w) ให้มีน้ำหนักรวม 25 กรัม โดยชั่งน้ำกลั่นใส่กระบอกอูลูมิเนียม จากนั้นจึงชั่งแป้งใส่ในกระบอกอูลูมิเนียม ใช้ใบพัดกวนผสมให้เข้ากัน ตั้งเครื่อง RVA โดยเริ่มต้นที่อุณหภูมิ 50 °C เพิ่มอุณหภูมิในอัตรา 13 °C/นาที จนถึงอุณหภูมิสูงสุด 95 °C จากนั้นกวนต่อแล้วลดอุณหภูมิลงเป็น 50 °C ด้วยอัตรา 7 °C/นาที และกวนต่อที่ 50 °C จนถึงสิ้นสุดเวลาที่ตั้งไว้

3. ศึกษาการทำผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือด

3.1 สูตรการผลิตผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือด คัดแปลงจาก Phongnoree & Thammasiri (2008) โดยใช้แป้งมันเลือดทดแทนแป้งมันสำปะหลังในท้องตลาด ในการทำผลิตภัณฑ์ข้าวเกรียบ สูตรที่ 1, 2, 3, 4, 5 และ 6 โดยทดแทนด้วยแป้งมันเลือดร้อยละ 0, 20, 40, 60, 80 และ 100 ตามลำดับ แสดงดังตารางที่ 1

สูตรมาตรฐานของข้าวเกรียบประกอบด้วย แป้งมันสำปะหลัง เกลือ กระเทียม และพริกไทย ร้อยละ 85, 3, 9 และ 3 ตามลำดับ

ตารางที่ 1 ข้าวเกรียบมันเลือดที่ทดแทนด้วยแป้งมันสำปะหลังแต่ละสูตร

วัตถุดิบ	ส่วนผสมหลัก (%)	อัตราส่วนแป้งมันเลือดต่อแป้งมันสำปะหลัง (%)					
		สูตรที่ 1	สูตรที่ 2	สูตรที่ 3	สูตรที่ 4	สูตรที่ 5	สูตรที่ 6
		0 : 100	20 : 80	40 : 60	60 : 40	80 : 20	100 : 0
แป้งมันเลือด	-	0	17	34	51	68	85
แป้งมันสำปะหลัง	85	85	68	51	34	17	0
เกลือ	3	3	3	3	3	3	3
กระเทียม	9	9	9	9	9	9	9
พริกไทย	3	3	3	3	3	3	3
น้ำร้อน (30%ของแต่ละสูตร)	-	30	30	30	30	30	30

3.2 วิธีการผลิตข้าวเกรียบจากแป้งมันเลือด ซึ่งส่วนผสมทุกอย่างใส่ลงในชามผสม ใช้น้ำร้อน (อุณหภูมิ 95 °C) ร้อยละ 30 ของสูตร ค่อยๆ เทลงในส่วนผสมแล้วคล่อมส่วนผสมทั้งหมดเข้าด้วยกัน นวดจนส่วนผสมเป็นเนื้อเดียวกัน ปั้นเป็นแท่งกลมเส้นผ่านศูนย์กลาง 1.5 นิ้ว นำไปนึ่งจนสุกและทิ้งไว้ให้เย็น เก็บเข้าสู่เย็น 20 ชั่วโมง นำออกมาหั่นให้มีขนาดหนา 2 มิลลิเมตร และนำไปตากแดด 5 ชั่วโมงจนแห้งสนิท นำไปทอดด้วยน้ำมันปาล์ม ที่อุณหภูมิ 175 °C ใช้เวลา 3-5 วินาที เมื่อข้าวเกรียบพอง ตักขึ้นสะเด็ดน้ำมัน รอให้เย็น และเก็บรักษาในถุงพลาสติกชนิดโพลิโพรพิลีน (polypropylene) ที่ปิดสนิท

4. การทดสอบคุณภาพทางประสาทสัมผัส

โดยใช้แบบทดสอบ 9 point Hedonic scale scoring test (1=ไม่ชอบมากที่สุด ถึง 9=ชอบมากที่สุด) กับผู้ทดสอบที่ไม่ผ่านการฝึกฝนจำนวน 15 คน โดยให้ผู้ทดสอบชิมประเมินและทำการให้คะแนน คุณลักษณะต่างๆ คือ ลักษณะปรากฏ สี กลิ่นรส ความกรอบ และความชอบโดยรวม และวิเคราะห์ทางสถิติด้วยโปรแกรมสำเร็จรูป SPSS ทดสอบความแตกต่างของคะแนนที่ระดับความเชื่อมั่นร้อยละ 95 วิเคราะห์ความแปรปรวน (Analysis of variance) และทดสอบความแตกต่างด้วยค่า F-test เปรียบเทียบความแตกต่างของค่าเฉลี่ยด้วย Duncan's New Multiple Range Test (DMRT)

5. การวิเคราะห์องค์ประกอบทางเคมีและพลังงานที่ได้รับ

โดยนำผลิตภัณฑ์ข้าวเกรียบที่ผู้ทดสอบยอมรับมากที่สุดมาวิเคราะห์เพื่อหาปริมาณความชื้น โดยวิธี Direct Heating Method อบด้วยตู้อบร้อน (Hot air oven) รุ่น Fd WTB binder วิเคราะห์ปริมาณเถ้าโดยการเผาด้วยเตาเผา อุณหภูมิสูง (furnace) ยี่ห้อ Nabertherm รุ่น L/LT วิเคราะห์ปริมาณไขมันด้วยเครื่องวิเคราะห์ไขมัน ยี่ห้อ Solvent Extractions รุ่น Model SER 148 วิเคราะห์ปริมาณโปรตีนด้วยเครื่องมือย่อยโปรตีน ยี่ห้อ Gerhardt รุ่น Kjeldatherm หาปริมาณเส้นใยด้วยเครื่องวิเคราะห์เส้นใย ยี่ห้อ Gerhardt รุ่น FT12 ตามวิธีของ AOAC (1995) หาปริมาณคาร์โบไฮเดรตโดยการคำนวณจากผลต่างของสารอาหารทั้งหมด และหาพลังงานที่ได้รับทั้งหมดจากเครื่องวิเคราะห์พลังงานในอาหารยี่ห้อ IKA รุ่น C600

ผลการทดลองและวิจารณ์ผล

1. การผลิตแป้งจากมันเลือด

แป้งมันเลือดที่ผลิตได้มีสีขาวกว่าแป้งมันเลือดที่ไม่ผ่านการลวกก่อนตาก (Chansri et al., 2019) เนื่องจากน้ำร้อนที่ใช้ลวกมีผลต่อการยับยั้งปฏิกิริยาการเกิดสีน้ำตาลจากเอนไซม์ (enzymatic browning reaction) ซึ่งมีหลายงานวิจัยที่ใช้วิธีการนี้เพื่อลดการเกิดสีน้ำตาลในกระบวนการผลิตแป้งจากพืชหัวชนิดต่างๆ นอกจากนี้ ยังมีรายงานจาก Metsdagh et al. (2008) ทำการศึกษาเรื่องการลวก พบว่าการลวกสามารถลดน้ำตาลรีดิวซ์ในแป้งลงได้ และทำให้มันฝรั่งทอดมีสีที่สว่างขึ้น และการลวกมันเทศสามารถป้องกันการเกิดปฏิกิริยาสีน้ำตาลที่เกิดขึ้นจาก

กระบวนการผลิตแป้งได้ เช่นเดียวกันกับการงานวิจัยของ Chen et al. (2017) ที่ได้ศึกษาการป้องกันการเกิดปฏิกิริยาการเกิดสีน้ำตาลในมันฝรั่งผ่านการลวกสามารถป้องกันการเกิดสีน้ำตาลจากเอนไซม์ได้ แต่อย่างไรก็ตามมีการศึกษาในปีต่อมาโดย Capotorto et al. (2018) พบว่า การลวกโดยการเติมกรดซิตริกลงไปใต้น้ำที่ใช้ลวกจะลดการทำงานของเอนไซม์ฟีนอลเลสที่เป็นสาเหตุของการเกิดปฏิกิริยาสีน้ำตาลได้ดีกว่าใช้น้ำร้อนเพียงอย่างเดียว แต่การเติมกรดซิตริกอาจมีผลต่อคุณภาพทางประสาทสัมผัสโดยเฉพาะรสชาติของผลิตภัณฑ์ที่อาจจะเกิดขึ้นได้ ทำให้ผู้วิจัยไม่เลือกวิธีการดังกล่าวมาใช้ในการผลิตแป้งมันเลือดสำหรับแปรรูปในครั้งนี้

แป้งมันเลือดที่ได้จากการผลิตแป้งที่ผ่านกรรมวิธีทำแห้งหลังจากการลวก มีองค์ประกอบทางเคมี (Proximate content) ดังตารางที่ 2 พบว่า แป้งมันเลือดมีความชื้น (moisture) ร้อยละ 7.05 ± 0.06 ซึ่งอยู่ในเกณฑ์คุณภาพของแป้งทั่วไปที่มีความชื้นในปริมาณต่ำ ปริมาณของแข็งทั้งหมด (Total Solid) ร้อยละ 87.35 ± 0.06 โดยน้ำหนักแห้ง มีปริมาณไขมันร้อยละ 0.39 ± 0.03 มีปริมาณเถ้าร้อยละ 3.61 ± 0.02 และมีปริมาณโปรตีน 1.60 ± 0.02 เปรียบเทียบกับองค์ประกอบทางเคมีของแป้งมันสำปะหลังที่ Charoenkul et al. (2011) วิเคราะห์ไว้ พบว่า แป้งมันสำปะหลังมีความชื้นร้อยละ 12.12 ± 0.09 ปริมาณของแข็งทั้งหมด 87.69 ± 0.08 โดยน้ำหนักแห้ง มีปริมาณเถ้าร้อยละ 0.19 ± 0.07 ส่วนปริมาณไขมันและปริมาณโปรตีนวิเคราะห์ไม่พบในแป้งมันสำปะหลัง

ตารางที่ 2 องค์ประกอบทางเคมีของแป้งมันเลือดเปรียบเทียบกับแป้งมันสำปะหลัง

องค์ประกอบทางเคมี (Proximate content) (n= 3)		
	แป้งมันเลือด	แป้งมันสำปะหลัง*
ความชื้น (Moisture) ร้อยละ	7.05 ± 0.06	12.12 ± 0.09
ปริมาณของแข็งทั้งหมด (Total Solid) ร้อยละ	87.35 ± 0.06	87.69 ± 0.08
เถ้า (Ash) ร้อยละ	3.61 ± 0.02	0.19 ± 0.07
ไขมัน (Crude Fat) ร้อยละ	0.39 ± 0.03	-
โปรตีน (Crude Protein) ร้อยละ	1.60 ± 0.02	-

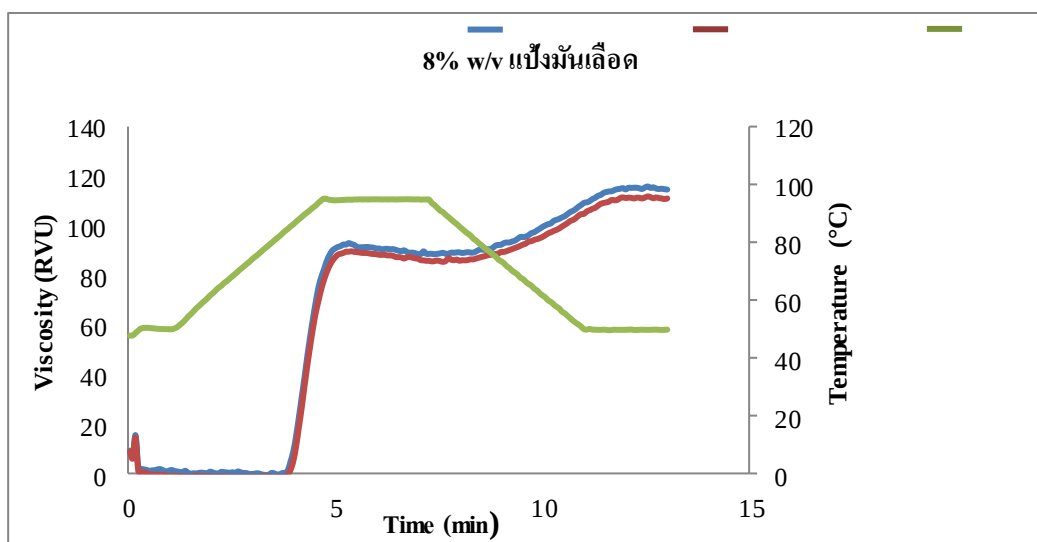
* หมายถึงอ้างอิงจากการวิเคราะห์ของ Charoenkul et al. (2011)

- ไม่พบ

จากผลการศึกษาองค์ประกอบทางเคมีของแป้งมันเลือดพบว่า แป้งมันเลือดมีเถ้า ไขมัน และโปรตีน ที่มากกว่าแป้งจากมันสำปะหลัง ซึ่งองค์ประกอบเหล่านี้ทำให้แป้งมีคุณสมบัติบางประการที่แตกต่างกันได้ เช่น เถ้า (ash content) บ่งบอกถึงปริมาณของสารอนินทรีย์หรือแร่ธาตุ (mineral content) ที่มีในแป้งซึ่งมีผลต่ออาหาร ทางด้านเนื้อสัมผัส รสชาติ รวมทั้งคุณค่าทางโภชนาการของอาหาร และเป็นสาเหตุในการทำให้เกิดคุณลักษณะ ทางเคมีเชิงฟิสิกส์ขึ้นระหว่างการปรุงอาหาร โดยเฉพาะค่า peak viscosity และ set back ของแป้ง ซึ่งเป็นความหนืด สูงสุดและการคืนตัวของแป้งหลังจากแป้งเย็นตัวลง (Charoenkul et al., 2011) ซึ่งส่งผลต่อลักษณะของอาหารที่ผลิต ได้ ส่วนโปรตีนเมื่อจับกับแป้งจะทำให้เกิดแรงยึดเหนี่ยวภายในเม็ดแป้งจากพันธะไดซัลไฟด์ในโปรตีนที่เกาะ เกี่ยวข้องกับเม็ดแป้งจึงมีผลต่อการกระจายตัวของแป้ง ความหนืด อัตราการดูดซับน้ำ อัตราการพองตัว และอัตราการเกิด เจลาติไนเซชันลดลง (Hamaker & Griffin, 1993; Srirot & Piyajomkwan, 2000) และในส่วนของไขมันในแป้งจะ รวมตัวกับโมเลกุลของอะไมโลสเกิดเป็นสารประกอบเชิงซ้อนเป็น โครงสร้างผลึกอย่างอ่อนที่เสริมสร้างความ แข็งแรงให้กับเม็ดแป้ง (Kasemsuwan et al., 1998) ซึ่งจะส่งผลต่อลักษณะและคุณสมบัติของแป้งโดยจะลด ความสามารถในการพองตัว การละลาย และการจับกับน้ำของแป้ง เมื่อแป้งผ่านความร้อนจะทำให้เกิดเป็น สารประกอบเชิงซ้อนที่ทำให้ได้ฟิล์มหรือแป้งเปียกที่มีลักษณะทึบแสงและขุ่น (Srirot & Piyajomkwan, 2000) องค์ประกอบทางเคมีเหล่านี้จะถูกนำไปใช้อธิบายคุณลักษณะของข้าวเกรียบมันเลือดที่ผลิตได้

2. ผลการวิเคราะห์ความหนืดของแป้งมันเลือด

จากการศึกษาคุณสมบัติทางเคมีกายภาพด้านความหนืดด้วยเครื่อง Rapid Visco Analyzer ได้กราฟแสดง ความหนืดของแป้งมันเลือดดังรูปที่ 1 และแสดงคุณสมบัติด้านต่างๆของความหนืด ดังแสดงในตารางที่ 2



รูปที่ 1 กราฟแสดงความหนืดของแป้งมันเลือด (M1(example 1), M2 (example 2), Temperature)

จากลักษณะของกราฟที่ได้เป็นกราฟแสดงความหนืดที่สัมพันธ์กับการเปลี่ยนแปลงของเม็ดแป้งดังนี้ ในช่วงแรกของการให้ความร้อน กราฟจะยังไม่แสดงความหนืดเนื่องจากเม็ดแป้งยังพองตัวน้อย เมื่ออุณหภูมิสูงขึ้น จนถึงจุดหนึ่งเรียกว่า pasting temperature กราฟจะเริ่มปรากฏความหนืดแสดงว่าโมเลกุลแป้งเกิดการรับน้ำเพิ่มมากขึ้น ทำให้เม็ดแป้งพองตัวขึ้นและเกิดความข้นหนืดจนถึงจุดที่เครื่องสามารถวัดได้ กราฟแสดงความหนืดสูงขึ้น จนถึงจุด peak viscosity จากนั้นความหนืดจะลดลง แสดงให้เห็นว่าเม็ดแป้งพองตัวมากขึ้นเรื่อยๆ จนเริ่มแตก เมื่อการแตกของเม็ดแป้งมีมากกว่าการพองตัวที่เพิ่มขึ้นและมีโมเลกุลอิสระละลายออกมาจะทำให้กราฟแสดงความหนืดลดลง ถ้าเม็ดแป้งไม่คงตัวและแตกง่าย ความหนืดจะยังคงลดลงเมื่ออุณหภูมิการหุงต้มไว้ที่ 95 °C ในช่วงการทำให้เย็นจากอุณหภูมิ 95 °C ไปเป็น 50 °C กราฟแสดงความหนืดสูงขึ้น การที่น้ำแป้งสุกมีความหนืดเพิ่มขึ้นนั้น เนื่องจากเกิดการคืนตัวของน้ำแป้งสุกโดยกลุ่มไฮดรอกซิลอิสระของเม็ดแป้งที่แตกหักและโมเลกุลอิสระที่ละลายออกมาในน้ำแป้งสุกโดยเฉพาะโมเลกุลอะไมโลสที่มีขนาดที่เหมาะสมจะจับกันด้วยพันธะไฮโดรเจนใหม่และกักน้ำเอาไว้ได้ (Saattrat, 2015)

ตารางที่ 2 ผลการวิเคราะห์ความหนืดของแป้งมันเลือดเปรียบเทียบกับแป้งมันสำปะหลัง

Pasting properties (n=2)		
	Purple yam flour	Cassava starch*
Peak viscosity (RVU)	91.92 ± 2.24	96-138
Holding strength (RVU)	87.54 ± 2.06	63- 84
Breakdown (RVU)	4.38 ± 0.18	11-78
Final viscosity (RVU)	113.42 ± 2.48	29-133
Peak time (min)	5.37 ± 0.05	-
Pasting temperature (°C)	82.38 ± 0.04	67.4-70.4
Setback (RVU)	25.88 ± 0.41	18-55

* อ้างอิงจากผลการวิเคราะห์ความหนืดแป้งมันสำปะหลังสายพันธุ์ต่างๆ ของ Charoenkul et al. (2011)

จากตารางที่ 2 เป็นค่าเฉลี่ยที่ได้จากการวิเคราะห์ความหนืดของแป้งด้วยวิธี RVA พบว่า กราฟความหนืดที่ได้จากการวิเคราะห์แป้งมันเลือดมีลักษณะจำเพาะที่สำคัญคือค่า pasting temperature ซึ่งเป็นอุณหภูมิที่เกิดการเปลี่ยนแปลงค่าความหนืด มีค่าค่อนข้างสูงประมาณ 82.38 ± 0.04 °C สูงกว่าแป้งข้าวยาคู แป้งมันสำปะหลัง และแป้งข้าวที่มีอุณหภูมิอยู่ในช่วง 69.40 ถึง 74.22 °C (Rittiruangdej & Suwannasichin, 2007) อาจเนื่องจากลักษณะการจัดเรียงตัวของโมเลกุลทั้งอะไมโลสและอะไมโลเพกตินที่อยู่ภายในเม็ดแป้งมีการจัดเรียงตัวกันอย่างเป็นระเบียบและหนาแน่น พันธะเคมีภายใน (พันธะไฮโดรเจน) จึงมีความแข็งแรง ทำให้ต้องใช้ความร้อนสูงในการทำลายโครงสร้างที่เป็นระเบียบนั้น ส่วนค่า peak viscosity ของแป้งมันเลือดเท่ากับ 91.92 ± 2.24 RVU ซึ่งเป็นค่าที่ต่ำกว่า

แป้งมันเทศที่มีค่า 101 RVU แต่ใกล้เคียงกับแป้งมันสำปะหลังมีค่า 96 RVU (Charoenkul et al., 2011) และสูงกว่าแป้งข้าวโพดที่มีค่า 49 RVU (Suzuki et al., 1993) จึงจัดได้ว่าแป้งมันเลือดเป็นแป้งที่มีความหนืดค่อนข้างสูงเนื่องจากเม็ดแป้งที่พองตัวมีความคงทนต่อการกวน จึงทำให้แป้งมันเลือดมีค่า breakdown ต่ำ มีค่าเท่ากับ 4.38 ± 0.18 RVU ค่า setback ซึ่งเป็นค่าที่บ่งชี้ถึงการคืนตัวของแป้งสุกพบว่าแป้งมันเลือดมีค่าเท่ากับ 25.88 ± 0.41 RVU ซึ่งค่า setback ที่วิเคราะห์ได้มีค่าใกล้เคียงกับแป้งมันสำปะหลังที่มี setback 31 RVU ซึ่งต่างจากแป้งถั่วเขียวที่มี setback 115 RVU ที่จัดเป็นแป้งที่มีปริมาณอะไมโลสสูงตามที่ Perez et al. (1997, 1998) ได้เคยรายงานเอาไว้

ลักษณะกราฟที่ได้จากการวิเคราะห์ความหนืดของแป้งมันเลือดมีลักษณะกราฟรูปแบบซี (C-Type) ตามที่ Schoch & Maywald (1968) ได้จำแนกเอาไว้ แป้งที่ให้กราฟรูปแบบซีเป็นแป้งที่มีการพองตัวจำกัด (Restricted-swelling starch) และเป็นแป้งที่ทนทานต่อแรงเฉือนจากการกวน จัดเป็นแป้งที่มีปริมาณอะไมโลสต่ำกว่าปริมาณอะไมโลเพคติน เช่นเดียวกับงานวิจัยของ Harijono et al. (2013) พบว่า ปริมาณของอะไมโลสในแป้งเป็นคุณสมบัติที่สำคัญอย่างหนึ่งของแป้ง ซึ่งเป็นส่วนสำคัญที่ทำให้เกิดลักษณะของอาหารที่แตกต่างกัน เช่นเดียวกับคุณสมบัติด้านอื่นๆ ของแป้ง เช่น การพองตัว การดูดซับน้ำ การละลายของแป้ง และความหนืดของแป้งแล้วแต่มีความจำเพาะต่อการเกิดลักษณะของผลิตภัณฑ์อาหารทั้งสิ้น จากลักษณะความหนืดของแป้งมันเลือดที่ได้ มีความหนืดใกล้เคียงกับค่าความหนืดของแป้งมันสำปะหลังที่ Charoenkul et al. (2011) วิเคราะห์ได้ จึงจัดแป้งมันเลือดเป็นแป้งที่เหมาะสมต่อการนำมาทำเป็นผลิตภัณฑ์ขนมปังกรอบ (snack) เช่นเดียวกับแป้งมันสำปะหลัง (Sesang, 2020) ผู้วิจัยจึงนำแป้งมันเลือดมาศึกษาเพื่อพัฒนาเป็นผลิตภัณฑ์ข้าวเกรียบต่อไป

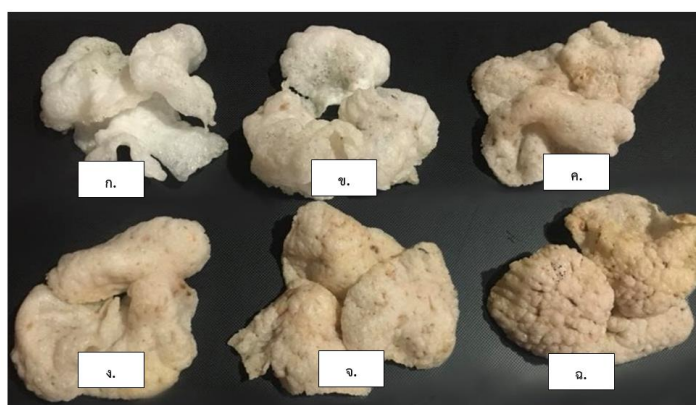
3. การผลิตข้าวเกรียบจากแป้งมันเลือด

จากการทดลองทำผลิตภัณฑ์ข้าวเกรียบทั้ง 6 สูตร พบว่า เมื่อใช้แป้งมันเลือดทดแทนแป้งมันสำปะหลังในอัตราส่วนที่เพิ่มมากขึ้นจะทำให้การนวดแป้งยากขึ้น แป้งไม่จับตัวกันเนื่องจากแป้งมันเลือดไม่เกิดการพองตัว หรือมีการพองตัวที่ค่อนข้างจำกัด (ตามลักษณะความหนืดที่วิเคราะห์ได้) เมื่อใช้แป้งมันเลือดปริมาณมากขึ้นในข้าวเกรียบแต่ละสูตรทำให้แป้งที่นวดได้มีความแข็งมากขึ้นเรื่อยๆ จนเกือบร่วน นอกจากนี้แป้งมันเลือดไม่มีกลูเตน เช่นเดียวกับแป้งจากมันพื้นบ้านชนิดอื่นจึงไม่เกิดโครงสร้างที่จับกันเป็นก้อนเหมือนแป้งสาลี (Chen et al., 2019) ที่สามารถนวดและขึ้นรูปเป็นฟิล์มได้ เมื่อนำแป้งที่นวดแล้วไปนึ่งจนสุกพบว่าก้อนแป้งที่มีอัตราส่วนของแป้งมันเลือดในปริมาณมากจะให้ก้อนแป้งแน่นและแข็งกว่าก้อนแป้งที่มีปริมาณแป้งมันเลือดน้อยกว่า ไม่มีความยืดหยุ่นและไม่ใส เมื่อนำมาหั่นเป็นแผ่นและตากแดดจนแห้งสนิทจะได้ผลิตภัณฑ์ข้าวเกรียบพร้อมนำไปทอด ดังแสดงในรูปที่ 2



รูปที่ 2 ข้าวเกรียบจากแป้งมันเลือดสูตรต่างๆ ก่อนทอด ในอัตราส่วนแป้งมันเลือดต่อแป้งมันสำปะหลัง (ก.) 0 : 100 (ข.) 20 : 80 (ค.) 40 : 60 (ง.) 60 : 40 (จ.) 80 : 20 และ (ฉ.) 100 : 0

จากรูปที่ 2 พบว่า ข้าวเกรียบมีสีน้ำตาลอมม่วงโดยมีสีเข้มขึ้นตามปริมาณที่มากขึ้นของแป้งมันเลือดที่ทดแทนลงไป และความใสลดลงตามลำดับเช่นกันตามปริมาณของไขมันที่มีอยู่ในแป้งมันเลือดที่มีเพิ่มมากขึ้นเพราะไขมันเป็นส่วนที่ทำให้เกิดความทึบแสงหรือความขุ่นนั่นเอง (Srirot & Piyajomkwan, 2000) ความละเอียดของเนื้อข้าวเกรียบจะหยาบขึ้นตามปริมาณแป้งมันเลือดที่เพิ่มมากขึ้น ซึ่งโดยปกติแล้วแป้งฟลาวจะมีส่วนประกอบของเส้นใยมากกว่าแป้งสตราซ์ ซึ่งมีผลต่อความเนียนละเอียดของเนื้อข้าวเกรียบ และพบว่าข้าวเกรียบที่ทำจากแป้งมันเลือดร้อยละ 100 เนื้อแป้งข้าวเกรียบที่ได้จะไม่ใส เนื้อแป้งหยาบมาก และเนื้อแป้งไม่เกาะกัน เมื่อนำข้าวเกรียบแต่ละสูตรไปทอดพบว่าข้าวเกรียบจะมีการพองตัวลดลงเรื่อยๆ เนื้อสัมผัสจะแข็งขึ้นเรื่อยๆ ตามปริมาณที่เพิ่มขึ้นของแป้งมันเลือดที่ทดแทนลงไป ลักษณะของข้าวเกรียบแต่ละสูตรแสดงดังรูปที่ 3



รูปที่ 3 ลักษณะของข้าวเกรียบจากแป้งมันเลือดสูตรต่าง ๆ หลังทอด ในอัตราส่วนแป้งมันเลือดต่อแป้งมันสำปะหลัง (ก.) 0 : 100 (ข.) 20 : 80 (ค.) 40 : 60 (ง.) 60 : 40 (จ.) 80 : 20 และ (ฉ.) 100 : 0

จากลักษณะของข้าวเกรียบที่ทอดแล้วพบว่าปริมาณแป้งมันเลือดที่ทดแทนลงไปมีผลต่อการผลิตข้าวเกรียบในแต่ละสูตร ทั้งนี้เนื่องจากแป้งมันเลือดจัดเป็นแป้งฟลาว (flour) ซึ่งเป็นแป้งที่ได้จากกระบวนการบดแห้งทำให้แป้งชนิดฟลาวมีองค์ประกอบทางเคมี เช่น แร่ธาตุ คากโบ และสารประกอบทางธรรมชาติอื่นๆ อยู่สูงกว่าแป้งชนิดที่เป็นสตาร์ช (starch) เช่นเดียวกับแป้งจากหัวมันชนิดอื่นตามที่ Asiyanbi-Hammed (2016) ได้ศึกษาไว้ และพบว่าองค์ประกอบทางเคมีต่างๆ ในแป้งจะส่งผลต่อเนื้อสัมผัสและลักษณะปรากฏที่เกิดขึ้นในผลิตภัณฑ์ ดังนั้นองค์ประกอบของแป้งมันเลือดที่วิเคราะห์ได้ ทั้งปริมาณเถ้า ไขมัน และ โปรตีน ที่มีมากกว่าในแป้งมันสำปะหลังจึงส่งผลต่อลักษณะปรากฏของข้าวเกรียบมันเลือดที่ผลิตได้ และคุณภาพทางประสาทสัมผัส ซึ่งผลการทดสอบคุณภาพทางประสาทสัมผัสแสดงในหัวข้อที่ 5

5. ผลการทดสอบคุณภาพทางประสาทสัมผัสของผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือด

ผลการทดสอบทางประสาทสัมผัสของผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือด โดยการทดสอบความชอบหรือระดับความพอใจโดยใช้ผู้ทดสอบชิมจำนวน 15 คน โดยทดสอบทางประสาทสัมผัสแบบ 9-Point Hedonic Scoring test และวิเคราะห์ผลทางสถิติ โดยวิธีวิเคราะห์ความแปรปรวน (Analysis Of Variance : ANOVA) ที่ระดับความเชื่อมั่นร้อยละ 95 ($P < 0.05$) ทางด้านลักษณะปรากฏ สี กลิ่นรส ความกรอบ และความชอบโดยรวมของผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือด แสดงดังตารางที่ 3 พบว่า คุณลักษณะทางประสาทสัมผัสด้านลักษณะปรากฏมีความแตกต่างกันโดยพบว่า ผู้บริโภครอบลักษณะปรากฏสูตรที่ 3 มากที่สุด และ ชอบสูตรที่ 6 น้อยที่สุด ด้านสีของผลิตภัณฑ์ข้าวเกรียบมีความชอบแตกต่างกัน โดยพบว่าสีของข้าวเกรียบสูตรที่ 4 ผู้บริโภครอบมากที่สุด และชอบสูตรที่ 6 น้อยที่สุด ทางด้านกลิ่นรสผู้บริโภครอบสูตรที่ 2 มากที่สุดและชอบสูตรที่ 6 น้อยที่สุด คุณลักษณะด้านความกรอบ ผู้บริโภครอบสูตรที่ 3 มากที่สุด และชอบสูตรที่ 6 น้อยที่สุด และคุณลักษณะทางประสาทสัมผัสด้านความชอบโดยรวม พบว่าสูตรที่ให้การยอมรับมากที่สุดคือสูตรที่ 3 และให้การยอมรับน้อยที่สุดคือสูตรที่ 6 ดังตารางที่ 3

ตารางที่ 3 ผลการทดสอบคุณภาพทางประสาทสัมผัสจากการทดสอบชิมข้าวเกรียบจากแป้งมันเลือด

สูตร (แป้งมันเลือด:แป้งมันสำปะหลัง)	คุณลักษณะต่างๆ				
	ลักษณะปรากฏ	สี	กลิ่น	ความกรอบ	ความชอบโดยรวม
1 (0:100)	5.06±2.74 ^c	5.38±2.81 ^{de}	3.63±1.45 ^b	6.29±1.68 ^{cd}	5.73±2.31 ^c
2 (20:80)	4.44±2.09 ^b	4.92±2.02 ^c	5.50±2.55 ^d	6.14±1.92 ^c	4.87±2.36 ^b
3 (40:60)	5.81±2.20 ^d	5.00±2.58 ^d	4.81±2.37 ^c	6.57±2.31 ^d	6.40±2.23 ^d
4 (60:40)	4.44±2.30 ^b	5.69±2.69 ^c	4.56±1.46 ^c	5.93±1.81 ^c	6.13±2.03 ^{cd}
5 (80:20)	4.56±1.96 ^b	4.31±2.50 ^b	4.44±2.06 ^c	4.71±2.40 ^b	3.80±2.18 ^b
6 (100:0)	3.19±1.51 ^a	3.46±2.10 ^a	2.69±1.45 ^a	3.07±2.02 ^a	3.33±2.38 ^a

หมายเหตุ - ข้อมูลแสดงในรูปค่าเฉลี่ย+ส่วนเบี่ยงเบนมาตรฐาน

- abc ตัวอักษรที่ต่างกันในแนวตั้งหมายถึงค่าเฉลี่ยมีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ($P < 0.05$)

จะเห็นได้ว่าผลการทดสอบคุณภาพทางประสาทสัมผัสมีคะแนนค่อนข้างต่ำเนื่องจากสูตรต้นแบบได้ถูกดัดแปลงมาจากการวิจัยอื่นเพียงส่วนของวิธีทำและส่วนผสมหลัก ซึ่งอาจส่งผลให้สูตรต้นแบบได้คะแนนการยอมรับต่ำกว่า และเมื่อทดสอบการชิมข้าวเกรียบมันเลือดพบว่าผลการทดสอบคุณภาพทางประสาทสัมผัสโดยส่วนใหญ่มีความแตกต่างกันทางสถิติที่ระดับความเชื่อมั่นร้อยละ 95 ยกเว้นความกรอบและความชอบโดยรวม ข้าวเกรียบมันเลือดสูตรที่ 3 และ 4 ซึ่งมีปริมาณแป้งมันเลือดร้อยละ 40 และ 60 ตามลำดับ มีคะแนนความชอบค่อนข้างสูง เนื่องจากข้าวเกรียบมีลักษณะโดยรวมที่ดี มีความกรอบ มีสี และกลิ่นที่ดีกว่า โดยผู้บริโภคให้คะแนนความกรอบสูตรที่ 3 ที่มีแป้งมันเลือดผสมอยู่ร้อยละ 40 เพราะข้าวเกรียบมีลักษณะกรอบไม่แข็งกระด้างจนเกินไปเมื่อเทียบกับสูตรที่ 4, 5 และ 6 ความกรอบที่เกิดจากการการพองตัวของแป้งซึ่งการพองตัวที่เกิดขึ้นนั้นมาจากปริมาณแป้งมันสำปะหลังที่มีอยู่ในสูตรมากกว่า จากคุณสมบัติของแป้งมันสำปะหลังทำให้ผลิตภัณฑ์ที่ทอดเกิดการพองตัวมากกว่าแป้งชนิดอื่น เพราะมันสำปะหลังมีองค์ประกอบทางเคมีที่ขัดขวางการพองตัวของแป้งในปริมาณที่น้อยกว่า ซึ่งก็คือปริมาณของเถ้า โปรตีน และไขมัน ซึ่งมีผลต่อเนื้อสัมผัสของผลิตภัณฑ์ทำให้เกิดลักษณะดังกล่าวขึ้นได้ (Asiyanbi-Hammed & Simsek (2016) นอกจากนี้แป้งสตาร์ชเช่นแป้งมันสำปะหลังมีองค์ประกอบหลักที่สำคัญที่ทำให้ข้าวเกรียบเกิดลักษณะฟู และกรอบ นั่นก็คือ สัดส่วนของอะไมโลสและอะไมโลเพกตินในเมล็ดแป้งหากมีปริมาณอะไมโลเพกตินสูงจะทำให้ผลิตภัณฑ์มีการพองตัวดี แป้งที่มีอะไมโลเพกตินสูงเมื่อบดแป้งแตกตัวได้ง่ายเมื่อทำเป็นข้าวเกรียบจะพองตัวได้มากและยังช่วยให้ผลิตภัณฑ์พองตัวมีลักษณะโปร่งเบา และในทางกลับกันแป้งที่เกิดเจลมีความหนืดสูงเพราะแป้งมีโปรตีนประกอบอยู่มาก (Srirot & Piyajomkwan, 2000) เช่นในแป้งมันเลือด จะส่งผลให้ผลิตภัณฑ์มีเนื้อสัมผัสแข็งและมีข้อจำกัดในการพองตัว (Chainui, 2007) ข้าวเกรียบมันเลือดที่ผลิตได้จึงแข็งกระด้างและไม่กรอบหากมีแป้งมันเลือดปริมาณมากในสูตร ผู้บริโภคจึงให้คะแนนการยอมรับคุณลักษณะด้านความชอบโดยรวมในข้าวเกรียบสูตรที่ 3 ซึ่งมีแป้งมันเลือดผสมอยู่เพียงร้อยละ 40 มากที่สุด เนื้อสัมผัสไม่แข็งหรือโปร่งเบาจนเกินไป ส่วนความชอบของสีที่ได้จากการชิมข้าวเกรียบมันเลือดจะเห็นได้ว่าสูตรที่มีปริมาณแป้งมันเลือดมากกว่าจะได้รับคะแนนการยอมรับต่ำกว่าเนื่องจากมีสีคล้ำขึ้น ส่วนกลิ่นรสของสูตรที่ 2 ผู้ชิมให้การยอมรับมากที่สุดที่มีปริมาณแป้งมันเลือดเพียงร้อยละ 20 เพราะแป้งมันเลือดที่ผลิตได้มีกลิ่นค่อนข้างอ่อนเช่นเดียวกับแป้งฟลาวาที่ได้จากมันพื้นบ้านหลายชนิด เช่น แป้งแห้ว หรือ แป้งสาลี เพราะมีส่วนประกอบทางเคมีค่อนข้างสูงผสมอยู่ด้วยจึงทำให้แป้งมีกลิ่นไม่พึงประสงค์ (Healthline, 2020) เมื่อใส่แป้งมันเลือดในปริมาณที่มากขึ้นจึงทำให้ข้าวเกรียบมีกลิ่นไม่ดีตามไปด้วยโดยเฉพาะกลิ่นที่เกิดจากปฏิกิริยาการหืนของไขมัน (rancidity) ที่มีอยู่ในแป้งมันเลือด (Alandete L.U., 2018)

6. ผลการศึกษาองค์ประกอบทางเคมีและพลังงานทั้งหมดของข้าวเกรียบจากแป้งมันเลือด

หลังจากทดสอบคุณภาพทางประสาทสัมผัสของข้าวเกรียบจากแป้งมันเลือดจะทำให้ผู้วิจัยสามารถเลือกสูตรที่มีคะแนนการยอมรับสูงที่สุดมาใช้เพื่อวิเคราะห์องค์ประกอบทางเคมีและพลังงานทั้งหมดที่ได้รับ โดยพบว่าผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือดในสูตรที่ 3 ที่ใช้แป้งมันเลือดย่อยละ 40 ทดแทนแป้งมันสำปะหลังในสูตรเหมาะสมที่สุด ดังแสดงในตารางที่ 4

ตารางที่ 4 ผลการวิเคราะห์องค์ประกอบทางเคมีและพลังงานที่ได้รับของข้าวเกรียบจากแป้งมันเลือดหลังทอด

องค์ประกอบทางเคมี (g/100g) (n = 3)	
ความชื้น	1.73 ± 0.04
คาร์โบไฮเดรต	65.37 ± 1.49
เถ้า	12.42 ± 0.62
ไขมัน	17.87 ± 0.77
โปรตีน	2.61 ± 0.32
เส้นใย	1.41 ± 0.47
พลังงานที่ได้รับ (kcal/100g)	466.39 ± 0.96

ผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือดหลังจากทอดแล้ว 100 กรัม มีปริมาณความชื้นร้อยละ 1.73 ± 0.04 ซึ่งเป็นไปตามมาตรฐานที่สำนักงานมาตรฐานผลิตภัณฑ์อุตสาหกรรม (2554) (Thai Industrial Standards Institute, 2011) กำหนดไว้ว่าข้าวเกรียบพร้อมบริโภคควรมีความชื้นไม่เกินร้อยละ 4 ปริมาณคาร์โบไฮเดรตร้อยละ 65.37 ± 1.49 มีปริมาณมากที่สุดตามสัดส่วนที่มีอยู่ในสูตร เถ้าร้อยละ 12.42 ± 0.62 จะเห็นว่าปริมาณเถ้าสูงทั้งนี้เนื่องจากแป้งจากพืชหัวจะมีแร่ธาตุในปริมาณมากกว่าแป้งชนิดอื่นเช่นเดียวกับที่ Akinola (2020) ได้ศึกษาเอาไว้ พบว่าปริมาณเถ้าที่ได้จากมันพื้นบ้านมีค่อนข้างสูงกว่าเถ้าจากแป้งชนิดอื่นทั่วไป แป้งจากมันพื้นบ้านมีแร่ธาตุอยู่ค่อนข้างสูงทั้งโพแทสเซียม เหล็ก และ สังกะสี ซึ่งแร่ธาตุเหล่านี้จะส่งผลให้มีปริมาณเถ้าสูง และเมื่อนำมาทำผลิตภัณฑ์ที่มีการใช้ส่วนผสมอื่น เช่น เกลือแกง (NaCl) เป็นส่วนผสมอยู่ด้วยจะทำให้ได้เถ้าที่เป็นโซเดียมออกไซด์ (Na₂O) (Zhang et al., 2015) เพิ่มมากขึ้นไปอีกส่งผลให้ปริมาณเถ้ามีมากขึ้นตามไปด้วย ไขมันร้อยละ 17.87 ± 0.77 ได้จากน้ำมันที่ใช้ทอดซึ่งมีผลต่อพลังงานที่ได้รับในข้าวเกรียบรองลงมาจากคาร์โบไฮเดรตที่มีอยู่ในข้าวเกรียบ มีโปรตีนอยู่ร้อยละ

2.61 ± 0.32 ถือได้ว่ามีปริมาณค่อนข้างต่ำ จะเห็นได้ว่าแป้งจากพืชหัวส่วนใหญ่จะมีโปรตีนน้อยกว่าที่ได้จากแป้งข้าวหรือธัญพืชอื่นๆ ซึ่งปริมาณโปรตีนในข้าวเกรียบจากแป้งมันเลือดที่ได้มีค่าใกล้เคียงกับ Ayodele et al. (2013) ได้ศึกษาเอาไว้ และมีเส้นใยร้อยละ 1.41 ± 0.47 โดยส่วนใหญ่แป้งฟลาวจะให้ปริมาณเส้นใยสูงกว่าแป้งสตาร์ช ในสูตรที่มีแป้งมันเลือดและแป้งมันสำปะหลังผสมกันอยู่ในอัตราส่วน 40 : 60 ทำให้ปริมาณเส้นใยที่ได้จากข้าวเกรียบมีปริมาณต่ำ ข้าวเกรียบมันเลือดหลังจากทอดแล้ว 100 กรัม ให้พลังงานทั้งหมด 466.39 ± 0.96 กิโลแคลอรี ถือเป็นผลิตภัณฑ์อาหารว่างที่ให้พลังงานสูงเนื่องจากส่วนประกอบส่วนใหญ่เป็นคาร์โบไฮเดรตและไขมันซึ่งเป็นสารอาหารที่ให้พลังงานสูงจึงควรบริโภคในปริมาณน้อยหรือควรแบ่งบริโภค

สรุปผลการวิจัย

การพัฒนาผลิตภัณฑ์ข้าวเกรียบจากแป้งมันเลือดทำขึ้นจากแป้งมันเลือดที่ผลิตโดยวิธีการทำแห้งหลังจากการลวก เมื่อวิเคราะห์คุณสมบัติด้านความหนืดของแป้งจะให้กราฟเป็นรูปแบบซี (c-type) ที่แสดงว่ามีปริมาณอะไมโลสต่ำ ลักษณะของแป้งดังกล่าวเหมาะต่อการนำมาทำผลิตภัณฑ์ขนมปังกรอบ และเมื่อนำมาผลิตข้าวเกรียบพบว่าข้าวเกรียบจากแป้งมันเลือดสูตรที่เหมาะสมที่สุดคือสูตรที่ 3 ที่ใช้แป้งมันเลือดทดแทนแป้งมันสำปะหลังในอัตราส่วนร้อยละ 40 เป็นสูตรที่ได้คะแนนการยอมรับทางประสาทสัมผัสสูงกว่าสูตรอื่นๆ ทั้งด้านลักษณะปรากฏ ความกรอบ สี และความชอบโดยรวม และเมื่อนำไปทอดพร้อมบริโภคจะให้พลังงานที่ได้รับทั้งหมดเท่ากับ 466.39 กิโลแคลอรีต่อ 100 กรัม ซึ่งมีปริมาณความชื้น ปริมาณของแข็งทั้งหมด เถ้า ไขมัน โปรตีน และ เส้นใย ร้อยละ 1.73 ± 0.04, 65.37 ± 1.49, 12.42 ± 0.62, 17.87 ± 0.77, 2.612 ± 0.32 และ 1.41 ± 0.47 ตามลำดับ

เอกสารอ้างอิง

- Akinola, O. (2020). Nutritional composition and physio-chemical properties of peeled and unpeeled yam flour (White Yam, *Dioscorea rotundata*). *Current Developments in Nutrition*, 4(2), 735. https://doi.org/10.1093/cdn/nzaa052_004
- Alandete, L. U. (2018). *Lipid degradation during grain storage: markers, mechanisms and shelf-life extension treatments* [Doctoral's thesis, The University of Queensland]. UQ eSpace. <https://espace.library.uq.edu.au/view/UQ:6e94099>
- AOAC. (1995). *Official Methods of Analysis* (16th ed.). Association of Official Analytical Chemists.
- Asiyanbi-Hammed, T. T., & Simsek, S. (2016). Comparison of physical and chemical properties of wheat flour, fermented yam flour, and unfermented yam flour. *Journal of Food Processing and Preservation*, 42(12), e13844. <https://doi.org/10.1111/jfpp.13844>

- Asiyanbi-Hammed, T. T. (2016). *Characteristict of yam composite flour: properties and function of bread and tortilla making* [Doctoral's thesis, University of Agriculture and Applied Science].
[https://www.scirp.org/\(S\(351jmbntv-nsjt1aadkposzje\)\)/reference/referencespapers.aspx?referenceid=2711224](https://www.scirp.org/(S(351jmbntv-nsjt1aadkposzje))/reference/referencespapers.aspx?referenceid=2711224)
- Ayodele, B., Bolade, M., & Usman, M. (2013). Quality characteristics and acceptability of 'amala' (yam- based thick paste) as influenced by particle size categorization of yam (*Dioscorea rotundata*) flour. *Food Science and Technology International*, 19, 35-43. <https://doi.org/10.1177/1082013212442181>
- Cakrawati, D., Srivichai, S., & Hongsprabhas, P. (2021). Effect of steam-cooking on (poly) phenolic compounds in purple yam and purple sweet potato tubers. *Food Research*, 5(1), 330-336.
[https://doi.org/10.26656/fr.2017.5\(1\).407](https://doi.org/10.26656/fr.2017.5(1).407) (in Thai)
- Capotorto, I., Amodio, M. L., Diaz, M. T. B., de Chiara, M. L. V., & Colelli, G. (2018). Effect of anti-browning solutions on quality of fresh-cut fennel during Storage. *Postharvest Biology and Technology*, 137, 21-30.
- Chainui, C. (2007). *Effects of Physical chemistry of starch mixtures (cassava starch cassava and sago flour) on the quality of Rice cracker* [Unpublished Master's thesis]. Prince of Songkla University (in Thai)
- Chansri, R., Chatthongphisut, R., & Chanthakhet, R. (2019). The study on physical and chemical properties of Dioscorea alata L. flour and The value added of Dioscorea alata L. In *the 10th Rajamangala Surin National Conferenc: "Research and Innovation for Sustainable Development"* (pp. 251-261). Rajamangala University of Technology Isan Surin Campus, Surin Province.
- Charoenkul, N., Uttapap, D., Pathipanawat, W., & Takeda, Y. (2011). Physicochemical characteristics of starches and flours from cassava varieties having different cooked root textures. *LWT - Food Science and Technology*, 44, 1774-1781. <https://doi.org/10.1016/j.lwt.2011.03.009>
- Chen, G., Ehmke, L., Sharma, C., Miller, R., Faa, P., Smith, G., & Li, Y. (2019). Physicochemical properties and gluten structures of hard wheat flour doughs as affected by salt. *Food chemistry*, 275, 569-576.
<https://doi.org/10.1016/j.foodchem.2018.07.157>
- Chen, X., Lu, J., Li, X., Wang, Y., Miao, J., Mao, X., & Gao, W. (2017). Effect of blanching and drying temperatures on starch-related physicochemical properties, bioactive components and antioxidant activities of yam flours. *LWT-Food Science and Technology*, 82, 303-310. <https://doi.org/10.1016/j.lwt.2017.04.058>

- Chen, Y., Fan, J., Yi, F., Luo, Z., & Fu, Y. (2003). Rapid clonal propagation of *Dioscorea zingiberensis*. *Plant Cell Tissue and Organ Culture*, 73(1), 75-80.
- Cory, H., Passarelli, S., Szeto J., Tamez, M., & Mattei, J. (2018). The role of polyphenols in human health and food systems: A mini-review. *Frontiers in Nutrition*, 5, e00087. <https://doi.org/10.3389/fnut.2018.00087>
- Hamaker, B. R., & Griffin, V. K. (1993). Effect of Disulfide Bond-Containing Protein on Rice Starch Gelatinization and Pasting. *Cereal Chemistry*, 70, 377-380.
- Harijono, T. E., Saputri, D., & Kusnadi, J. (2013). Effect of Blanching on Properties of Water Yam (*Dioscorea alata*) Flour. *Advance Journal of Food Science and Technology*, 5(10), 1342-1350. <https://doi.org/10.19026/ajfst.5.3108>
- Healthline. (2020). Does Flour Go Bad? <https://www.ecowatch.Com/three-films-for-earth-week2652609187.html>
- Homhuan, R. (2017). Useful of the local yam. *Journal of Kaset A-pirom*, 18(3), 36-38 (in Thai)
- Kasemsuwan T., Bailey T., & Jane J. (1998). Preparation of clear noodles with mixtures of tapioca and high-amylose starches. *Carbohydrate Polymers*, 32, 301-312. [https://doi.org/10.1016/S0144-8617\(97\)002567](https://doi.org/10.1016/S0144-8617(97)002567)
- Larief, R., Dirpan, A., & Theresia. (2018). Purple Yam Flour (*Dioscorea alata* L.) processing effect on Anthocyanin and Antioxidant Capacity in Traditional Cake “Bolu Cukke” Making. In *IOP Conference Series: Earth and Environmental Science*. <https://iopscience.iop.org/article/10.1088/17551315/207/1/012043>
- Lebot, V. (2020). *Tropical Root and Tuber Crops: Cassava, Sweet Potato, Yams and Aroids* (2nd ed.). CABI publisher.
- Liu, X., Lu, K., Yu, J., Copeland, L., Wang, S., & Wang, S. (2019). Effect of purple yam flour substitution for wheat flour on in vitro starch digestibility of wheat bread. *Food Chem*, 284, 118-124. <https://doi.org/10.1016/j.foodchem.2019.01.025>
- Metsdagh, F., Wilde, T. D., Delporte, K., Peteghem, C. V., & Meulenaer, B. D. (2008). Impact of chemical pre-treatments on the acrylamide formation and sensorial quality of potato crisps. *Food Chemistry*, 106, 914 - 922. <https://doi.org/10.1016/j.foodchem.2007.07.001>

- Oko, A. O., & Famurewa, A. C. (2015). Estimation of nutritional and starch characteristics of *Dioscorea alata* (Water Yam) varieties commonly cultivated in the south-eastern Nigeria. *Current Journal of Applied Science and Technology*, 145-152. <https://doi.org/10.9734/BJAST/2015/14095>
- Perez, E., Breene, W. M., & Bahnassey, Y. A. (1998), Variations in the Gelatinization Profiles of Cassawa, Sagu and Arrowroot Native Starch as Measured with Different Thermal and Mechanical Methods. *Starch/Staerke*, 50(2-3), 70-72.
- Perez, E., Lares, M., & Gonzalez, Z., (1997). Some Characteristics of sago (*Canna edulis* Kerr) and Zulu (*Maranta* sp.) Rhizome. *Journal of Agriculture Food Chemistry*, 45(7), 2546-2549.
- Phongnoree, J., & Thammasiri, B. (2008). Product development of lotus root cracker. In *the 2nd Kamphaeng Phet Rajanhat Student University Student National Conference* (pp. 847-853). Valaya Alongkorn Rajabhat University. Pathum Thani. (in Thai)
- Promdang, S., Homhuan, R., Wongmaneroaj, M., & Jamjumras, S. (2021). *The nutritional essence of yam*. In *The 18th National Academic Conference*, Kasetsart University, Kamphaeng Saen Campus. (in Thai)
- Rittiruangdej, P., & Suwannasichin, T. (2007). Component analysis of side behavior Viscosity and texture properties of gels in different powders. In *Proceedings of 45th Kasetsart University Annual Conference: Agricultural Extension and Home Economics*, Agro-Industry, Kasetsart University, Bangkok. (in Thai)
- Rosales, A. L., & Barrion, A. S. A. (2019). Effect of boiling on carbohydrate profile, starch digestibility, and glycemic index of Purple Yam (*Dioscorea alata* L.) powder. In *The 2nd International Conference on Nutrition*, Food Science and Technology, University of the Philippines, Philippines.
- Saartrat, S., Uttapap, D., Puttanlek, C., & Rungsardthong, V. (2015). Edible canna (*Canna edulis*) starch modified by acetylation. *KMUTT Research and Development Journal*, 28(1), 87-102. (in Thai)
- Satour, M., Mitaine-Offer, A. C., & Lacaille-Dubois, M. A. (2007). The *Dioscorea* genus: A review of bioactive steroid saponins. *Journal of Natural Medicines*, 61, 91-101.
- Schoch, T. J. & Maywald, F. C. (1968). Preparation and Properties of Various Legume Starch. *Cereal Chemistry*, 45, 564 - 573.
- Suriya, M., Baranwal, G., Bashir, M., Reddy, C. K., & Haripriya, S. (2016). Influence of blanching and drying methods on molecular structure and functional properties of elephant foot yam (*Amorphophallus paeoniifolius*) flour. *LWT-Food Science and Technology*, 68, 235-243. <https://doi.org/10.1016/j.lwt.2015.11.060>

- Sesang, S. (2020). *Properties of starch affecting the processing of aquatic products*. Aquaculture Industry Technology Research and Development Division, Department of Fisheries. (in Thai)
- Srirot, K., & Piyajomkwan, K. (2000). *Starch technology* (2nd ed). Kasetsart University Press.
- Suzuki, A., Kaneyama, M., Shibamura, K., Takeda, Y., Abe, J., & Hizukuri, S. (1993). Physicochemical Properties of Japanese Arrowhead (*Sagittaria trifolia* L. var *sinensis* Makino) Starch. *Denpun Kagaku*, 40(1), 41-48. <https://doi.org/10.5458/jag1972.40.41>
- Thai Industrial Standards Institute. (2011). *Community product standard of cracker, CMU. 107/2554*. Ministry of Industry. Bangkok. (in Thai)
- Zhang, X., Zhang, H., & Na, Y. (2015). Transformation of Sodium during the Ashing of Zhundong Coal. *Procedia Engineering*, 102, 305 - 314. <https://doi.org/10.1016/j.proeng.2015.01.147>

Research Article

การศึกษาสมบัติและการบำบัดสันเท้าแตกของแผ่นโฟมรองสันเท้าจากน้ำยางธรรมชาติที่มียูเรียเป็นสารตัวเติม

Study on properties and cracked heels treatment of insole foam made from urea filled natural rubber latex

ฮาซัน ดอโป^{1*}, ซิตีไซยิดาห์ สายวารี², สุกรี เส็มหมาด³ และอัจมาน อาแด⁴

Hasan Daupor^{1*}, Sitisaiyidah Saiwari², Sukree Semmard³ and Ajaman Adair⁴

¹สาขาวิชาเคมี คณะวิทยาศาสตร์เทคโนโลยีและการเกษตร มหาวิทยาลัยราชภัฏยะลา ประเทศไทย

²ภาควิชาเทคโนโลยียางและพอลิเมอร์ คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยสงขลานครินทร์ ประเทศไทย

³โรงพยาบาลศูนย์ยะลา ประเทศไทย

⁴สาขาวิชาวิทยาศาสตร์เครื่องสำอางและความงาม คณะวิทยาศาสตร์เทคโนโลยีและการเกษตร มหาวิทยาลัยราชภัฏยะลา ประเทศไทย

¹Chemistry Program, Faculty of Science Technology and Agriculture, Yala Rajabhat University, Yala 95000, Thailand

²Department of Rubber Technology and Polymer Science, Faculty of Science and Technology, Prince of Songkla University, Pattani campus, Pattani 94000, Thailand

³Yala Reginal Hospital, Thailand

⁴Cosmetic Science and Beauty Program, Faculty of Science Technology and Agriculture, Yala Rajabhat University, Yala 95000, Thailand

*E-mail: hasan.d@yru.ac.th

Received: 18/04/2021; Revised: 31/07/2021; Accepted: 17/08/2021

บทคัดย่อ

งานวิจัยนี้ศึกษาการเตรียมฟองน้ำจากยางธรรมชาติผสมสารยูเรียเกรดเครื่องสำอางนำไปขึ้นรูปเป็นผลิตภัณฑ์ต้นแบบแผ่นรองสันเท้า และศึกษาสมบัติในการบำบัดอาการสันเท้าแตกกับกลุ่มอาสาสมัครจำนวน 35 คน สำหรับสูตรและสภาวะที่ใช้ในการเตรียมฟองน้ำสำหรับแผ่นรองสันเท้าในงานวิจัยชิ้นนี้ คือ ใช้น้ำยางคอมพาวนด์ผสมแคลเซียมคาร์บอเนต 30 phr ยูเรีย 12 phr และ DPG 1.2 phr ใช้เวลาในการตีฟองน้ำ 15 นาทีที่อุณหภูมิห้อง ให้ค่าความหนาแน่นและค่าความดันที่ใช้ในการกดฟองน้ำให้ยุบตัว 25% เท่ากับ $0.80 \pm 0.002 \text{ g/cm}^3$ และ $1.71 \pm 0.20 \text{ kPa}$ ตามลำดับ นอกจากนี้ พบว่าการเติมยูเรีย ค่าความหนาแน่นจะลดลง เนื่องจากมีแก๊สคาร์บอนไดออกไซด์เกิดขึ้นมากวิเคราะห์หาปริมาณยูเรียด้วยเทคนิค Attenuated total reflectance-Fourier transform infrared (ATR-FTIR) โดยการ

สร้างกราฟมาตรฐานของยูเรียในโหมด one peak และ one base พบว่าปริมาณยูเรียที่คงอยู่ในผลิตภัณฑ์หลังจากผ่านการใช้งานจริงกับกลุ่มอาสาสมัคร 1 สัปดาห์จนถึง 4 สัปดาห์ ปริมาณยูเรียจะค่อย ๆ ลดลง คือ 133.41 ± 11.87 , 77.23 ± 24.86 , 59.21 ± 20.56 และ 26.07 ± 8.99 มิลลิกรัม ตามลำดับ จากการทดลองใช้ผลิตภัณฑ์โฟมรองส้นเท้ากับกลุ่มอาสาสมัครบุคลากรทางการแพทย์และสาธารณสุข โรงพยาบาลศูนย์ยะลา ที่มีปัญหาส้นเท้าแตกจำนวน 35 คน ผลการวิเคราะห์ทางสถิติประกอบกับภาพถ่าย พบว่าสามารถบำบัดอาการส้นเท้าแตกหลังจากการใช้ 1 สัปดาห์ และผิวหนังบริเวณส้นเท้าแตกจะดีขึ้นเรื่อย ๆ ในสัปดาห์ถัดไปเมื่อเทียบกับกลุ่มอาสาสมัครอ้างอิงที่ใช้แผ่นโพนที่ไม่มีสารยูเรีย โดยสอดคล้องกับผลการวิเคราะห์ทางสถิติด้วย t-test แบบ 2 กลุ่มตัวอย่างที่สัมพันธ์กัน ที่ระดับนัยสำคัญ 0.05 ในด้านความพึงพอใจโดยรวมต่อผิวหนังบริเวณส้นเท้าก่อนใช้และหลังใช้ 4 สัปดาห์ พบว่าก่อนใช้และหลังใช้ 1 สัปดาห์ไม่แตกต่างกัน ($t\text{-prob}=0.173$) แต่หลังจากการใช้ 2, 3 และ 4 สัปดาห์ มีความแตกต่างอย่างมีนัยสำคัญ

คำสำคัญ: ยูเรีย, น้ำยางธรรมชาติ, โพนยาง, แผ่นรองส้นเท้า, ส้นเท้าแตก

Abstract

This research studied preparation of natural rubber foam adding urea cosmetic grade for producing heel pads prototype and further investigation a foot care treatment with 35 volunteers. It was found that the optimum latex compound formulation consisted of CaCO_3 30 phr, urea 12 phr and DPG 1.2 phr. They were blended for 15 minutes at room temperature. Density and 25% compressive stress of foam were $0.80 \pm 0.002 \text{ g/cm}^3$ and $1.71 \pm 0.20 \text{ kPa}$, respectively. Moreover, an increase of urea led to decreasing trends of both density and stress of the prototype products which is due mainly to a more releasing of carbon dioxide gas. The ATR-FTIR technique was used to determine the quantity of urea by generating calibration curve in one peak and one base modes. The prototype products were used by the volunteer for 1- 4 weeks. It was found that the urea contents slightly decreased from $161.30 \pm 16.89 \text{ mg}$ to 133.41 ± 11.87 , 77.23 ± 24.86 , 59.21 ± 20.56 and $26.07 \pm 8.99 \text{ mg}$ 1 to 4 weeks usages, respectively. The practical satisfaction tests were done with 35 volunteers who have suffered from heel broken. The statistical analysis results combined with photos was found that the cracked heels were able to be healed after having the treatment for 1 week and the skin have been better in the next week compared to that of using the controlled samples without urea. This result was confirmed by the statistical analysis, 2 type relation t-test, at a level of significance 0.05 in overall satisfaction before and after using for 4 weeks. The results found that there was no difference ($t\text{-prob}=0.173$) after 1 week of treatment but there were significant differences after treating for 2, 3 and 4 weeks.

Keywords: Urea, Natural rubber latex, Rubber foam, Insole, Heel cracked

จากการศึกษาอาการส้นเท้าแตกพบว่ามักเกิดในเพศหญิงมากกว่าเพศชาย เมื่ออายุมากขึ้น อาการส้นเท้าแตกมีระดับความรุนแรงถึงขั้นอักเสบ ปวดแสบ ร้อนแดง บางรายเจ็บจนไม่สามารถเหยียบบนพื้นได้ สาเหตุของส้นเท้าแตกเกิดจากการที่ไขมันทั้งที่เป็น Skin fat และ Sebum ที่เชื่อมระหว่างเซลล์ผิวหนังหรือระหว่างเยื่อผิว โดยกลไกของผิวหนังที่เรียกว่า Rein's barrier เสื่อมสภาพไป เนื่องจากแรงกดหรือแรงขัดถูจากการใช้เท้ารับน้ำหนัก ร่างกาย ส่งผลให้ไขมันที่ทำหน้าที่กักเก็บความชุ่มชื้นแก่ผิวหนังบริเวณนี้เสื่อมสภาพไป ทำให้น้ำในผิวหนังชั้นนอกที่ปกคลุมเท้า (stratum corneum) ซึ่งเป็นหนังกำพร้าชั้นบนสุดระเหยออกไป จึงไม่สามารถรักษาความชุ่มชื้นให้แก่ผิวหนังได้ตามธรรมชาติได้ (Kingkaew, 2009) ปัจจัยเสริมที่ทำให้ส้นเท้าแตก ได้แก่ น้ำหนักตัวที่กดทับบริเวณส้นเท้า ที่เสริมให้ผิวหนังที่แห้งอยู่แล้วปริแตก สภาพผิว อายุ รวมทั้งการสวมใส่รองเท้าที่มีขนาดไม่พอดีกับขนาดของเท้า ลักษณะของรองเท้าที่ใช้พื้นแข็งเกินไป ประกอบกับสภาพอากาศของประเทศไทยมีลักษณะร้อนชื้น จึงส่งผลให้ผิวหนังเกิดการแตกได้ง่ายยิ่งขึ้น การนํายางพารา มาทำเป็นพื้นรองเท้า จะได้พื้นรองเท้าที่มีความนุ่มและยืดหยุ่น เป็นการลดปัจจัยที่จะส่งเสริมให้ส้นเท้าแตกได้ การรักษาอาการส้นเท้าแตก มีวิธีการหลายวิธี ได้แก่ วิธีการตามธรรมชาติและ การใช้ยารักษา ดังนี้ (Wolverton, 2012) 1) สารเพคตินจากเปลือกกล้วย 2) การใช้ขี้เถ้ามะละกอ 3) น้ำมันมะกอก น้ำมันงา ขี้ผึ้ง หรือ วาสลีน 4) แวกซ์เท้าด้วยพาราฟิน 5) นํามะนาวผสมดินสอพอง 6) สารส้ม เป็นต้น

- 3 -

สารรักษาเส้นเท้าแตกรูปของครีม จะมียูเรียเป็นองค์ประกอบหลักอยู่ 10-45 เปอร์เซ็นต์ ยูเรียจะทำหน้าที่กักเก็บความชุ่มชื้นให้แก่ผิวและช่วยเร่งทำให้เกิดการผลิตของเซลล์ผิวเก่า (Jeffrey, 1997) ส่งผลให้ร่างกายสร้างเซลล์ผิวใหม่ขึ้นมาทดแทน

ดังนั้น การใช้ยูเรียเป็นองค์ประกอบในแผ่นรองรองเท้าจึงมีความเป็นไปได้ที่อนุภาคยูเรียจะซึมเข้าไปยังผิวหนังบริเวณส้นเท้าที่มีรอยแตกในขณะสวมใส่รองเท้า จนเกิดการรักษาสภาพผิวของส้นเท้าแตกนั้นให้กลับมาเรียบเนียนเหมือนเดิมได้ ผู้วิจัยจึงสนใจที่จะนำยูเรียเกรดที่ใช้ในเครื่องสำอางมาผสมกับน้ำยางคอมพาวนด์และขึ้นรูปเป็นฟองน้ำ สำหรับรองเฉพาะบริเวณส้นเท้า โดยมีการศึกษาคุณลักษณะทางกายภาพและทางเคมีของแผ่นยางคอมพาวนด์หลังขึ้นรูป และนำแผ่นรองเท้าที่ได้ไปทดลองใช้กับกลุ่มอาสาสมัครที่มีอาการส้นเท้าแตก โดยสอดเข้าไปด้านในของส้นรองเท้า และศึกษาติดตามประสิทธิภาพในการบำบัดอาการส้นเท้าแตกหลังจากสวมใส่ในระยะเวลาหนึ่ง คาดว่าผลิตภัณฑ์ลักษณะดังกล่าวนี้จะเพิ่มทางเลือกและเพิ่มโอกาสในการรักษาอาการส้นเท้าแตกให้มากยิ่งขึ้น

วัสดุอุปกรณ์และวิธีการทดลอง

สูตรน้ำยางคอมพาวนด์และการเตรียมโฟม

สูตรน้ำยางคอมพาวนด์สำหรับงานวิจัยนี้ แสดงในตารางที่ 1 โดยมีวิธีการผสมและขึ้นรูป ดังนี้ นำน้ำยางคอมพาวนด์ มาปั่นให้เกิดฟอง เป็นเวลา 5 นาที เติม 10% TSPP 1.0 phr กวนเป็นเวลา 2 นาที เติม 50% Urea 12 phr กวนเป็นเวลา 2 นาที เติม 20% DPG 1.2 phr กวนเป็นเวลา 2 นาที เติม 50% ZnO 2.0 phr กวนเป็นเวลา 2 นาที เติม 20% SSF 1.2 phr กวนเป็นเวลา 1 นาที จากนั้นเทโฟมยางลงในแม่พิมพ์ ปล่อยให้โฟมยางเกิดการเจลสมบูรณ์แล้ว นำโฟมยางไปนึ่งไอน้ำเป็นเวลา 60-90 นาที ตามด้วยนำโฟมยางออกจากแม่พิมพ์ล้างสารเคมีที่เหลือจากปฏิกิริยาออกให้สะอาด สุดท้ายนำโฟมยางอบในตู้อบ ที่อุณหภูมิ 70 °C จนแห้ง

ตารางที่ 1 สูตรน้ำยางคอมพาวนด์

สารเคมี	น้ำหนักยางแห้ง (phr)
60% HA natural rubber latex	100.0
20% Potassium oleate	2.0
50% Sulfur dispersion	2.5
50% ZMBT dispersion	1.0
50% ZDEC dispersion	1.0
50% Wingstay-L dispersion	1.0
50% CaCO ₃ dispersion	30

การวิเคราะห์หาปริมาณการคงอยู่ของยูเรียในแผ่นโพร่งสันเท้า

ให้ตัวแทนอาสาสมัครจำนวน 5 คน สำหรับทดสอบการคงอยู่ของยูเรีย โดยการคัดเลือกตัวแทนอาสาสมัครแบบจำเพาะเจาะจง (purposive sampling) จากกลุ่มที่ทดลองใช้แผ่นโพร่งสันเท้าที่มียูเรีย 30 คน โดยกำหนดคุณสมบัติตามช่วงน้ำหนัก ดังนี้ 1) กลุ่มที่มีอายุ 40-50 กิโลกรัม จำนวน 1 คน 2) กลุ่มที่มีอายุ 51-60 กิโลกรัม จำนวน 1 คน 3) กลุ่มที่มีอายุ 61-70 กิโลกรัม จำนวน 1 คน 4) กลุ่มที่มีอายุ 71 กิโลกรัมขึ้นไป จำนวน 2 คน โดยให้ตัวแทนอาสาสมัครสวมใส่แผ่นโพร่งสันเท้าจนครบ 1 สัปดาห์ จะนำแผ่นโพร่งสันเท้ากลับมาวิเคราะห์ที่ห้องปฏิบัติการ โดยวางทั้งแผ่นบนแผ่นรองรับตัวอย่างบนเครื่องอินฟราเรดสเปกโตรสโกปี ทำการวัด 3 จุด ด้วยจำนวนสแกน 16 ครั้ง/นาที่บนพื้นผิว จะนำแผ่นโพร่งสันเท้า หลังจากวิเคราะห์เสร็จสิ้น คืนแผ่นโพร่งสันเท้าให้ตัวแทนอาสาสมัครนำกลับไปใช้ต่อ เพื่อจะใช้วิเคราะห์ในสัปดาห์ที่ 2 ทำเช่นนี้ จนครบ 4 สัปดาห์ ทำการวิเคราะห์ในโหมด 1 Peak ที่ตำแหน่ง 3332 cm^{-1} และ Peak method ได้แก่ 1 Base (Height) ที่ตำแหน่ง 3380 cm^{-1} อาศัยความสูงของสเปกตรัมจาก 2 ตำแหน่งนี้ คำนวณความเข้มข้นของยูเรียด้วยกราฟมาตรฐานของยูเรีย

การสร้างกราฟมาตรฐานยูเรียด้วยเครื่องอินฟราเรดสเปกโตรสโกปี

ผสมสารมาตรฐานยูเรียกับโพแทสเซียมโบรไมด์ในอัตราส่วน 1: 5 อัดเป็นแผ่นบางหนัก 30 มิลลิกรัม ที่ความเข้มข้น 20 40 60 80 120 160 และ 200 มิลลิกรัม ทำการทดลอง 3 ซ้ำ นำไปวัดค่าการดูดกลืนแสงอินฟราเรดที่ย่านความถี่ $4000\text{--}500\text{ cm}^{-1}$

ประชากรกลุ่มตัวอย่างทดลองใช้ผลิตภัณฑ์

ทดลองใช้ผลิตภัณฑ์กับกลุ่มอาสาสมัครบุคลากรทางการแพทย์และสาธารณสุข ในโรงพยาบาลศูนย์ยะลา จากแผนกต่าง ๆ ได้แก่ แผนกห้องคลอด แผนกห้องบัตรและคัดกรอง แผนกไอซียู แผนกฉุกเฉิน ห้องฟอกไต แผนกกายภาพบำบัด แผนก MRI และแผนกศัลยกรรมชาย รวมทั้งหมด 8 แผนกงานจำนวน 35 คน แบ่งอาสาสมัครออกเป็น 2 กลุ่ม กลุ่มแรกใช้แผ่นโพร่งสันเท้าที่มียูเรียจำนวน 30 คน ที่เหลืออีก 5 คน ทดลองใช้แผ่นโพร่งสันเท้าที่ไม่มียูเรีย

วิธีการทดลองใช้ผลิตภัณฑ์

- ในขั้นแรกให้อาสาสมัครทำความสะอาดเท้าทั้ง 2 ข้างด้วยสบู่อ่อน แต่ถ้าบริเวณสันเท้าสกปรกมากก็สามารถใช้แปรงสีฟันเก่าขัดถูเบา ๆ จนสะอาด แล้วจึงล้างด้วยน้ำสะอาด เช็ดเท้าให้แห้ง
- ทำการถ่ายภาพบริเวณสันเท้าทั้ง 2 ข้างเอาไว้เพื่อใช้เปรียบเทียบและประเมินผลการเปลี่ยนแปลงของสันเท้าหลังจากใช้ผลิตภัณฑ์ หลังจากนั้นให้อาสาสมัครกรอกแบบสอบถามประเมินความพึงพอใจต่อผิวของตนเองก่อนใช้ผลิตภัณฑ์
- สวมใส่ผลิตภัณฑ์โดยสอดไว้ใต้สันเท้าภายในรองเท้าหรือสอดเข้าไปในถุงเท้าและสวมใส่รองเท้าตามปกติ ดังภาพที่ 1 ตั้งแต่เช้าจนถึงเย็น สัปดาห์ละ 5 วัน ใช้เดินไปตามปกติของการใช้ชีวิตประจำวัน หลังเลิกงานสามารถถอดออกได้

- หลังจากใช้งานผลิตภัณฑ์ในแต่ละวันสามารถดึงออกมาฝึ่งให้อากาศถ่ายเทได้ หลังจากใช้ผลิตภัณฑ์ครบ 1 สัปดาห์ จะถ่ายรูปเพื่อเปรียบเทียบการเปลี่ยนแปลงของผิวหนังบริเวณสันเท้า และประเมินความพึงพอใจต่อผลการบำบัดโดยให้อาสาสมัครกรอกแบบสอบถามด้วยตนเอง โดยมีระดับคะแนนดังนี้ 1.00 – 1.50 หมายถึง น้อยที่สุด 1.51 – 2.50 หมายถึง น้อย 2.51 – 3.50 หมายถึง ปานกลาง 3.51 – 4.50 หมายถึง มาก และ 4.15 – 5.00 หมายถึง มากที่สุด
- หลังจากใช้ผลิตภัณฑ์ครบ 2 3 4 และ 5 สัปดาห์ จะทำการเก็บภาพและแบบสอบถามเหมือนข้างต้น
- เก็บรวบรวมข้อมูลทั้งภาพถ่ายและแบบประเมิน เพื่อนำมาวิเคราะห์ผลทางสถิติด้วยโปรแกรม SPSS



รูปที่ 1 วิธีการสวมใส่ผลิตภัณฑ์แผ่นโฟมรองสันเท้า

เครื่องมือที่ใช้ในการวิจัย

- ภาพถ่ายก่อนและหลังการใช้ผลิตภัณฑ์ ที่บันทึกภาพด้วยกล้อง แล้วนำมาเปรียบเทียบการเปลี่ยนแปลงของผิวหนังก่อนและหลังใช้ผลิตภัณฑ์
- แบบประเมินความพึงพอใจของอาสาสมัคร ทั้งก่อนและหลังการใช้ผลิตภัณฑ์ โดยแต่ละคนจะได้ประเมินความพึงพอใจต่อผลการรักษาคณะ 5 ครั้ง คือ ครั้งที่ 1 ประเมินก่อนใช้ผลิตภัณฑ์ ครั้งที่ 2 3 4 และ 5 ประเมินหลังจากเริ่มใช้ผลิตภัณฑ์แล้วในสัปดาห์ที่ 2 3 4 และสัปดาห์ที่ 5 ในระยะเวลา 1 เดือน หลังจากนั้นผู้วิจัยจะนำแบบประเมินมาวิเคราะห์ทางสถิติ ซึ่งแบบประเมินประกอบด้วย 4 ส่วน คือ 1) เป็นข้อมูลส่วนตัวทั่วไป ประวัติการแพ้ยาและสารเคมีต่าง ๆ 2) เป็นแบบประเมินความพึงพอใจต่อผลการรักษา 3) เป็นแบบประเมินระดับความรุนแรงของสันเท้าแตก โดยใช้ Rating scale ที่ผู้วิจัยได้กำหนดไว้เป็นเครื่องมือในการวัดค่าคะแนนการเปลี่ยนแปลง โดยระดับความรุนแรงของสันเท้าแตก คัดแปลงจาก Classification of striae based on clinical appearance (Adatto & Deprez, 2003) และวิธีประเมินทางผิวหนัง (Pornnida, 2008) ซึ่งเป็นความคิดเห็นและข้อเสนอแนะต่อวิธีการใช้และผลิตภัณฑ์

สถิติที่ใช้ในการวิเคราะห์ข้อมูล

- ข้อมูลทั่วไปและข้อมูลเชิงบรรยาย ในส่วนที่ 1 ใช้สถิติเชิงพรรณนา (Descriptive statistics) โดยใช้การแจกแจงความถี่ ร้อยละ ค่าเฉลี่ย
- การวิเคราะห์ข้อมูลเพื่อเปรียบเทียบประสิทธิภาพและความพึงพอใจต่อผลการรักษาด้วยวิธีการเดียวกัน (ทำข้างเดียวกัน) ใช้สถิติ Pair sample T-test

ผลและวิจารณ์ผลการทดลอง

สูตรที่เหมาะสมในการเตรียมแผ่นโฟมรองสันเท้า

สูตรที่เหมาะสมในการเตรียมแผ่นโฟมรองสันเท้าจากน้ำยางคอมพาวนด์ที่มียูเรียเป็นสารตัวเติม คือ การใช้ยูเรียปริมาณ 12 phr กับแคลเซียมคาร์บอเนต 30 phr และส่วนผสมอื่น ๆ ดังตารางที่ 1 ข้างต้น จะได้ตัวอย่างผลิตภัณฑ์แผ่นโฟมรองสันเท้า ดังภาพที่ 2



(ก)



(ข)

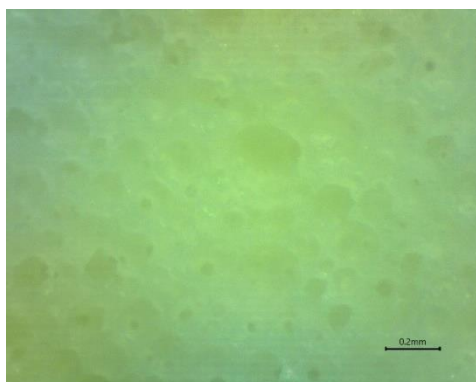
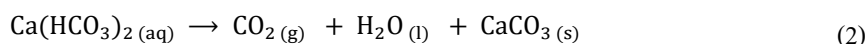
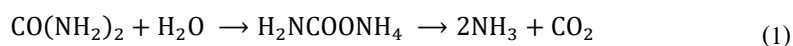
รูปที่ 2 ตัวอย่างแผ่นโฟมรองสันเท้า (ก) แผ่นโฟมรองสันเท้าบรรจุถุงเพื่อแจกให้อาสาสมัครทดลองใช้ (ข)

ความหนาแน่นของโฟม

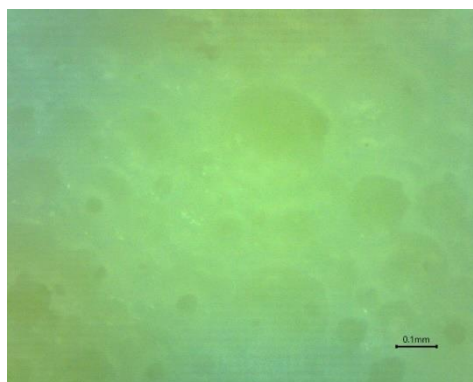
จากผลการทดสอบความหนาแน่นของโฟมที่ไม่เติมยูเรียและเติมยูเรีย 12 phr มาตรฐาน ASTM D3574-95 พบว่ามีค่าความหนาแน่นเท่ากับ 0.10 ± 0.001 และ 0.08 ± 0.002 g/cm³ ตามลำดับ ความหนาแน่นลดลงตามปริมาณการเพิ่มของยูเรีย ทำให้โฟมมีความยืดหยุ่นและความหนาแน่นลดลงกว่าโฟมที่ไม่มียูเรีย อัตราเร็วในการบั่นให้เกิดฟองใช้อัตราเร็วคงที่ตลอด เพื่อควบคุมฟองน้ำให้มีความหนาแน่นพอดี แต่ถ้าบั่นโดยใช้ระดับความเร็วรอบที่สูงกับเวลานาน จะทำให้ฟองยางมีความหนาแน่นต่ำ ในขณะที่การตีฟองด้วยความเร็วรอบต่ำ ๆ กับเวลานานขึ้น จะทำให้ฟองยางมีความหนาแน่นสูงขึ้นได้

ความเค้นที่ใช้ในการกดฟองน้ำให้ยุบตัว 25%

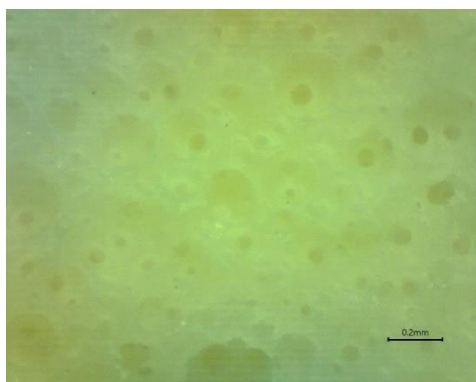
จากผลการทดสอบความสามารถในการรับแรงของฟองน้ำให้ยุบตัว 25% ตามมาตรฐาน มอก. 173-2529 พบว่าโฟมที่ไม่มียูเรียและมียูเรีย 12 phr มีค่าความเค้นเท่ากับ 1.70 ± 0.16 และ 1.71 ± 0.20 KPa แสดงให้เห็นว่าโฟมที่มียูเรียจะมีการยุบตัวที่ใกล้เคียงกับโฟมที่ไม่มียูเรีย เนื่องจากหมู่เอมีนในโครงสร้างของยูเรียทำให้เกิดก๊าซคาร์บอนไดออกไซด์ ดังสมการที่ 1 และคาร์บอนไดออกไซด์สามารถทำปฏิกิริยากับแคลเซียมคาร์บอเนตเกิดเป็นแคลเซียมไบคาร์บอเนต ซึ่งสามารถให้ก๊าซคาร์บอนไดออกไซด์อีกเช่นกัน ดังสมการที่ 2 (Patnaik, 2002) จึงทำให้เซลล์โฟมมีรูพรุนเกิดขึ้นจำนวนมาก เป็นลักษณะของโฟมแบบ open-cell ดังภาพที่ 3ก-ข ส่วนโฟมที่ไม่มียูเรียเซลล์จะมีรูพรุนอยู่บางบริเวณ แต่ส่วนใหญ่จะไม่มีรูพรุน จึงจัดเป็นโฟมแบบ semi-open cell ดังภาพที่ 3ค-ง



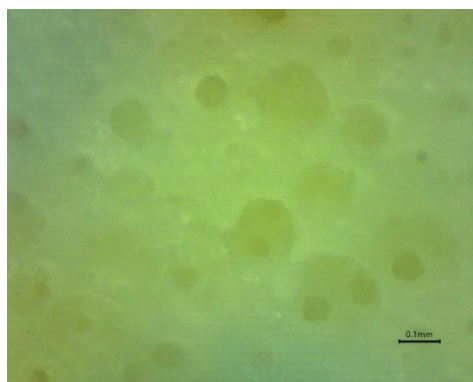
(ก)



(ข)



(ค)

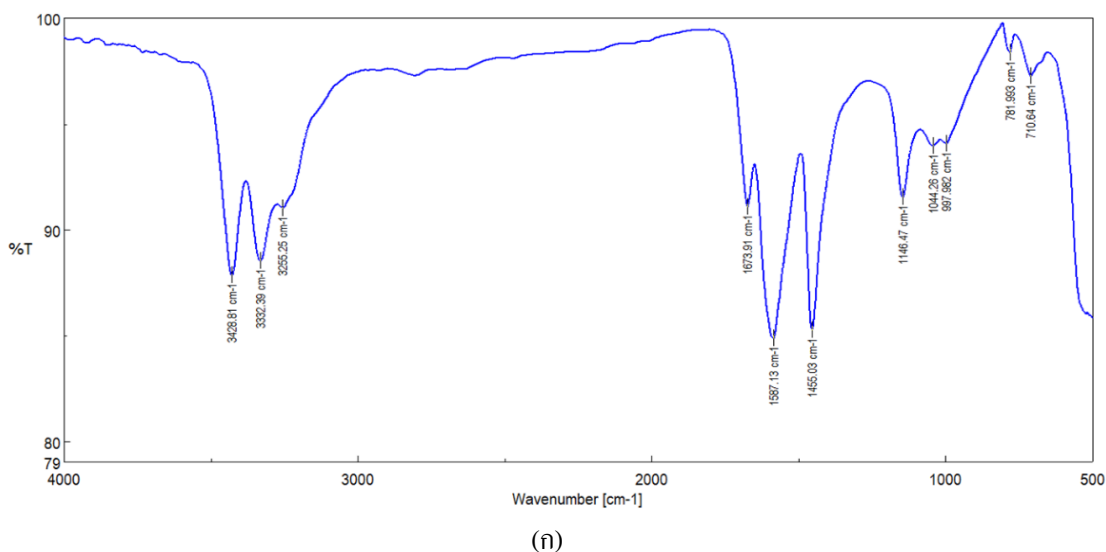


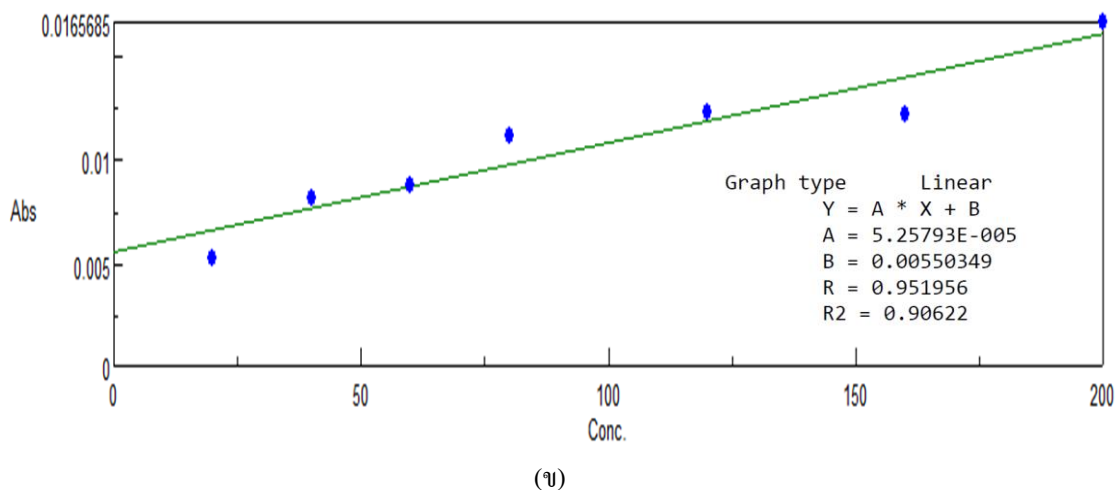
(ง)

รูปที่ 3 ภาพถ่ายลักษณะพื้นผิวโพลีด้วยกล้องจุลทรรศน์ Stereo (ก) โพลีที่ไม่มียูเรียกำลังขยาย 20 เท่า (ข) โพลีที่ไม่มียูเรียกำลังขยาย 30 เท่า (ค) โพลีที่มียูเรียกำลังขยาย 20 เท่า (ง) โพลีที่มียูเรียกำลังขยาย 30 เท่า

สเปกตรัมการดูดกลืนอินฟราเรดและกราฟมาตรฐานของสารยูเรีย

สารมาตรฐานยูเรียจะปรากฏสเปกตรัมการดูดกลืนแสงของหมู่ฟังก์ชันเอมีน ($-NH_2$) ที่เด่นชัดที่ตำแหน่ง 3428 และ 3332 cm^{-1} ดังภาพที่ 4ก เป็นการสั่นแบบ out-of-plane stretching และ in-plane stretching ตามลำดับ (Grdadolnik and Marechalb, 2002) เมื่อทำการขึ้นชั้นการบดและการคงอยู่ของสารยูเรียในผลิตภัณฑ์พบว่า สารยูเรียหลังจากบดออกมาจากผลิตภัณฑ์และที่ยังคงอยู่ในผลิตภัณฑ์ ตำแหน่งการดูดกลืนแสงของสารยูเรียไม่ปรากฏที่ตำแหน่ง 3444 cm^{-1} แต่ที่ตำแหน่ง 3332 cm^{-1} จะเด่นชัดขึ้น คาดว่าเป็นอันตรกิริยาของหมู่คาร์บอนิลในสารยูเรียไปเกิดพันธะไฮโดรเจนกับองค์ประกอบอื่นในแผ่นโพลีเอทิลีนหรือเกิดพันธะไฮโดรเจนกับความชื้น (Jeffrey, 1997) ดังนั้น ในการสร้างกราฟมาตรฐานยูเรียและการวิเคราะห์ปริมาณสารยูเรีย จึงเลือกความยาวคลื่นที่ตำแหน่ง 3332 cm^{-1} และมีความเข้มข้นของยูเรียตั้งแต่ 20, 40, 60, 80, 120, 160 และ 200 มิลลิกรัม ได้สมการเส้นตรงเป็น $y = 0.0000549x + 0.00453$ โดยมีค่าความเป็นเชิงเส้น (R^2) เท่ากับ 0.9062 ดังภาพที่ 3ข จะใช้กราฟมาตรฐานนี้สำหรับการวิเคราะห์หาปริมาณยูเรียในสารตัวอย่างอื่น ๆ ต่อไป

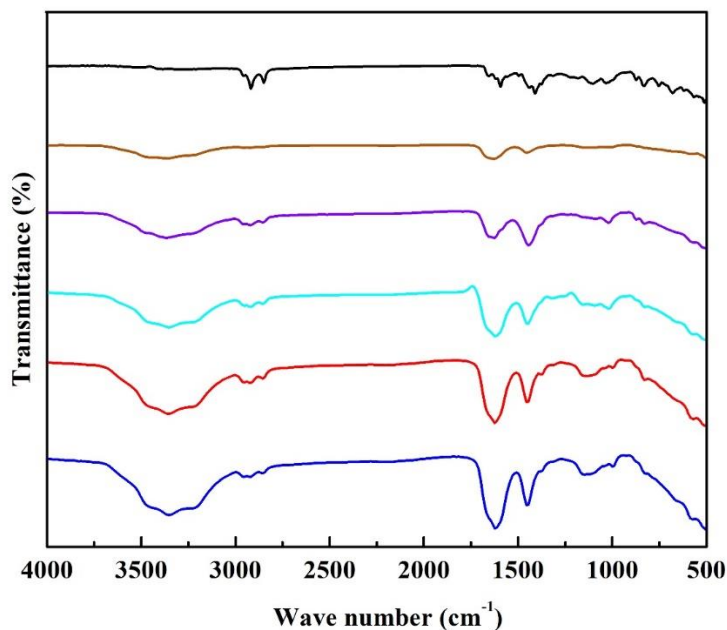




รูปที่ 4 สเปกตรัมการดูดกลืนแสงอินฟราเรดของสารมาตรฐานยูเรีย (ก) และกราฟมาตรฐานของสารยูเรียที่ความเข้มข้นต่าง ๆ (ข)

ปริมาณการคงอยู่ของยูเรียในผลิตภัณฑ์ (Shelf life)

ผลิตภัณฑ์แผ่นโฟมรองสันเท้าเมื่อผ่านการใช้งาน สารยูเรียที่ค่อย ๆ บวมออกมา จะทำให้ปริมาณสารยูเรียปริมาณ 161.3038 มิลลิกรัม (กำหนดให้เทียบเท่า 12 phr หรือเทียบเท่า 100 เปอร์เซ็นต์) ที่คงอยู่ในผลิตภัณฑ์ค่อย ๆ ลดลง หลังจากการใช้งานผ่านไปสัปดาห์แรก เมื่อนำไปวิเคราะห์หาปริมาณยูเรียพบว่ายังคงมีสารยูเรียอยู่ประมาณ 82.71 เปอร์เซ็นต์ เมื่อมีการใช้งานไปเป็นเวลา 2 3 และ 4 สัปดาห์ พบว่าปริมาณของสารยูเรียคงเหลือในผลิตภัณฑ์ประมาณ 47.88 39.71 และ 16.16 เปอร์เซ็นต์ ตามลำดับ ดังตารางที่ 2 ที่ความเข้มของสเปกตรัมอินฟราเรดค่อย ๆ ลดลง ดังภาพที่ 5 ขณะเดียวกันหลังจากผ่านการใช้งาน 4 สัปดาห์ แผ่นโฟมมีรอยละของการหดตัวเฉลี่ยเพียง 5 เปอร์เซ็นต์ ทำให้สามารถใช้งานต่อได้ แต่ปริมาณยูเรียคงเหลือจะน้อยลง



รูปที่ 5 ตัวอย่างสเปกตรัมอินฟราเรดของสารยูเรียที่คงอยู่ในแผ่นรองสันเท้า

หมายเหตุ โดยสเปกตรัมสีดำ (-) คือ แผ่นโฟมที่ไม่เติมยูเรีย สเปกตรัมสีน้ำเงิน (-) ก่อนใช้ผลิตภัณฑ์ สเปกตรัมสีแดง (-) คือยูเรียที่ตรวจพบหลังใช้ 1 สัปดาห์ สเปกตรัมสีฟ้า (-) คือ 2 สัปดาห์ สเปกตรัมสีม่วง (-) คือ 3 สัปดาห์ และสเปกตรัมสีน้ำตาล (-) คือ 4 สัปดาห์

ตารางที่ 2 ปริมาณการคงอยู่ของยูเรียในผลิตภัณฑ์แผ่นโฟมรองสันเท้า

สารตัวอย่าง	ปริมาณยูเรียคงเหลือที่วัดได้ (mg)	ปริมาณยูเรียคงเหลือเทียบเท่า (%)
ก่อนใช้ผลิตภัณฑ์	161.30 ± 16.89	100
หลังใช้ 1 สัปดาห์	133.41 ± 11.87	82.71
หลังใช้ 2 สัปดาห์	77.23 ± 24.86	47.88
หลังใช้ 3 สัปดาห์	59.21 ± 20.56	39.71
หลังใช้ 4 สัปดาห์	26.07 ± 8.99	16.16

ผลการเปรียบเทียบความพึงพอใจระหว่างกลุ่มอ้างอิงกับกลุ่มทดสอบในการใช้แผ่นโฟมรองสันเท้า

ตารางที่ 3 แสดงผลการศึกษาเปรียบเทียบความพึงพอใจระหว่างกลุ่มอ้างอิงกับกลุ่มทดสอบในการใช้แผ่นโฟมรองสันเท้า พบว่าความพึงพอใจโดยรวมและความพึงพอใจต่อการเปลี่ยนแปลงความชุ่มชื้นนุ่มนวลของผิวหนังสันเท้าระหว่างกลุ่มทดสอบกับกลุ่มอ้างอิง มีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ที่ระดับ 0.05 โดยพึงพอใจมากที่สุดในช่วงสัปดาห์ที่ 3 สำหรับความพึงพอใจโดยรวม (0.011*) และสัปดาห์ที่ 4 สำหรับการเปลี่ยนแปลงความชุ่มชื้นนุ่มนวลของผิวหนังสันเท้า (0.037*) และกลุ่มทดสอบมีแนวโน้มความพึงพอใจทั้งสองด้านเพิ่มขึ้นทุกสัปดาห์จนถึงสัปดาห์ที่ 4 ความพึงพอใจต่อการเปลี่ยนแปลงสภาพความลึกร่องรอยแตกสันเท้า สภาพความเรียบเนียนของสันเท้า สภาพความยืดหยุ่นของผิวหนังสันเท้า ระหว่างกลุ่มทดสอบกับกลุ่มอ้างอิง ไม่มีความแตกต่างกันทางสถิติ แต่มีค่าเฉลี่ยแตกต่างกันมากที่สุดในช่วงสัปดาห์ที่ 2 และมีแนวโน้มความพึงพอใจต่อการเปลี่ยนแปลงสภาพความยืดหยุ่นของผิวหนังสันเท้าเพิ่มขึ้นทุกสัปดาห์จนถึงสัปดาห์ที่ 4 ส่วนความพึงพอใจต่อระดับความรุนแรงของสันเท้าแตก ไม่มีความแตกต่างกันทางสถิติระหว่างกลุ่มทดสอบกับกลุ่มอ้างอิงและความพึงพอใจต่อการเปลี่ยนแปลงในแต่ละสัปดาห์ไม่แตกต่างกัน

ตารางที่ 3 แสดงค่าเฉลี่ยและค่าเบี่ยงเบนมาตรฐานและการทดสอบความแตกต่างระหว่างกลุ่มอ้างอิงกับกลุ่มทดสอบในการใช้แผ่นโฟมรองสันเท้า

ลำดับ	หัวข้อ	ค่าเฉลี่ย ก่อนใช้	ค่าเฉลี่ยหลังใช้		t-value	t-prob
			สัปดาห์	กลุ่มอ้างอิง	กลุ่มทดสอบ	
1	ความพึงพอใจโดยรวมต่อ ผิวหนังบริเวณสันเท้า	2.3143	1	2.200	2.6000	-1.756
			2	2.200	2.6286	-2.520
			3	2.200	2.8571	-3.026
			4	2.800	3.4000	-0.975
2	สภาพความลึกของร่องรอย แตก	3.2571	1	2.000	2.3714	-0.487
			2	2.000	2.6286	-2.514
			3	2.400	2.7714	-1.604
			4	3.000	3.2571	-0.165
3	ความเรียบเนียนของสันเท้า	3.0857	1	1.800	2.4286	-1.416
			2	2.000	2.5429	-2.127
			3	2.600	2.8857	-1.127
			4	3.000	3.2571	-0.628

4	ความยืดหยุ่นของผิวหนัง สันเท้า	2.4571	1	2.000	2.4857	-1.390	0.209
			2	2.000	2.6286	-2.054	0.083
			3	2.600	2.8000	-0.809	0.443
			4	2.600	3.0857	-1.842	0.097
5	ความชุ่มชื้นนุ่มนวลของ ผิวหนังสันเท้า	2.6857	1	2.000	2.4857	-1.523	0.169
			2	2.200	2.6000	-2.054	0.065
			3	3.000	3.0000	-0.360	0.722
			4	3.000	3.3429	-2.183	0.037*
6	ระดับความรุนแรงของสัน เท้าแตก	3.9143	1	3.8286	3.900	0.627	0.549
			2	3.8000	4.000	0.475	0.652
			3	3.8000	4.200	1.068	0.330
			4	3.4857	4.200	0.081	0.938

หมายเหตุ

- เครื่องหมาย * หมายถึง ความพึงพอใจที่แตกต่างกันอย่างมีนัยสำคัญทางสถิติ ที่ระดับ 0.05 ระหว่างกลุ่มทดสอบกับกลุ่มอ้างอิง

- ระดับคะแนนความพึงพอใจ กำหนดช่วงคะแนนดังนี้ 1.00 – 1.50 หมายถึง น้อยที่สุด 1.51 – 2.50 หมายถึง น้อย 2.51 – 3.50 หมายถึง ปานกลาง 3.51 – 4.50 หมายถึง มาก และ 4.15 – 5.00 หมายถึง มากที่สุด

ผลจากการสัมภาษณ์และการสังเกต

จากการที่ได้พูดคุยระหว่างลงพื้นที่เก็บข้อมูลจากอาสาสมัครบางราย เช่น รหัส 035 ที่มีความรุนแรงในระดับที่ 4 (คือ ระดับที่เห็นเป็นรอยแตกชัดเจน แต่ยังไม่มีเลือดออก) ที่ได้ใช้ผลิตภัณฑ์แผ่นโพนเป็นประจำ พบว่าเกิดการผลิตเซลล์ผิวหนังใน 1 สัปดาห์ จะสังเกตเห็นเป็นขุยขาว ๆ หลุดออกมาอย่างชัดเจน จนกระทั่งสัปดาห์ที่ 4 รอยแตกดีขึ้นกว่าเดิมมาก ผิวเรียบเนียนขึ้น และค่อย ๆ ดีขึ้นเป็นลำดับ ดังภาพที่ 6 ก-ค และในอาสาสมัครรหัส 028 มีความรุนแรงระดับที่ 4 เช่นกัน ก่อนใช้ผลิตภัณฑ์ผิวบริเวณสันเท้าเห็นเป็นรอยแตกชัดเจนและแห้งกร้าน หลังจากทดลองใช้ไปเพียง 1 สัปดาห์ ผิวส่วนที่แห้งกร้านจะหลุดลอกออกไป และค่อย ๆ ดีขึ้นเป็นลำดับในสัปดาห์ถัดไป ดังภาพที่ 6 ง-จ เมื่อได้เปรียบเทียบกับอาสาสมัครรหัส 014 ที่มีความรุนแรงในระดับ 4 เช่นกัน เป็นอาสาสมัครกลุ่มอ้างอิง ที่ได้ทดลองใช้แผ่นโพนรองสันเท้าสูตรที่ไม่มีสารยูเรีย ระยะเวลา 4 สัปดาห์เท่ากัน พบว่าไม่มีหลุดลอกของผิวที่แตก จึงไม่ทำให้ผิวเรียบเนียนขึ้นแต่อย่างใด โดยไม่มีการเปลี่ยนแปลงใด ๆ ต่อสันเท้าแตกในการทดลองใช้แผ่นโพนที่ไม่มีสารยูเรีย



(ก)



(ข)



(ค)



(ง)



(จ)



(ฉ)



(ช)



(ซ)

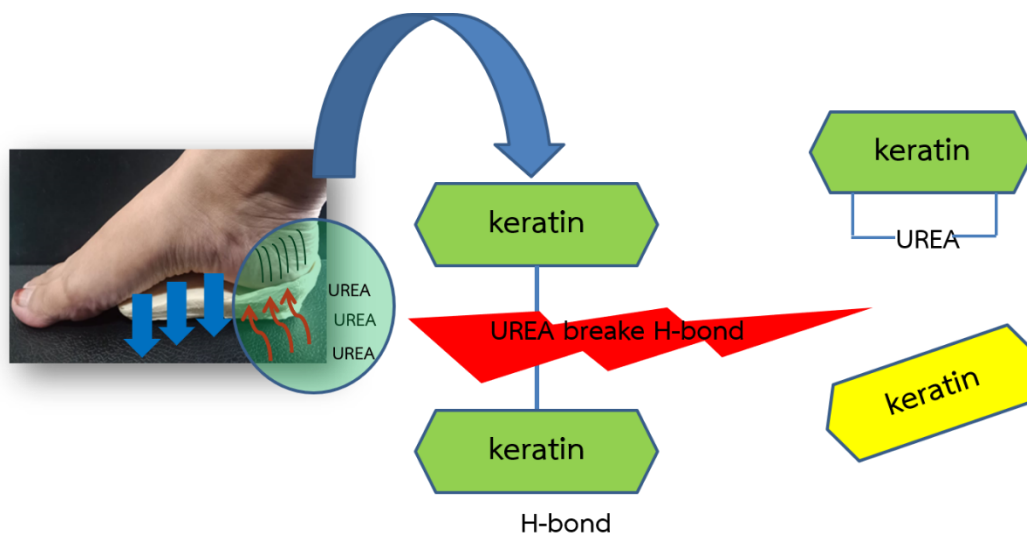


(ณ)

รูปที่ 6 ภาพก่อนใช้ หลังใช้ 1 สัปดาห์และหลังใช้ 4 สัปดาห์ ของอาสาสมัครรหัส 035 (ก-ค) อาสาสมัครรหัส 028 (ง-ฉ) และอาสาสมัคร 014 (ช-ณ)

กระบวนการทำงานของสารยูเรียที่ปล่อยออกมาสู่ผิวหนัง

ขณะที่มีการสวมใส่ จะมีการลงน้ำหนักกดทับแผ่นโพรตีสันเท้า จะทำให้ยูเรียที่อยู่ภายในแผ่นโพรตีสันที่รูพรุนซึมออกมาและเข้าเข้าสู่ผิวหนังผ่านรอยแตกของสันเท้า ยูเรียจะเข้าไปทำลายพันธะไฮโดรเจนที่ยึดเหนี่ยวแต่ละแผ่นของโปรตีนเคราติน ซึ่งเป็นองค์ประกอบหลักของหนังกำพร้า ที่กำลังเสื่อมสภาพ เมื่อยูเรียไปทำลายพันธะไฮโดรเจนให้หลุดไป และยูเรียกลับไปสร้างพันธะไฮโดรเจนเสียเอง ในแผ่นโปรตีนเคราตินอันเดียวกัน ไม่ได้สร้างพันธะระหว่างแผ่นเคราตินแผ่นอื่น ๆ ทำให้แผ่นเคราตินหลุดออก หรือเกิดการลอกของผิวหนังนั่นเอง (คัดแปลงจาก Wolverton, 2012; Rossky, 2008; Bennion & Daggett, 2003) ด้วยกลไกทางธรรมชาติของร่างกายมนุษย์จะสร้างผิวหนังใหม่ที่เรียบเนียนกว่าขึ้นมาทดแทน ดังภาพที่ 7



รูปที่ 7 กระบวนการปล่อยยูเรียและการกระตุ้นการผลัดเซลล์ผิวด้วยยูเรีย

สรุปผลการทดลอง

แผ่นโพรตีสันเท้าที่เตรียมด้วยสูตรดังกล่าวนี้ มีค่าความหนาแน่นเท่ากับ 0.08 g/cm^3 จะลดลงเมื่อเติมยูเรีย 12 phr เนื่องจากมีแก๊สคาร์บอนไดออกไซด์เกิดขึ้นเมื่อเติมยูเรีย ส่วนความเค้นในการกดฟองน้ำให้ยุบตัว 25 % เท่ากับ $1.71 \pm 0.20 \text{ kPa}$ พบว่ามีค่าการยุบตัวใกล้เคียงกับโพรตีสันที่ไม่เติมยูเรีย แต่สภาพของโพรตีสันที่มีค่าความหนาแน่น 0.08 g/cm^3 และความเค้นในการกดฟองน้ำเท่ากับ 1.70 kPa จะทำให้ยูเรียปล่อยออกมาได้ดี จนสามารถรักษาสันเท้าแตกได้ อีกทั้งแผ่นโพรตีสันยังคงรูปเดิมยังสามารถใช้งานได้ปกติหลังจากผ่านการใช้งานแล้ว 4 สัปดาห์ ปริมาณของตัวยายูเรียจากการวิเคราะห์เชิงปริมาณด้วยเครื่อง ATR-FTIR ยูเรียจะยังคงอยู่ในผลิตภัณฑ์หลังจากผ่านการทดลองใช้จริงสวมใส่ในรองเท้าโดยอาสาสมัครเป็นระยะเวลา 1 เดือน ถึงแม้ว่าปริมาณจะค่อย ๆ ลดน้อยลงก็ตาม สอดคล้องกับปริมาณยูเรียที่ตรวจพบในกระดาษซังสารที่ปล่อยออกมาจากชั้นทดสอบ การปล่อยยูเรียจากชั้นทดสอบหรือ

ผลิตภัณฑ์แผ่นโฟมรองส้นเท้า เป็นการจำลองที่จะแสดงให้เห็นว่ายูเรียสามารถซึมเข้าสู่รอยแตกของส้นเท้าได้ ผลการทดลองใช้กับอาสาสมัครที่มีปัญหาส้นเท้าแตกจำนวน 35 คน หลังจากได้ทดลองใช้เป็นเวลา 1 สัปดาห์ ส้นเท้าที่แตกมีการเปลี่ยนแปลง ผิวที่แตกเป็นรอยเกิดการผลัดของเซลล์ผิวออกมาเป็นขุยขาว ๆ เนื่องจากยูเรียที่สัมผัสกับรอยแตกเข้าไปเร่งให้เกิดการผลัดเซลล์ผิวโดยการทำให้พันธะไฮโดรเจนระหว่างแผ่นเคราตินหลุดออกจากกัน รอยแตกจึงดีขึ้น หลังจากสัปดาห์ที่ 2-4 ผิวเรียบและนุ่มกว่าเดิม ไม่แห้งกร้าน มีความชุ่มชื้น แสดงว่าสามารถบำบัดอาการส้นเท้าแตกได้ ถึงแม้ว่าตัวยูเรียที่คงอยู่ในผลิตภัณฑ์จะน้อยลงในสัปดาห์ที่ 4 แต่ผิวที่แตกเกิดการผลัดเซลล์ผิวแล้วในสัปดาห์ที่ 1 ดังนั้น คาดว่าไม่มีความจำเป็นที่ผลิตภัณฑ์แผ่นโฟมรองส้นเท้าจะต้องมีตัวยูเรียอยู่อีก เพราะยูเรียไปเร่งให้เกิดการผลัดของเซลล์ผิวและให้ความชุ่มชื้นแล้ว หลังจากนั้นคงเป็นกระบวนการทำงานของผิวหนังตามธรรมชาติของแต่ละคนที่จะสร้างผิวหนังใหม่ที่เรียบเนียนกว่าขึ้นมาทดแทน แผ่นโฟมรองส้นเท้าที่ตัวยูเรียมีปริมาณน้อยลงหรืออาจจะหมดไป เปรียบเสมือนพื้นรองเท้าที่มีความนุ่มที่สามารถป้องกันปัจจัยที่อาจจะทำให้ส้นเท้ากลับมาแตกซ้ำได้อีก

กิตติกรรมประกาศ

ขอขอบคุณสำนักงานคณะกรรมการส่งเสริมวิทยาศาสตร์ วิจัยและนวัตกรรม (สกสว.) กลุ่มมุ่งเป้าทางการแพทย์ ประจำปีงบประมาณ พ.ศ. 2561 ที่สนับสนุนทุนวิจัย

เอกสารอ้างอิง

- Adatto, M. A., & Deprez, P. (2003). Striae treated by a novel combination treatment-sand abrasion and a patent mixture containing 15% trichloroacetic acid followed by 6-24 hrs. of a patent cream under plastic occlusion. *Journal of Cosmetic Dermatology*, 2(2), 61–67. <https://doi.org/10.1111/j.1473-2130.2004.00023.x>
- Arthit, S., & Wiriya, T. (2010). *Materials development and design of heel cushion for reducing plantar heel pressure* [Unpublished master's thesis]. Prince of Songkla University. (in Thai)
- Bennion, B. J., & Daggett, V. (2003). The molecular basis for the chemical denaturation of proteins by urea. *Proceeding of the National Academy of Sciences of the United State of America*, 100(9), 5142-5147. <https://doi.org/10.1073/pnas.0930122100>
- Grdadolnka, J., & Marechalb, Y. (2002). Urea and urea-water solutions-an infrared study. *Journal of Molecular Structure*, 615, 177-189. [https://doi.org/10.1016/S0022-2860\(02\)00214-4](https://doi.org/10.1016/S0022-2860(02)00214-4)
- Jeffrey, G. A. (1997). *An introduction to hydrogen bonding (Topics in physical chemistry)*. New York, United States of America: Oxford University Press.

- Kingkaew, H. (2009). *Development of heel's crack pad for commercial purpose* [Unpublished master's thesis]. Chulalongkorn University. (in Thai)
- Patnaik, P. (2002). *Handbook of Inorganic Chemicals*. New York, United States of America: McGraw-Hill.
- Pornnida, S. (2008). *Study the increasing effect on heels cracked therapy by papaya latex, case study of customers at Aree Drug Store, Khon Kaen Province* [Unpublished master's thesis]. Khon Kaen University. (in Thai)
- Rossky, P. J. (2008). Protein denaturation by urea: slash and bond. *Proceeding of the National Academy of Sciences of the United State of America*, 105(44), 16825-16826. <https://doi.org/10.1073/pnas.0809224105>
- Wolverton, S. E. (2012). *Comprehensive Dermatologic Drug Therapy*. China: Elsevier.
- Yamsri, C., Pensri, P., Janwattanakul, P., Boonyong, S., Romsai, W., & Rattanapongbundit, N. (2013). Effect of kinesio taping combined with stretching on heel pain and foot functional ability in persons with plantar fasciitis: a preliminary study. *Chulalongkorn Medical Journal*, 57(1), 61-78. (in Thai)

Research Article

เทคนิคทางปัญญาประดิษฐ์สำหรับการตรวจข่าวปลอมภาษาไทย Artificial Intelligent Techniques for Thai Fake News Detection

พยุ่ง มีสัจ^{1*}, พนม คลีฉายา², อองอาจ อุ่นอนันต์³, กมลรัตน์ กิจรุ่งไพศาล⁴

Phayung Meesad^{1*}, Phnom Kleechaya², Aongart Aun-a-nan³, Kamonrat Kijrungsaisarn⁴

^{1,3}คณะเทคโนโลยีสารสนเทศและนวัตกรรมดิจิทัล มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ ประเทศไทย

^{1,3}Faculty of Information Technology and Digital Innovation, King Mongkut's University of Technology North Bangkok, Thailand

^{2,4}หน่วยปฏิบัติการวิจัยและพัฒนาทรัพยากรมนุษย์ด้านความรู้ทางดิจิทัลและการรู้เท่าทันสื่อ (DIRU) จุฬาลงกรณ์มหาวิทยาลัย
ประเทศไทย

^{2,4}Digital Intelligence and Literacy Research Unit (DIRU), Chulalongkorn University, Thailand

*E-mail: phayung.m@itd.kmutnb.ac.th

Received: 13/06/2021; Revised: 24/10/2021; Accepted: 23/04/2022

บทคัดย่อ

ด้วยเทคโนโลยีดิจิทัลที่ทันสมัยเพิ่มความสะดวกแก่ผู้ใช้งานในการเผยแพร่ข่าวสารบนสื่อออนไลน์ โดยส่วนหนึ่งได้เผยแพร่ข่าวปลอมที่มีจำนวนมากซึ่งเป็นปัญหาทำให้ผู้หลงเชื่อข่าวปลอมว่าเป็นข่าวจริงและนำไปแชร์ต่อ ทำให้เกิดปัญหาอื่นๆ ต่อเนื่องตามมา ปัจจุบันยังไม่มีเครื่องมือที่จะช่วยตรวจสอบข่าวปลอมภาษาไทย ที่จะช่วยชะลอการแพร่กระจายข่าว การตรวจจับข่าวปลอมแบบอัตโนมัติถือว่าเป็นงานที่ยากและมีความท้าทายสูง เนื่องจากข่าวมีความเคลื่อนไหวอย่างมาก งานวิจัยนี้เสนอวิธีการตรวจสอบข่าวปลอมภาษาไทยด้วยเทคนิคปัญญาประดิษฐ์ ผู้วิจัยใช้สามเทคนิคหลักในการสร้างแบบจำลองการตรวจจับข่าวปลอม ได้แก่ การสืบค้นข้อมูลสารสนเทศ การประมวลผลภาษาธรรมชาติ และการเรียนรู้ของเครื่อง การสืบค้นข้อมูลสารสนเทศเป็นกลไกในการดึงข้อมูลจากเว็บไซต์ข่าวออนไลน์และโซเชียลมีเดีย การประมวลผลภาษาธรรมชาติจะวิเคราะห์ข่าวที่ได้กลับคืนมาเพื่อสกัดข้อมูลคุณลักษณะมีความโดดเด่นเป็นอย่างดี การเรียนรู้ของเครื่องทำการรับข้อมูลฟีดแบ็กและจัดประเภทบทความข่าวออกเป็นสามประเภท ได้แก่ ข่าวจริง ข่าวปลอม และข่าวน่าสงสัย ในการเตรียมข้อมูลสอน ผู้วิจัยใช้เว็บครอว์เลอร์เพื่อรวบรวมจัดเก็บข้อมูลจำนวน 53,220 ตัวอย่าง โดยทำการจัดประเภทไว้ล่วงหน้าเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย เพื่อป้องกันความโน้มเอียง จำนวนตัวอย่างข้อมูลในแต่ละกลุ่ม สุ่มเลือกให้มีความสมดุล ผู้วิจัยใช้วิธีการสอนแบบจำลองแบบไขว้ทดสอบข้อมูล 10 ชุด แบบจำลองการเรียนรู้ของเครื่องที่ใช้ในการศึกษา ได้แก่ การถดถอยโลจิสติกส์ เพื่อนบ้านใกล้เคียง นาอ็พเบย์ โครงข่ายประสาทเทียมชนิดเพอร์เซ็ปตรอนหลายชั้น ซัพพอร์ต

เวกเตอร์แมชชีน ป่าสุ่ม และโครงข่ายประสาทเทียมชนิดหน่วยความจำระยะสั้นแบบยาว ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยด้านค่าความถูกต้อง ค่าความแม่นยำ ค่าการเรียกคืน และค่าถ่วงดุล ของการเรียนรู้ของเครื่องชนิดต่าง ๆ กลุ่มโมเดลที่พบว่าดีที่สุดในงานวิจัยนี้ได้แก่ DT, RF, SVM, LSTM และ MLP ซึ่งเหมาะสำหรับการเรียนรู้ของเครื่องที่จะนำไปใช้เป็นระบบตรวจข่าวปลอม

คำสำคัญ: ข่าวปลอม, ปัญญาประดิษฐ์, การสืบค้นสารสนเทศ, การประมวลผลภาษาธรรมชาติ, การเรียนรู้ของเครื่อง

Abstract

With modern digital technology, it is more convenient for users to distribute news on online media. Some of them have spread a lot of fake news, which is a problem causing people to believe fake news as real news and share it with others. Currently, there is no tool to detect fake news and interrupt the spread of Thai fake news. Detecting fake news automatically is a difficult and challenging task due to the dynamics of the news. This research presents a method for detecting fake news in Thai with artificial intelligence techniques based on information retrieval, natural language processing, and machine learning. In preparing the training data, the researchers used web crawlers to collect 53,220 samples and pre-categorize them as real, fake, and suspicious news. We balanced the number of data samples in each group to avoid biasing. We trained machine learning models based on 10-fold cross validation. The machine learning models used in the study were Logistic Regression (LR), K-Nearest Neighbor (KNN), Naïve Bayesian (NB), Multilayer Perceptron (MLP), Support Vector Machine (SVM), Random Forest (RF) and Short-Term Long-Term Memory (LSTM). We compared the average performance of different types of machine learning based on accuracy, precision, recall, and f-measure. The models found best in this study were DT, RF, SVM, LSTM, and MLP, suitable for machine learning a fake news detection system.

Keywords: fake news, artificial intelligence, information retrieval, natural language processing, machine learning

บทนำ

ด้วยเทคโนโลยีสารสนเทศหรือดิจิทัลในปัจจุบัน การสื่อสารข้อมูลสามารถส่งถึงกันได้อย่างรวดเร็ว วัฒนาการของเทคโนโลยีสารสนเทศและการสื่อสารได้เพิ่มจำนวนผู้ใช้อินเทอร์เน็ตอย่างมากซึ่งเปลี่ยนวิธีการใช้ข้อมูลและข่าวสารจากแบบดั้งเดิมเป็นแบบดิจิทัล ทำให้เกิดความสะดวกรวดเร็วทั้งผู้เสนอข่าวและผู้อ่านข่าว ในความสะดวกรวดเร็วของระบบอินเทอร์เน็ตทำให้เกิดข่าวปลอมขึ้นจำนวนมากเช่นกัน ความนิยมของสาธารณชนในโซเชียลมีเดีย และเครือข่ายสังคม ทำให้เกิดการแพร่กระจายของข่าวปลอม โดยที่การบิดเบือน และมุมมองที่รุนแรง

ข่าวปลอมได้กลายเป็นหนึ่งในความกังวลที่สำคัญเนื่องจากมีศักยภาพที่จะทำให้รัฐบาลไม่มั่นคง หรืออาจทำให้เป็นอันตรายต่อสังคม ด้วยสื่อสมัยใหม่ ข่าวปลอมอาจกล่าวได้ว่าเป็นภัยคุกคามที่สำคัญ ส่งผลให้ความเชื่อมั่นของรัฐบาลลดลง กระทบประชาชนในทุก ๆ ด้าน

การตรวจจับและลดผลกระทบของข่าวปลอมเป็นหนึ่งในปัญหาพื้นฐานของยุคสมัยปัจจุบันและได้รับความสนใจอย่างกว้างขวาง ข่าวปลอมเป็นหัวข้อเกี่ยวข้องกับหลายสาขาซึ่งนักวิชาการสาขาต่าง ๆ ได้ทำการศึกษาแง่มุมต่าง ๆ ของข่าวปลอม เช่น การเรียนรู้ของเครื่อง ฐานข้อมูล วารสารศาสตร์ รัฐศาสตร์ และอื่น ๆ อีกมากมาย (Lakshmanan *et al.*, 2019) ข่าวปลอมที่เผยแพร่อยู่บนอินเทอร์เน็ตมีข้อมูลในรูปแบบที่หลากหลาย เช่น เอกสาร วิดีโอ และไฟล์เสียง ข่าวปลอมออนไลน์ในรูปแบบที่ไม่มีโครงสร้าง จึงค่อนข้างยากที่จะตรวจจับและจำแนกเนื่องจากต้องใช้ความเชี่ยวชาญของมนุษย์อย่างเคร่งครัด อย่างไรก็ตาม เทคนิคการคำนวณเชิงปัญญาประดิษฐ์ เช่น การประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) และการเรียนรู้ของเครื่อง (Machine Learning: ML) สามารถนำมาประยุกต์ใช้ตรวจจับและจำแนกข้อความหรือบทความที่เป็นประเภทหลอกลวงโดยธรรมชาติจากบทความที่อิงข้อเท็จจริง (Ahmad *et al.*, 2020) ในช่วงหลายปีที่ผ่านมาได้มีนักวิจัยเกี่ยวกับปัญญาประดิษฐ์ที่นำเสนอการตรวจข่าวปลอมจำนวนมาก

Granik & Mesyura (2017) ใช้วิธีการง่าย ๆ สำหรับการตรวจข่าวปลอมโดยใช้ตัวแยกประเภทนาอิวเบย์ (Naïve Baye) วิธีนี้ใช้เป็นระบบซอฟต์แวร์และทดสอบกับชุดข้อมูลของโพสต์ข่าวบนเฟซบุ๊ก โดยมีความแม่นยำในการจำแนกประเภทที่ประมาณ 74% ในชุดทดสอบ ซึ่งเป็นผลลัพธ์ที่ดีเมื่อพิจารณาจากความเรียบง่ายสัมพัทธ์ของแบบจำลอง อีกทั้ง Shu *et al.* (2019) ได้ศึกษาและนำเสนอวิธีการต่าง ๆ เพื่อจะตรวจสอบคุณสมบัติข่าวปลอมและวิธีการแพร่กระจายข่าวทางเครือข่ายสังคมออนไลน์ แนะนำประเภทเครือข่ายยอดนิยม และเสนอว่าเครือข่ายเหล่านี้สามารถใช้ในการตรวจจับและบรรเทาข่าวปลอมในสื่อสังคมออนไลน์ได้

สำหรับเทคนิคใหม่ ๆ Shishah (2021) ใช้เทคนิคตัวแทนเข้ารหัสแบบสองทิศทางจากตัวแปลง (Bidirectional Encoder Representations from Transformers: BERT) การเรียนรู้ร่วมกันที่รวมการจำแนกคุณสมบัติเชิงสัมพันธ์ (Relational Features Classification: RFC) และความสัมพันธ์ชื่อนิบุคคล (Name Entity Relation: NER) การทดลองกับชุดข้อมูลจริงสองชุด พบว่าประสิทธิภาพเป็นที่น่าพอใจ โดยวัดด้วยความถูกต้อง (Accuracy) ค่าถ่วงดุล (F-Measure) และพื้นที่ใต้เส้นโค้ง (Area Under Curve: AUC) เอกลักษณะของกรอบการทำงานร่วมกันทำให้น้ำหนักที่มีความหมายต่อคุณลักษณะ ซึ่งนำไปสู่ประสิทธิภาพที่ดีขึ้นเมื่อเทียบกับเทคนิคปัญญาประดิษฐ์อื่น ๆ นอกจากนี้ Birunda & Devi (2021) ได้เสนอการตรวจข่าวปลอมจากแหล่งข่าวหลายแหล่งแบบอัตโนมัติ โดยแยกคุณลักษณะแบบข้อความจากบทความข่าวจริงและข่าวปลอมโดยใช้ค่าน้ำหนักซึ่งเป็นการถ่วงดุลและค่าถ่วงเอกสารกลับด้าน (Term Frequency - Inverted Document Frequency: TF-IDF) และใช้คะแนนความน่าเชื่อถือของแหล่งที่มาคำนวณตามคุณลักษณะไชด์ยูอาร์แอล ร่วมกับคุณสมบัติแบบข้อความกับคะแนนความน่าเชื่อถือของแหล่งที่มาหลายแหล่ง ความน่าเชื่อถือของข่าวจะถูกประมาณการ กรอบการทำงานที่เสนอนำไปใช้กับตัวแยกประเภทการ

เรียนรู้ของเครื่องเพื่อตรวจสอบประสิทธิภาพในการตรวจหาข่าวปลอม ผลการทดลองได้ประสิทธิภาพของกรอบงานที่เสนอด้วยอัลกอริทึม Gradient Boosting ประมาณ 99.5% จนถึงระดับสูงสุด (Birunda & Devi, 2021) ในการตรวจจับข่าวปลอมภาษาอื่นๆ ที่ไม่ใช่ภาษาอังกฤษ DongHo Lee *et al.* (2019) เสนอสถาปัตยกรรมการเรียนรู้เชิงลึกสำหรับการตรวจจับข่าวปลอมที่เขียนเป็นภาษาเกาหลี โดยใช้สถาปัตยกรรมการเรียนรู้เชิงลึก และใช้รูปแบบการฝังคำที่เรียนรู้โดยหน่วยพยางค์ หลังจากการสอนและทดสอบการใช้งาน ระบบมีความถูกต้องที่มีความหมายสำหรับการจัดประเภทเนื้อหาและความคลาดเคลื่อนของบริบท แต่ความแม่นยำยังต่ำสำหรับการจำแนกประเภทพาดหัวและความคลาดเคลื่อนของเนื้อหา

ในประเทศไทยเองก็เริ่มมีงานวิจัยเกี่ยวกับการตรวจจับข่าวปลอมของไทยใช้เทคนิคปัญญาประดิษฐ์ Aphiwongsophon & Chongstitvatana (2018) ทำการศึกษาข่าวปลอมภาษาไทยจาก Twitter โดยใช้วิธีการยอคนิยมสามวิธีในการทดลอง ได้แก่ นาอ็ฟเบย์ (Naive Bayes: NB) โครงข่ายประสาท (Neural Network: NN) และซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) การศึกษาพบว่ากระบวนการทำให้ข้อมูลเป็นมาตรฐานเป็นสิ่งจำเป็นสำหรับการทำความสะอาดข้อมูลก่อนที่จะป้อนข้อมูลไปสู่แบบจำลองการเรียนรู้ของเครื่องเพื่อจำแนกประเภทข้อมูล แบบจำลองที่ดีที่สุดคือ SVM มีความถูกต้องถึง 99.90% ซึ่งงานวิจัยนี้เน้นที่ข้อมูล Twitter เท่านั้น ไม่มีการใช้งานสำหรับการตรวจจับออนไลน์ นอกจากนี้ Mookdarsanit & Mookdarsanit (2021) เสนอกรอบการเรียนรู้เชิงลึกสำหรับการตรวจจับข่าวปลอมของไทยเกี่ยวกับโควิด-19 โดยใช้ข้อความจากโซเชียล งานวิจัยนี้ได้สร้างแบบจำลองการเรียนรู้แบบถ่ายโอน ซึ่งรวมถึง BERT, Universal Language Model Fine-Tuning (ULMFIT) และ Generative Pre-trained Transformer (GPT) นักวิจัยใช้ชุดข้อมูลเปิดข่าวโควิด-19 ที่แปลจากภาษาอังกฤษเป็นภาษาไทยและเตรียมการสอนแบบจำลองการเรียนรู้เชิงลึกเกี่ยวกับโควิด-19 เพื่อปรับแต่งชุดข้อมูลให้เข้ากับประเทศไทย โดยนักวิจัยใช้ข้อมูลเพิ่มเติมโดยรวบรวมข้อมูลข้อความภาษาไทยจากโซเชียลมีเดียและระบุว่าเป็นประเภทข่าวปลอมและข่าวจริง ผลลัพธ์ที่ดีที่สุดจากการทดลองทำให้ประสิทธิภาพความถูกต้องดีที่สุดถึง 72.93% ไม่มีรายงานกรณีการใช้งานจริงสำหรับข่าวปลอมของไทยในการวิจัยนี้

สำหรับงานวิจัยครั้งนี้ผู้วิจัยเสนอวิธีทางปัญญาประดิษฐ์สำหรับการตรวจข่าวปลอมภาษาไทย วิธีที่นำเสนอในงานวิจัยนี้ผู้วิจัยเสนอกรอบการทำงานเพื่อสร้างระบบตรวจจับข่าวปลอมออนไลน์ของไทยโดยอัตโนมัติ กรอบงานที่เสนอประกอบด้วยสามโมดูล ได้แก่ การสืบค้นข้อมูลสารสนเทศ การประมวลผลภาษาธรรมชาติ และการเรียนรู้ของเครื่อง งานวิจัยนี้มีวัตถุประสงค์หลัก คือ 1) นำเสนอกรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์ 2) นำเสนอลักษณะเด่นจากการวิเคราะห์ภาษาธรรมชาติ 3) นำเสนอแบบจำลองการเรียนรู้ของเครื่องสำหรับการตรวจจับข่าวปลอมภาษาไทย และ 4) พัฒนาเว็บแอปพลิเคชันตรวจข่าวปลอมภาษาไทย

วรรณกรรมที่เกี่ยวข้อง

หลายปีที่ผ่านมา มีนักวิจัยนำเสนอการตรวจข่าวปลอมด้วยเทคนิคด้านปัญญาประดิษฐ์มากมายหลายกลุ่มด้วยกัน โดยส่วนมากใช้เทคนิคที่เกี่ยวข้องกับปัญญาประดิษฐ์ส่วนมากอยู่ในระดับการประยุกต์ใช้การเรียนรู้ของเครื่อง ซึ่งเป็นการสอนข้อมูลให้เครื่องรู้รูปแบบของข่าว ในการที่จะสอนข้อมูลข้อความข่าวให้เครื่องรู้จำนั้นต้องใช้ศาสตร์ที่เกี่ยวข้องหลัก ๆ คือการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) (Ayoub et al., 2021) ซึ่งเป็นปัญญาประดิษฐ์ (Artificial Intelligence: AI) ที่เกี่ยวข้องกับการให้คอมพิวเตอร์สามารถเข้าใจข้อความและคำพูดในลักษณะเดียวกับที่มนุษย์สามารถทำได้ เป็นกระบวนการวิเคราะห์ข้อมูลประเภทข้อความซึ่งเป็นภาษาที่มนุษย์ใช้กันแบบธรรมชาติ กระบวนการที่สำคัญในการประมวลผลภาษาธรรมชาติ คือการทำเหมืองข้อความ (Text Mining) Vijayarani & Janani (2016) ซึ่งเป็นการขุดข้อความระบุข้อเท็จจริง ความสัมพันธ์ และคำยืนยันที่จะถูกฝังอยู่ในมวลของข้อมูลขนาดใหญ่ที่เป็นข้อความ เมื่อดึงออกมาแล้วข้อมูลนี้จะถูกแปลงเป็นรูปแบบที่มีโครงสร้างที่สามารถนำไปวิเคราะห์เพิ่มเติม การทำเหมืองข้อความใช้วิธีการที่หลากหลายในการประมวลผลข้อความ ซึ่งเป็นหนึ่งในวิธีที่สำคัญที่สุด ขั้นตอนการทำเหมืองข้อความ มีขั้นตอนพื้นฐานที่เกี่ยวข้องในการเตรียมเอกสารข้อความที่ไม่มีโครงสร้างสำหรับการวิเคราะห์เชิงลึก ได้แก่ 1. การเตรียมข้อมูลข้อความ (Text Pre-processing) 2. การแปลงข้อความ (Text Transformation) 3. การเลือกลักษณะเด่น (Feature Selection) 4. การสร้างแบบจำลองการเรียนรู้ของเครื่อง (Machine Learning Modeling) 5. การประเมิน (Evaluation) และ 6. การนำไปประยุกต์ใช้ (Deployment)

ขบวนการสำคัญอีกประการหนึ่งในการเตรียมข้อมูลสำหรับการประมวลผลภาษาธรรมชาติคือการเก็บรวบรวมข้อมูล แหล่งข้อมูลมีมากมายหลายแหล่ง ทั้งที่เป็นข้อมูลที่จัดเก็บในองค์แบบในรูปแบบไฟล์อิเล็กทรอนิกส์ชนิดต่าง ๆ ข้อมูลที่จัดเก็บในฐานข้อมูลแบบมีโครงสร้างและแบบไม่มีโครงสร้าง ข้อมูลที่อยู่บนเว็บไซต์ทั่วไปและบนสื่อสังคมออนไลน์ วิธีการสืบค้นสารสนเทศ (Information Retrieval: IR) (Krallinger et al., 2017) เป็นเทคนิควิธีสำหรับการสืบค้นข้อมูล ซึ่งสามารถนำมาประยุกต์ใช้ในการเก็บรวบรวมข้อมูล สำหรับข้อมูลที่อยู่บนเว็บไซต์ทั่วไป เทคนิคการสืบค้นจะเป็นการใช้เว็บครอว์เลอร์ค้นหาข้อมูลและสกัดข้อมูลเฉพาะที่ต้องใช้งาน ในกรณีการสร้างระบบตรวจสอบข่าวปลอม การสืบค้นสารสนเทศจะเป็นการสืบค้นข่าวที่เกี่ยวข้องกับข่าวสืบค้น และเว็บครอว์เลอร์จะไปสืบค้นและรับข้อความข่าวที่มีความคล้ายกลับข่าวสืบค้น รวมทั้งข้อมูลเมตาที่เกี่ยวข้องกับแหล่งข่าวเพื่อใช้ในการจำแนกชนิดข่าว ข่าวสืบค้นที่ได้รับกลับมาก็จะถูกส่งต่อไปยังการประมวลผลภาษาธรรมชาติ เพื่อสกัดลักษณะเด่นของข่าวก่อนส่งต่อไปยังการเรียนรู้ของเครื่องในขั้นตอนถัดไป

การเรียนรู้ของเครื่อง (Machine Learning: ML) (Gori, 2018) เป็นเซตย่อยของปัญญาประดิษฐ์ ที่มีความสามารถในการเรียนรู้ข้อมูลสอนที่จัดเตรียมไว้ล่วงหน้า การเรียนรู้ของเครื่องแบ่งตามชนิดการสอนได้หลายกลุ่มหลัก ๆ ได้แก่ การเรียนรู้แบบมีผู้สอน (Supervised Learning) การเรียนรู้แบบไม่มีผู้สอน (Unsupervised Learning) การเรียนรู้แบบกึ่งมีผู้สอน (Semi-supervised Learning) และการเรียนรู้แบบเสริมแรง (Reinforcement Learning) สำหรับการเรียนรู้แบบมีผู้สอนยังสามารถแบ่งการทำงานได้อีกสองแบบ ได้แก่ การถดถอย (Regression)

และการจำแนกประเภทหรือกลุ่ม (Classification) สำหรับการสร้างแบบจำลองตรวจจำวปลอมจะอยู่ในกลุ่มการจำแนกประเภทซึ่งมีอยู่มากมายหลายเทคนิค ในที่นี้จะขอลำถึงดังต่อไปนี้ การถดถอยโลจิสติกส์ เทคนิคเพื่อนบ้านใกล้เคียง โครงข่ายประสาทเทียมชนิดเพอร์เซ็ปตรอนหลายชั้น ชัฟฟอร์ตเวกเตอร์แมชชีน นาอ์ฟเบย์ ต้นไม้ตัดสินใจ ป่าสุ่ม และโครงข่ายประสาทเทียมชนิดความจำระยะสั้นแบบยาว

การถดถอยโลจิสติกส์ (Logistics Regression: LR) (Kleinbaum & Klein, 2010) การถดถอยโลจิสติกส์เป็นแบบจำลองทางสถิติที่ในรูปแบบพื้นฐานใช้ฟังก์ชันโลจิสติกส์เพื่อสร้างแบบจำลองตัวแปรตามแบบไบนารี การถดถอยโลจิสติกส์เป็นการประมาณค่าพารามิเตอร์ของแบบจำลองโลจิสติกส์ การถดถอยโลจิสติกส์สามารถนำไปใช้การวิเคราะห์เชิงคาดการณ์ตัวแปรแบบไบนารี เพื่ออธิบายข้อมูลและอธิบายความสัมพันธ์ระหว่างตัวแปรอิสระอย่างน้อยหนึ่งตัวแปรกับตัวแปรตามที่เป็นค่าเอาต์พุตซึ่งมีค่าเป็นไปได้สองค่าคือ “0” และ “1” ตัวแปรอิสระหรืออินพุตซึ่งอาจจะเป็นค่าระดับ ค่าลำดับ ค่าเป็นช่วง อัตราส่วน ไบนารี หรือค่าต่อเนื่อง อย่างน้อยหนึ่งตัวแปร ในการตัดสินใจเอาต์พุตเป็นการรวมกันแบบเชิงเส้นของตัวแปรอิสระ ฟังก์ชันแบบโลจิสติกส์ใช้แปลงการรวมกันแบบเชิงเส้นของอินพุตเป็นความน่าจะเป็นซึ่งเป็นช่วงระหว่าง 0 ถึง 1 ที่สอดคล้องกับค่าเป้าหมายสองกลุ่ม “0” หรือ “1”

วิธีการเพื่อนบ้านใกล้เคียง (K-Nearest Neighbor: KNN) (Guo et al., 2003) เป็นการเรียนรู้ของเครื่องแบบมีผู้สอน และเป็นวิธีการที่ได้รับความนิยมในการใช้งานอย่างมาก เนื่องจากความง่ายในการประยุกต์ใช้งานแต่มีประสิทธิภาพสูง และสามารถนำไปประยุกต์ใช้กับงานได้อย่างหลากหลาย โดยเฉพาะอย่างยิ่งด้านการจำแนกข้อมูล ขั้นตอนการทำงานของวิธีการเพื่อนบ้านใกล้เคียง ได้แก่ 1. เก็บข้อมูลชุดสอนเป็นแบบจำลอง 2. กำหนดค่า K เป็นจำนวนสมาชิกเพื่อนบ้านใกล้เคียง 3. คำนวณหาระยะห่างระหว่างจุดด้วยวิธีต่าง ๆ อาจจะใช้ Euclidian Distance, Cosine Distance หรือวิธีอื่น ๆ เป็นการวัดค่าระยะห่างระหว่างข้อมูลทดสอบกับข้อมูลชุดสอน 4. เรียงลำดับระยะห่างระหว่างข้อมูลทดสอบและข้อมูลสอน โดยพิจารณาจากข้อมูลจำนวน K เรคคอร์ดที่ใกล้ที่สุด และ 5. ตัดสินใจโดยการโหวตผลลัพธ์จากค่าเป้าหมายของข้อมูลชุดสอนที่อยู่ใกล้ชุดทดสอบที่สุดจาก K เรคคอร์ด

โครงข่ายประสาทเทียมชนิดเพอร์เซ็ปตรอนหลายชั้น (Multilayer Perceptron: MLP) (Hagan et al., 2014) มีพื้นฐานมาจากการจำลองการทำงานของสมองมนุษย์ ด้วยโปรแกรมคอมพิวเตอร์ จุดมุ่งหมายของโครงข่ายประสาทเทียมคือต้องการให้คอมพิวเตอร์มีความชาญฉลาดในการเรียนรู้เหมือนที่มนุษย์มีการเรียนรู้ สามารถฝึกฝนได้ และสามารถนำความรู้และทักษะ รวมทั้งสามารถนำไปประยุกต์ใช้ได้ดีกับปัญหาด้านการจำแนกกลุ่ม การถดถอย และการจัดกลุ่ม เทคนิคนี้มักถูกเรียกว่า “Black Box” เนื่องจากการทำงานมีความซับซ้อนมากกว่าเทคนิคอื่นๆ ก่อนข้างมาก มีการนำมาประยุกต์ใช้ในงานด้านต่าง ๆ ข้อดีของการจำแนกกลุ่มด้วยวิธีโครงข่ายประสาทเทียม ได้แก่ ความมีประสิทธิภาพในการพยากรณ์ข้อมูลสูง และรองรับการจำแนกข้อมูลได้ทั้งข้อมูลที่เป็นลักษณะตัวเลข ต่อเนื่อง ข้อจำกัดของการจำแนกกลุ่มด้วยแบบจำลองโครงข่ายประสาทเทียม ได้แก่ 1) บางครั้งอาจเกิดเหตุการณ์เรียนรู้ข้อมูลสอนมากเกินไปทำให้ แบบจำลองเรียนรู้ได้ผลลัพธ์ที่ดีเฉพาะในข้อมูลสอน แต่กลับมีประสิทธิภาพไม่ดีเมื่อนำมาใช้กับข้อมูลในสภาพแวดล้อมจริง 2) กระบวนการทำงานของแบบจำลองเป็นแบบ Black block จึงไม่

สามารถทราบเหตุผลการตัดสินใจของแบบจำลอง 3) ใช้เวลาในการเรียนรู้และใช้หน่วยความจำมาก หากข้อมูลมีขนาดใหญ่ และ 4) ไม่คงทนต่อข้อมูลรบกวนและข้อมูลแปลกปลอม

ซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) (Cortes & Vapnik, 1995) เป็นเครื่องมือสำหรับการเรียนรู้แบบใหม่บนพื้นฐานของการเรียนรู้ทฤษฎีทางสถิติ มีประโยชน์อย่างมากในการนำไปใช้จำแนกข้อมูลหรือการรู้จำรูปแบบของข้อมูลโดยวิธี SVM ได้มีงานวิจัยหลายงานที่น่าวิธีการดังกล่าวไปประยุกต์ใช้ SVM มีแนวคิดของทฤษฎีดังนี้ 1) การลดความเสี่ยงเชิงโครงสร้างให้ต่ำที่สุด เป็นแนวคิดที่แสดงถึงขอบเขตของความเสียหายหรือความน่าจะเป็นที่จะเกิดความผิดพลาดของการเรียนรู้ มีประโยชน์สำหรับการนำมาใช้ในกระบวนการเรียนรู้เพื่อหาฟังก์ชันในการตัดสินใจเพื่อให้เกิดความผิดพลาดน้อยที่สุด 2) ปริภูมิแสดงลักษณะสำคัญกับส่วนของเคอร์เนลฟังก์ชัน ที่จะทำการแมปข้อมูลจากปริภูมินำเข้าไปสู่ปริภูมิแสดงลักษณะที่สำคัญ เพื่อประโยชน์ในการสร้างฟังก์ชันสำหรับการตัดสินใจที่มีลักษณะไม่เป็นไปในรูปแบบเชิงเส้นกับข้อมูลในปริภูมินำเข้า ทำให้การแบ่งกลุ่มข้อมูลมีความยืดหยุ่นมากขึ้น และ 3) ระบายหลายมิติที่ใช้แยกข้อมูลได้ดีที่สุด เป็นแนวคิดหลักของเทคนิค SVM เพื่อทำการหาระนาบที่มีระยะห่างจากข้อมูลทำการแบ่งให้มากที่สุด สำหรับการแบ่งกลุ่มข้อมูลสองกลุ่มออกจากกัน งานวิจัยทั่วไปพบว่า SVM มีข้อดีด้านความคงทนต่อข้อมูลรบกวน

นาอิวเบย์ (Naïve Bayesian: NB) (Yager, 2006) เป็นตัวแยกประเภทในกลุ่มความน่าจะเป็น ประยุกต์ใช้ทฤษฎีของเบย์ (Baye's Theorem) NB มีสมมติฐานความเป็นอิสระที่ระหว่างตัวแปรต่าง ๆ นาอิวเบย์ได้รับการศึกษาอย่างกว้างขวางตั้งแต่ทศวรรษ 1960 ยังคงเป็นวิธีที่นิยมถึงปัจจุบัน สำหรับการจัดหมวดหมู่ข้อความ การจำแนกเอกสาร เช่น การจำแนกอีเมล กฎหมาย กีฬา และการเมือง เป็นต้น โดยนาอิวเบย์จะใช้ข้อมูลสอนที่เป็นค่าจากข้อความแปลงเป็นค่าคุณลักษณะที่ถูกประมวลผลไว้ล่วงหน้า ด้วยความไม่ซับซ้อนของนาอิวเบย์จึงได้รับความสนใจถูกนำไปใช้ในแอปพลิเคชันมากมาย นาอิวเบย์สามารถปรับขนาดของงานขนาดใหญ่ โดยต้องการพารามิเตอร์จำนวนหนึ่งเป็นเชิงเส้นในจำนวนตัวแปรอินพุตและเอาต์พุต นาอิวเบย์เป็นเทคนิคที่นิยมใช้เปรียบเทียบกับเทคนิคอื่น ๆ ถือเป็นบรรทัดฐาน

วิธีต้นไม้ตัดสินใจ (Decision Tree: DT) (Myles et al., 2004) มีลักษณะโครงสร้างแบบต้นไม้หัวกลับ เป็นการนำข้อมูลมาสร้างแบบจำลองการพยากรณ์ในรูปของต้นไม้ตัดสินใจ โดยมีการใช้แอททริบิวต์ของข้อมูลกำหนดเป็นโหนดต่าง ๆ และกิ่งก้านของต้นไม้เป็นค่าของข้อมูลในแต่ละแอททริบิวต์ ต้นไม้ตัดสินใจเริ่มต้นด้วยโหนดราก แยกข้อมูลตามเงื่อนไขไปตามกิ่งก้าน เชื่อมต่อกับโหนดระหว่างกลาง และแยกข้อมูลตามกิ่งก้านตามเงื่อนไขต่อไปเรื่อย ๆ ไปสิ้นสุดที่โหนดใบซึ่งเป็นโหนดตัดสินใจ ซึ่งต้นไม้ตัดสินใจนั้นทำงานแบบเรียนรู้แบบมีผู้สอน สามารถประยุกต์ใช้ในการจำแนกประเภทได้ จากกลุ่มตัวอย่างของข้อมูลที่ได้กำหนดได้ก่อนล่วงหน้าซึ่งใช้เป็นข้อมูลสอน และสามารถพยากรณ์กลุ่มของรายการที่ยังไม่เคยพบมาก่อนได้ ต้นไม้ตัดสินใจเป็นแบบจำลองที่โครงสร้างไม่ซับซ้อน สร้างได้ง่าย และตีความเป็นกฎการตัดสินใจได้ จึงเป็นที่นิยมใช้ในกระบวนการทำเหมืองข้อมูลและนำไปใช้ในการตัดสินใจ

ป่าสุ่ม (Random Forest: RF) (Svetnik et al., 2003) หรือป่าตัดสินใจสุ่ม (Random Decision Forest) เป็นวิธีการเรียนรู้ทั้งมวล (Ensemble Learning) ถูกนำไปใช้สำหรับการจำแนกประเภท การถดถอย และงานอื่น ๆ มากมาย การสร้างแบบจำลองป่าสุ่มเป็นการสร้างต้นไม้การตัดสินใจจำนวนมากจากข้อมูลสอน โดยมีแนวคิดว่าการช่วยกันกีดกันต้นไม้มีโอกาสลดข้อผิดพลาดได้ ในการจำแนกประเภทผลลัพธ์ในการตัดสินใจของป่าสุ่มได้จากประเภทจากการตัดสินใจโดยต้นไม้ส่วนใหญ่ สำหรับงานการถดถอยใช้ค่าเฉลี่ยหรือการคาดการณ์ของต้นไม้แต่ละต้นในการให้คำตอบตัดสินใจ โดยทั่วไปแล้วป่าสุ่มจะมีประสิทธิภาพดีกว่าต้นไม้ตัดสินใจเดี่ยว โดยลักษณะข้อมูลอาจส่งผลต่อประสิทธิภาพการทำงานได้ด้วย

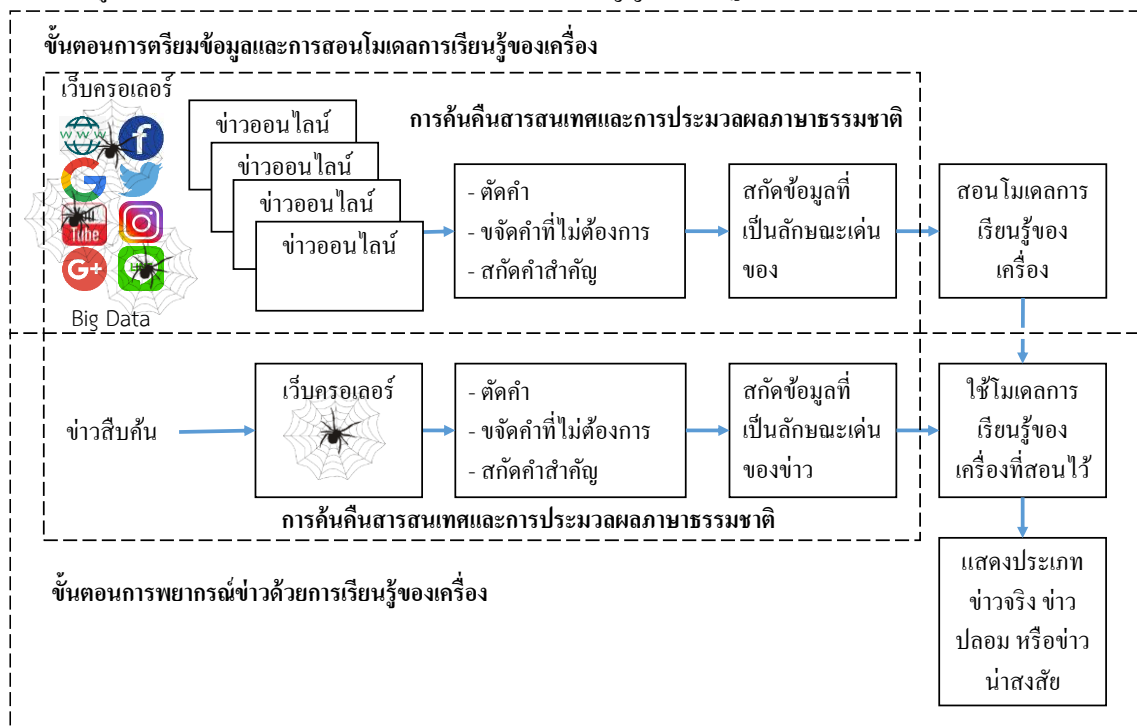
โครงข่ายประสาทเทียมแบบความจำระยะสั้นระยะยาว (Long Short-Term Memory Neural Network: LSTM) (Hochreiter & Schmidhuber, 1997) เป็นการเรียนรู้เชิงลึกกลุ่มเดียวกับโครงข่ายประสาทเทียมแบบวนซ้ำ (RNN) LSTM มีการย้อนกลับของข้อมูล จึงทำให้ไม่เพียงแต่สามารถประมวลผลข้อมูลได้แต่ยังสามารถประมวลผลข้อมูลแบบลำดับได้เป็นอย่างดี LSTM สามารถใช้ได้กับงานต่าง ๆ เช่น การจำแนกความรู้สึกจากเอกสาร การรู้จำลายมือที่เชื่อมต่อ การรู้จำเสียงพูด และการตรวจจับความผิดปกติของข้อมูลเครือข่าย นิวรอนของ LSTM ทั่วไปเรียกว่าเซลล์ ประกอบด้วยประตูเข้า ประตูส่งออก และประตูลืม เซลล์จะจดจำค่าในช่วงเวลาต่าง ๆ และประตูทั้งสามจะควบคุมการไหลของข้อมูลเข้าและออกจากเซลล์ LSTM เหมาะสมอย่างยิ่งกับการจำแนกการประมวลผลและการทำนายตามข้อมูลอนุกรมเวลา เนื่องจากอาจมีช่วงเวลาที่ไม่มีทราบระยะเวลาระหว่างเหตุการณ์สำคัญในอนุกรมเวลา LSTM ถูกการพัฒนาขึ้นเพื่อแก้ปัญหาการลู่สู่ค่าอนันต์ และการลู่เข้าสู่ศูนย์ของค่าความลาดชันค่าผิดพลาดที่พบในขณะสอนของ RNN แบบเดิม

การดำเนินงานวิจัย

การดำเนินงานวิจัย ประกอบด้วย 4 ขั้นตอนหลัก ได้แก่ 1) การออกแบบกรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์ 2) การนำเสนอลักษณะเด่นจากการวิเคราะห์ภาษาธรรมชาติ 3) การสร้างแบบจำลองการเรียนรู้ของเครื่องสำหรับการตรวจข่าวปลอมภาษาไทย และ 4) การพัฒนาเว็บแอปพลิเคชันตรวจข่าวปลอมภาษาไทย

เนื้อหาข่าวปลอมเกิดขึ้นจากมีผู้ที่ตั้งใจให้มีความหมายเปลี่ยนไปจากข้อมูลเดิมโดยมีวัตถุประสงค์แตกต่างกันไป ข่าวปลอมจึงเป็นข้อมูลที่กว้างและมีความเป็นพลวัตสูง ทำให้ยากต่อการสร้างการเรียนรู้ของเครื่องให้เป็นระบบตรวจข่าวปลอมที่สามารถเรียนรู้ตามเนื้อหาที่เปลี่ยนไป อย่างไรก็ตามมีความเป็นไปได้ที่จะสร้างระบบตรวจข่าวปลอมมีประสิทธิภาพสูง จากการศึกษาพบว่าผู้ผลิตข่าวนิยมเผยแพร่ข่าวออนไลน์เพื่อผู้คนสามารถเข้าถึงข่าวได้อย่างรวดเร็วผ่านอินเทอร์เน็ตจากทุกที่ อีกทั้งสื่อสังคมออนไลน์ได้เก็บทั้งข่าวจริงและข่าวปลอมบนคลาวด์เซิร์ฟเวอร์ มีผู้อ่านข่าวได้แสดงความคิดเห็นในเว็บโซเชียลมีเดียจำนวนมาก จึงเป็นแหล่งข้อมูลที่สำคัญในการตรวจข่าวได้แบบอัตโนมัติและมีประสิทธิภาพสูงได้

ในงานวิจัยนี้ผู้วิจัยเสนอกรอบงานการตรวจจับข่าวปลอมด้วยการเรียนรู้ด้วยเทคนิคปัญญาประดิษฐ์ ดังรูปที่ 1 โดยกรอบงานตรวจสอบข่าวปลอมที่นำเสนอ แบ่งเป็นสองขั้นตอน คือ ขั้นตอนการเตรียมข้อมูลและการสอนแบบจำลองการเรียนรู้ของเครื่อง และขั้นตอนการพยากรณ์ข่าวด้วยการเรียนรู้ของเครื่อง ทั้งสองขั้นตอนมีความเกี่ยวข้องกับโมดูลสามโมดูลหลัก ได้แก่ การค้นคืนสารสนเทศ (IR) การประมวลผลภาษาธรรมชาติ (NLP) และการเรียนรู้ของเครื่อง (ML) โมดูล IR ดึงข้อมูลข่าวสารจากอินเทอร์เน็ตตามข่าวสืบค้นที่ป้อนโดยผู้ใช้ ผลลัพธ์จากการค้นคืนเป็นเนื้อหาข่าวที่เกี่ยวข้องจากแหล่งข่าวออนไลน์มากมาย โมดูล NLP จะวิเคราะห์เอกสารที่ได้รับกลับคืนมา โดยดำเนินการตัดคำ ขจัดคำที่ไม่ต้องการ และสกัดคุณลักษณะข่าว โมดูล ML แบ่งประเภทข่าวออกเป็นสามประเภท ได้แก่ ข่าวจริง ข่าวปลอม และข่าวน่าสงสัย และในส่วนของ “การแสดงผลข้อมูลข่าวที่เกี่ยวข้องกับข่าวสืบค้น” รูปที่ 1 แสดงกรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์



รูปที่ 1 กรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์

'ข่าวไม่ปลอม', 'ข่าวจริง', 'ข่าวไม่เท็จ', 'ข่าวไม่บิดเบือน', 'ไม่ปลอม', 'เป็นข่าวจริง', 'เป็นเรื่องจริง', 'ไม่เท็จ', 'ไม่บิดเบือน', 'เรื่องจริง', 'ข้อความจริง', 'ข้อมูลจริง', 'เป็นข้อมูลจริง', 'ข้อมูลถูก', 'แชร์ต่อได้', 'เป็นความจริง', 'มีมูล'

รูปที่ 2 ตัวอย่างคำที่มักปรากฏในข่าวจริงที่แตกต่างจากคำในข่าวปลอม

'การติดต่อ', 'การหมิ่นประมาท', 'การหลอกลวง', 'การเล่นตลก', 'การโฆษณาชวนเชื่อ', 'ข่าวขยะ', 'ข่าวที่ผิดพลาด', 'ข่าวที่ไม่ถูกต้อง', 'ข่าวที่ไม่น่าเชื่อถือ', 'ข่าวบิดเบือน', 'ข่าวปลอม', 'ข่าวปลอม', 'ข่าวผิด', 'ข่าวมั่ว', 'ข่าวลบ', 'ข่าวเท็จ', 'ข่าวไม่จริง', 'ข่าวไม่ถูกต้อง', 'ข้อความข่าวปลอม', 'ข้อความบิดเบือน', 'ข้อความปลอม', 'ข้อความเท็จ', 'ข้อมูลปลอม', 'ข้อมูลผิด', 'ข้อมูลผิดพลาด', 'ข้อมูลเท็จ', 'ข้อมูลไม่จริง', 'ข้อเท็จจริงเท็จ', 'ข้อเท็จจริงไม่ถูกต้อง', 'ข้อเท็จจริงไร้ค่า', 'ความเชื่อผิดๆ', 'ความเท็จ', 'คำกล่าวอ้าง', 'คำพูดเหลวไหล', 'คำลวง', 'คำเท็จ', 'คำโกหก', 'จับโป๊ะ', 'ชวนเชื่อ', 'ต่อแหล', 'บิดเบือน', 'ปลอม', 'ปลอม', 'มั่ว', 'มโน', 'รายงานเท็จ', 'สื่อสนั่น', 'สื่อหึ่ง', 'ข่าวหลอกลวง', 'อวดอ้าง', 'เท็จ', 'เป็นข้อมูลเท็จ', 'เป็นเฟคนิส', 'เรื่องราวปลอม', 'เรื่องเท็จ', 'เหตุการณ์หลอก', 'โฆษณาชวนเชื่อ', 'โยงมั่ว', 'ไม่น่าเชื่อถือ', 'ไม่มีมูลความจริง', 'ไม่มีอยู่จริง', 'ไม่เป็นความจริง', 'ไม่ใช่ข่าวจริง', 'ไม่ใช่เรื่องจริง'

รูปที่ 3 ตัวอย่างคำที่มักปรากฏในข่าวปลอมที่แตกต่างจากคำในข่าวจริง

ผู้วิจัยทำการเก็บข้อมูลด้วยวิธีการค้นคืนสารสนเทศด้วยเว็บครอว์เลอร์เพื่อเก็บข้อมูลตั้งแต่วันที่ 1 สิงหาคม 2563 ถึงวันที่ 31 พฤษภาคม 2564 ข้อมูลที่ได้รับคืนกลับมาจะใช้สำหรับการสอนแบบจำลองการเรียนรู้ของเครื่อง จำนวน 44,271 ข้อความข่าว แบ่งเป็น 3 ประเภท ได้แก่ ข่าวจริง (real) ข่าวปลอม (fake) และข่าวน่าสงสัย (suspicious) การระบุประเภทเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย นั้น ผู้วิจัยนำข่าวที่ได้รับจากค้นคืนช่วงเริ่มต้นในการเก็บ มาวิเคราะห์เพื่อหาองค์ประกอบของข่าวทั้งสามประเภท เพื่อหาฟีเจอร์หรือลักษณะเด่นของข่าว 5 ค่า ได้แก่ จำนวนคำที่เป็นลักษณะข่าวปลอม (score_fake) จำนวนคำที่เป็นลักษณะข่าวจริง (score_real) ค่าความคล้ายของข่าว (sim_matched) จำนวนโดเมนที่พบระบุข่าวปลอม (domain_fake) จำนวนโดเมนที่พบระบุข่าวจริง (domain_real) สำหรับการคำนวณค่าลักษณะเด่นที่เป็นค่าความคล้ายของข่าว (sim_matched) วัดจากค่าความเหมือนแบบโคไซน์ (cosine similarity) (Soyusiwaty & Zakaria, 2018) ผู้เขียนเป็นสคริปต์เป็นกฎเพื่อจัดเก็บข้อมูลข่าวตามฟีเจอร์และทำการระบุค่าเป้าหมายประเภทของข่าวโดยอัตโนมัติ ซึ่งระบุประเภทเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย จากนั้นผู้วิจัยตรวจสอบความถูกต้องอีกครั้งเพื่อพร้อมใช้สำหรับการสอนให้เครื่องเรียนรู้ รูปที่ 2 แสดงตัวอย่างคำที่มักปรากฏในข่าวจริงที่แตกต่างจากคำในข่าวปลอม และรูปที่ 3 แสดงตัวอย่างคำที่มักปรากฏในข่าวปลอมที่แตกต่างจากคำในข่าวจริง

ผู้วิจัยทำให้สมดุลข้อมูลในแต่ละกลุ่มโดยวิธีการสุ่มเพิ่มแบบ SMOTE (Chawla *et al.*, 2002) เป็นกลุ่มละ 17,740 ข้อความ รวมเป็นจำนวน 53,220 ข้อความข่าว ดังแสดงในตารางที่ 1

ตารางที่ 1 ข้อมูลที่จัดเก็บด้วยวิธีการค้นคืนสารสนเทศด้วยเว็บครอว์เลอร์

ชนิดข่าว	จำนวนข่าว	จำนวนสมมูล (คู่เพิ่ม)
ปลอม (Fake)	17,740	17,740
จริง (Real)	13,072	17,740
น่าสงสัย (Suspicious)	13,459	17,740
รวม	44,271	53,220

งานวิจัยนี้ข้อมูลข่าวสืบค้นถูกสกัดเป็นลักษณะเด่นของข้อมูลจำนวน 5 ค่า ได้แก่ score_fake, score_real, sim_matched, domain_fake และ domain_real ค่าเป้าหมายประกอบด้วยข่าวปลอม (fake) ข่าวจริง (real) และน่าสงสัย (suspicious) ข่าวสืบค้นจะถูกนำไปเทียบกับข่าวออนไลน์ที่มีความคล้ายคลึงกัน ผู้วิจัยใช้หลักการการประมวลผลภาษาธรรมชาติเพื่อสกัดลักษณะเด่น 5 ค่า วิธีการสกัดลักษณะเด่น เริ่มจากการผลการค้นคืนสารสนเทศที่เป็นข่าวคล้ายคลึงกับข่าวสืบค้น ทำการตัดค่าและนับจำนวนค่าที่เป็นลักษณะข่าวปลอม (score_fake) นับจำนวนค่าที่เป็นลักษณะข่าวจริง (score_real) คำนวณค่าความคล้ายของข่าว (sim_matched) นับจำนวนโดเมนที่พบระบุข่าวปลอม (domain_fake) นับจำนวนโดเมนที่พบระบุข่าวจริง (domain_real) ตัวอย่างข้อมูลข่าวสืบค้นและลักษณะเด่นแสดงดังตารางที่ 2

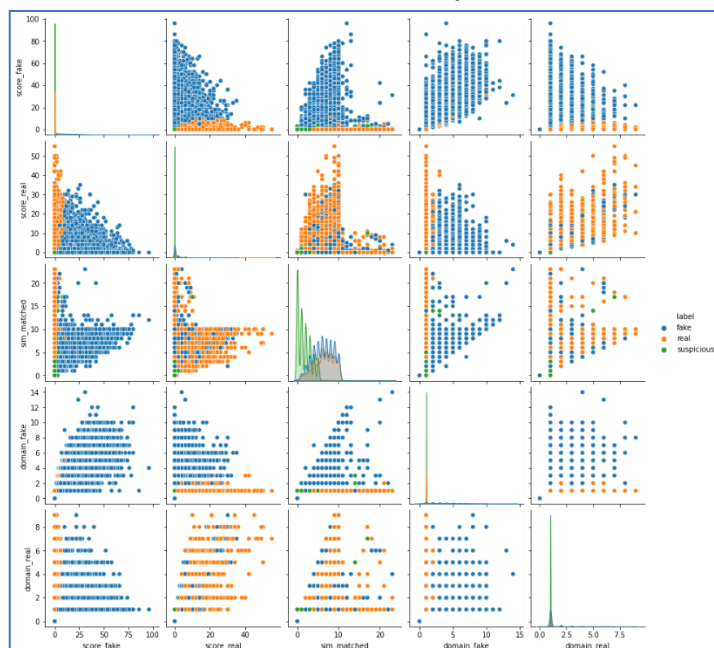
ตารางที่ 2 ตัวอย่างข้อมูลข่าวสืบค้นและลักษณะเด่น

ข่าวสืบค้น	ประเภท	score_fake	score_real	sim_matched	domain_fake	domain_real
กรม. เกาวันที่ 4 และ 7 ก.ย. 63 เป็นวันหยุดชดเชยวันสงกรานต์	real	2	4	21	1	3
ฆ่าเชื้อโควิด-19 ด้วยการกินผลไม้ ที่มีฤทธิ์เป็นด่าง	fake	7	0	12	6	1
ขุนพลเคนมาร์กับการสร้าง กำแพงปกป้องเอริเคนในวันที่ หมดสติบนสนามของศึกยูโร 2020	suspicious	0	0	0	0	0
หมู-ไก่ เป็นโรคเอดส์	fake	31	4	23	14	4
กรมอนามัย รุกโรงงานขนาด ใหญ่ ประเมิน Thai Stop COVID Plus ครบ 100 % 15 มิถุนายนนี้	suspicious	0	0	2	0	0
รพท. เปิดมหรหรรษ์ใกล้เกลี่ย สินเชื่อเช่าซื้อรถ ผ่านทาง ออนไลน์	real	4	6	9	1	2

ตารางที่ 3 ค่าความสัมพันธ์ของตัวแปรในข้อมูลสำหรับการสอนและทดสอบแบบจำลอง

	score_fake	score_real	sim_matched	domain_fake	domain_real	fake	Real	suspicious
score_fake	1.000	0.081	0.424	0.901	0.103	0.724	-0.377	-0.398
score_real	0.081	1.000	0.201	0.097	0.785	0.042	0.223	-0.266
sim_matched	0.424	0.201	1.000	0.440	0.236	0.360	0.303	-0.683
domain_fake	0.901	0.097	0.440	1.000	0.121	0.693	-0.363	-0.378
domain_real	0.103	0.785	0.236	0.121	1.000	0.056	0.129	-0.187
Fake	0.724	0.042	0.360	0.693	0.056	1.000	-0.529	-0.540
Real	-0.377	0.223	0.303	-0.363	0.129	-0.529	1.000	-0.428
suspicious	-0.398	-0.266	-0.683	-0.378	-0.187	-0.540	-0.428	1.000

ตารางที่ 3 แสดงค่าลักษณะเด่นที่แยกออกมาที่มีความสัมพันธ์กับเป้าหมายโดยลักษณะเด่น sim_matched มีค่าสหสัมพันธ์เชิงบวกกับทั้งข่าวปลอมและข่าวจริง ลักษณะเด่น score_fake และ domain_fake มีค่าสหสัมพันธ์เท่ากับ 0.724 และ 0.693 ตามลำดับ ซึ่งแสดงถึงความสามารถในการคาดการณ์ข่าวปลอม นอกจากนี้ลักษณะเด่น score_real และ domain_real มีค่าสหสัมพันธ์เชิงบวกกับข่าวจริง โดยเท่ากับ 0.223 และ 0.129 ตามลำดับ ในภาพรวมข้อมูลที่สกัดลักษณะเด่นจากงานวิจัยนี้สามารถนำไปสร้างการเรียนรู้ด้วยเครื่องเพื่อทำนายข่าวปลอมและข่าวจริงได้



รูปที่ 3 การพล็อตแบบกระจายแสดงความสัมพันธ์ของข้อมูล

เพื่อการมองข้อมูลเป็นภาพและวิเคราะห์ความสัมพันธ์ ผู้วิจัยใช้วิธีการแสดงเป็นการพล็อตแบบกระจายรูปที่ 3 แสดงการพล็อตข้อมูลแบบกระจาย เป็นการแสดงตำแหน่งของข้อมูลตามความสัมพันธ์ระหว่าง 2 ตัวแปร โดยหนึ่งตัวแปรในแนวนอนแต่ละแถว และแนวดิ่งเป็นอีกหนึ่งตัวแปรในแต่ละคอลัมน์ ในแต่ละภาพย่อยเป็นความสัมพันธ์ของ 2 ตัวแปร การพล็อตแบบกระจายทำให้พบข้อสังเกตได้ว่าข้อมูลสามารถจำแนกได้ง่ายในหลาย ๆ ภาพที่เป็นคู่ความสัมพันธ์ระหว่างลักษณะเด่น เป็นการยืนยันว่าข้อมูลที่ผ่านการประมวลผลสามารถนำไปสร้างการเรียนรู้ของเครื่องได้

ในการสร้างแบบจำลองการเรียนรู้ของเครื่องเพื่อจำแนกข่าวจริงและข่าวปลอม จะต้องจัดเตรียมข้อมูลจำนวนมากเพียงพอสำหรับสอนแบบจำลองการเรียนรู้ของเครื่อง งานวิจัยนี้ผู้วิจัยใช้ข้อมูลประเภทข้อความข่าวก่อนการนำไปสอนให้กับตัวแบบการเรียนรู้ของเครื่องต้องทำการประมวลผลคำด้วยเทคนิควิธีการประมวลผลภาษาธรรมชาติ ส่วนอัลกอริทึมการเรียนรู้ของเครื่องที่เลือกใช้มีหลายเทคนิค เช่น LR, KNN, MLP, SVM, NB, DT, RF และ LSTM เทคนิคเหล่านี้เป็นเทคนิคการเรียนรู้ของเครื่องชนิดมีผู้สอน โดยมีรายละเอียดดังนี้ 1) การเตรียมข้อมูล ดำเนินการหลังจากขั้นตอนการรวบรวมข้อมูลด้วยการสืบค้นสารสนเทศในเว็บด้วยเว็บเบราว์เซอร์ จากนั้นนำข้อมูลที่ได้อิงประมวลผลภาษาธรรมชาติ ได้คุณลักษณะเด่นอยู่ในรูปแบบที่สามารถวิเคราะห์ได้ด้วยเทคนิคการเรียนรู้ของเครื่อง การประมวลผลภาษาธรรมชาติ 2) การสร้างแบบจำลองหรือตัวแบบการเรียนรู้ของเครื่อง โดยทำการนำข้อมูลที่ผ่านการเตรียมข้อมูลแล้ว เข้าสู่การสร้างแบบจำลองโดยใช้เทคนิคการเรียนรู้แบบมีผู้สอน โดยการสอนแบบจำลองใช้ข้อมูล 53,220 เรคคอร์ด สอนด้วยวิธีการทดสอบแบบไขว้ 10 ชุดข้อมูล (10-Fold Cross Validate) 3) การทดสอบและประเมินประสิทธิภาพแบบจำลอง โดยนำชุดข้อมูลที่เครื่องคอมพิวเตอร์ยังไม่เคยเรียนรู้มาทำการทดสอบแบบจำลองว่ามีค่าประสิทธิภาพในการวิเคราะห์ โดยประเมินจาก ค่าความถูกต้อง (Accuracy) ความแม่นยำ (Precision) ค่าการเรียกคืน (Recall) และค่าถ่วงดุล (F-Measure) และ 4) การนำแบบจำลองที่ดีที่สุดไปประยุกต์ใช้งาน โดยการพัฒนาเป็นเว็บแอปพลิเคชันระบบตรวจข่าวปลอม

การพัฒนาระบบตรวจข่าวปลอม ผู้วิจัยเน้นการพัฒนาให้ทำงานบนคลาวด์เพื่อรองรับความพร้อมใช้งาน เน้นการเลือกใช้ซอฟต์แวร์แบบเปิดทั้งหมด เช่น ระบบปฏิบัติการ Ubuntu 20.04 LTS ใช้ Python เวอร์ชัน 3.8 เป็นภาษาในการพัฒนาเป็นหลัก ใช้เว็บเฟรมเวิร์กเป็น Django เวอร์ชัน 3.1.5 ใช้ฐานข้อมูลเป็นแบบ MongoDB เวอร์ชัน 4.4.6 สำหรับจัดเก็บข่าว และใช้ MySQL เวอร์ชัน 8.0 สำหรับจัดเก็บข้อมูลผู้ใช้งานและเมตาดาตาของระบบ สำหรับระบบสืบค้นข้อมูลใช้เทคโนโลยีที่สำคัญ ๆ ได้แก่ Selenium, Requests, Google, Chrome, และ Firefox สำหรับเครื่องมือประมวลผลภาษาธรรมชาติ Pythainlp เวอร์ชัน 2.3.1 (Phatthiyaphaibun et al., 2016) เครื่องมือในการสร้างการเรียนรู้ด้วยเครื่องประกอบด้วย Scikit-learn และ TensorFlow ด้านความปลอดภัยของระบบ ใช้เครื่องมือเพื่อรักษาความปลอดภัยและป้องกันการบุกรุก ได้แก่ Fail2ban, Uncomplicated Firewall และ ClamAV

ผลการดำเนินการวิจัย

การสร้างแบบจำลองการเรียนรู้ของเครื่องสำหรับตรวจจับข่าวปลอม ผู้วิจัยทำการเปรียบเทียบแบบจำลองเพื่อเลือกแบบจำลองที่ดีที่สุดในระบบตรวจจับข่าวปลอม ผู้วิจัยเลือกซอฟต์แวร์แบบเปิด ได้แก่ Python, Numpy, Pandas, Scikit-Learn, Tensorflow และเครื่องมือที่จำเป็นอื่น ๆ ในการสร้างระบบตรวจจับข่าวปลอม เครื่องมือวิเคราะห์ข้อมูลและการสร้างแบบจำลองการเรียนรู้ของเครื่อง ได้แก่ LR, MLP, DT, RF, SVM, NB และ KNN ในแพ็คเกจ Scikit-learn และ LSTM ใน TensorFlow ตัวชี้วัดประสิทธิภาพที่ใช้รวมถึงค่าความถูกต้อง (Accuracy) ค่าความแม่นยำ (Precision) ค่าการเรียกคืน (Recall) และค่าถ่วงดุล (F-Measure)

จากผลการวิจัยในการสร้างแบบจำลองการเรียนรู้ของเครื่องประกอบไปด้วย สำหรับเทคนิคการเรียนรู้ด้วยเครื่องที่ใช้ในการทดลองเปรียบเทียบในการวิจัยครั้งนี้ พบว่าทุกเทคนิคมีผลการทดสอบความสามารถในการจำแนกข่าวได้ดี ซึ่งมีค่าประสิทธิภาพที่สูงกว่า 90% ในทุกเทคนิค ยกเว้น NB ซึ่งมีการจำแนกข้อมูลทดสอบถูกต้องเพียง 78% เท่านั้น เมื่อพิจารณาถึงความถูกต้องการจำแนกข้อมูลเรียงลำดับจากสูงสุดไปต่ำสุด มีลำดับดังต่อไปนี้ DT (94.2%), RF (94.2%), SVM (94.0%), LSTM (93.81%), MLP (93.6%), LR (91.9%), KNN (91.5%) และ NB (78.2%) ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยของการเรียนรู้ของเครื่องชนิดต่าง ๆ แสดงดังตารางที่ 4

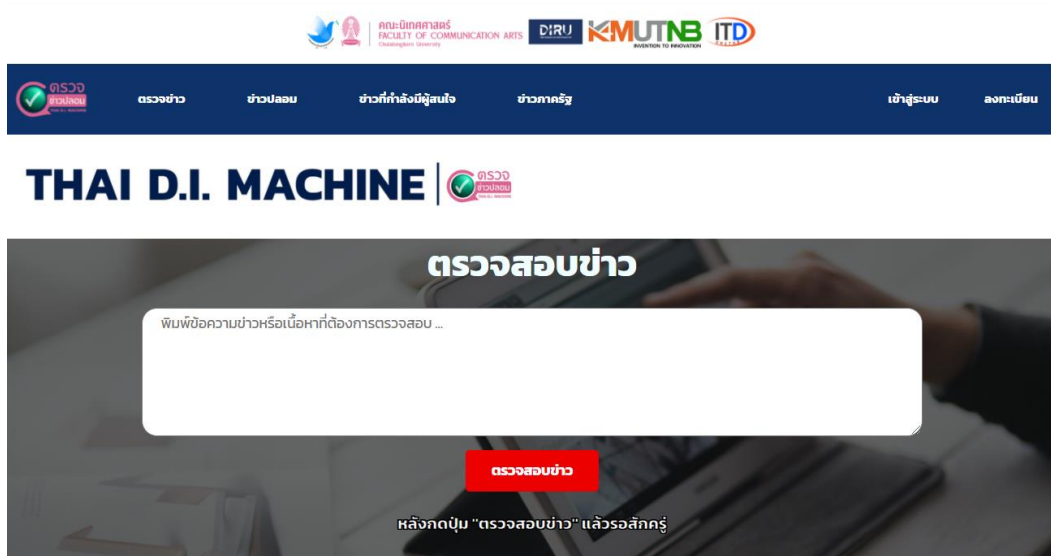
ตารางที่ 4 ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยของการเรียนรู้ของเครื่องชนิดต่าง ๆ

Metrics	LR	MLP	DT	RF	SVM	NB	KNN	LSTM
Accuracy	0.919±0.028	0.936±0.047	0.942±0.052	0.942±0.052	0.940±0.052	0.785±0.050	0.915±0.032	93.81±0.006
Precision	0.923±0.026	0.948±0.036	0.952±0.037	0.953±0.037	0.950±0.038	0.814±0.041	0.925±0.027	94.44±0.006
Recall	0.919±0.028	0.936±0.047	0.942±0.052	0.942±0.052	0.940±0.052	0.785±0.050	0.915±0.032	93.81±0.006
F-Measure	0.919±0.028	0.939±0.055	0.940±0.055	0.940±0.055	0.938±0.055	0.776±0.061	0.914±0.033	93.78±0.006

เมื่อวิเคราะห์ถึงเหตุผลในการจำแนกข้อมูลของแต่ละเทคนิค เริ่มจากเทคนิค สำหรับ 5 เทคนิคที่มีประสิทธิภาพสูงสุดด้านความถูกต้อง ตั้งแต่ 93.81% ถึง 94.2% ได้แก่ MLP, LSTM, SVM, RF และ DT ซึ่งมีความแตกต่างกันไม่มาก MLP เป็นโครงข่ายประสาทเทียมแบบหลายชั้นมีความเป็นไปได้ที่จะปรับปรุงให้มีประสิทธิภาพสูงขึ้นด้วยการเพิ่มจำนวนชั้นและพารามิเตอร์แต่จะเกิดปัญหาการจำข้อมูลสอนมากเกินไป (Overfitting) ได้เช่นกัน ส่วน LSTM เป็นเทคนิคโครงข่ายประสาทเทียมประเภทการเรียนรู้เชิงลึกซึ่งมีจำนวนพารามิเตอร์จำนวนมาก จึงมีโอกาสสูงในการปรับปรุงพารามิเตอร์ให้มีขอบเขตข้อมูลได้ดีและมีความคงเส้นคงวาทนทานต่อข้อมูลแปลกปลอมได้ดีที่สุดโดยมีส่วนเบี่ยงเบนมาตรฐานต่ำที่สุด ในขณะที่ SVM นอกจากจะมีประสิทธิภาพสูงยังมีความทนทานต่อข้อมูลรบกวนและข้อมูลแปลกปลอมได้ดีเช่นกัน สำหรับ RF ซึ่งได้ผลเทียบเท่ากับ DT แต่ RF ใช้ทรัพยากรที่มากกว่า ทั้งนี้ RF อาศัยหลักการจำแนกข้อมูลด้วยชุดต้นไม้ตัดสินใจจำนวนมาก เป็นการเพิ่มจำนวนแบบจำลองย่อยร่วมกันตัดสินใจให้มีความถูกต้องในการจำแนกได้ดียิ่งขึ้น ในส่วน KNN เป็นการจำแนกข้อมูลโดยไม่ต้องมีการสร้างแบบจำลองโดยตรงแต่จะใช้ข้อมูลชุดสอนทั้งหมดเป็นแบบจำลอง และจำแนกข้อมูลด้วยการวัดความคล้ายคลึงระหว่างข้อมูลที่จะจำแนก เทียบกับข้อมูลชุดสอน ด้วยการกำหนดค่า $K=5$ เป็นการเลือกเรคคอร์ดที่ใกล้เคียงที่สุด 5

เรคคอร์ดเพื่อโหวตผลการจำแนก KNN เองก็ถือว่าเป็นที่นิยมเพราะง่ายในการสร้าง ให้ผลลัพธ์เป็นที่ยอมรับได้เช่นกัน สำหรับ LR เป็นเทคนิคการถดถอยที่สามารถทำงานสำหรับจำแนกข้อมูลเนื่องจากใช้ฟังก์ชันลอจิสติกมอดเป็นแปลงเอาต์พุตเป็น 0 และ 1 แทนที่จะเป็นค่าต่อเนื่อง เทคนิค LR ได้ผลการจำแนกสูงกว่า 90% เช่นกันในงานวิจัยครั้งนี้ สำหรับ NB ซึ่งเป็นเทคนิคที่ใช้หลักการความน่าจะเป็น ซึ่งเป็นเทคนิคเดียวที่มีประสิทธิภาพต่ำกว่า 80% ทั้งนี้ในงานวิจัยส่วนใหญ่พบว่า NB จะเป็นเทคนิคที่เป็นรองเทคนิคอื่น ๆ จึงนิยมใช้เป็นตัวหลักในการเปรียบเทียบ ซึ่งเทคนิคใหม่ ๆ คาดหวังว่าจะต้องดีกว่า NB

เมื่อพิจารณาถึงเหตุผลของการจัดเตรียมข้อมูล ถือว่าเป็นสิ่งสำคัญก่อนป้อนข้อมูลเข้าสู่แบบจำลองการเรียนรู้ของเครื่อง หากการจัดเตรียมข้อมูลได้ไม่ดี ก็ไม่อาจสร้างแบบจำลองการเรียนรู้ของเครื่องในการจำแนกข้อมูลที่มีประสิทธิภาพสูงได้ ซึ่งในงานวิจัยครั้งนี้ใช้วิธีการสกัดลักษณะเด่นของคำที่มักปรากฏในข่าวจริง และข่าวปลอมโดยการนับความถี่ของคำเหล่านั้นเพื่อใช้เป็นข้อมูลประจำของแต่ละข่าวที่สืบค้น อีกหนึ่งลักษณะเด่นคือค่าความคล้ายของข่าวสืบค้นกับเนื้อหาในแหล่งข่าวก็เป็นปัจจัยหลักทำให้เกิดอำนาจจำแนกข้อมูลได้ดี ประกอบกับการใช้จำนวนโดเมนข่าวที่เผยแพร่ข่าวปลอมและโดเมนข่าวที่เผยแพร่ข่าวจริง ก็เป็นอีก 2 ปัจจัย ที่ช่วยให้การจำแนกข่าวได้ดีขึ้น ข้อมูลที่ได้ทำให้เกิดการจำแนกได้เป็นอย่างดีเห็นได้ชัดจากการพล็อตค่าสหสัมพันธ์และการวิเคราะห์ความสัมพันธ์ระหว่างคุณลักษณะกับประเภทข่าว ทั้ง 5 ลักษณะเด่นที่ผู้วิจัยใช้ส่งผลให้การเรียนรู้ของเรียนรู้ข้อมูลได้อย่างมีประสิทธิภาพดังจะเห็นจากหลายแบบจำลองมีความถูกต้องสูงกว่า 90% สำหรับงานวิจัยนี้เลือก LSTM สำหรับเป็นตัวจำแนกข้อมูลในแอปพลิเคชันเนื่องจากโมเดลอยู่ในกลุ่มได้ผลทดลองด้านประสิทธิภาพสูงสุดและมีความเสถียรสูงสุดจากการวัดประสิทธิภาพแบบไขว้ 10 ชุดข้อมูล



รูปที่ 4 ระบบตรวจข่าวปลอมด้วยการเรียนรู้ของเครื่อง <https://thaidimachine.org>

ผลการพัฒนาระบบตรวจสอบข่าวปลอม ผู้วิจัยได้ทำการพัฒนาและติดตั้งระบบบนคลาวด์ที่ประชาชนสามารถเข้าใช้งานได้ที่ <https://thaidimachine.org> ผลการประเมินระบบโดยผู้ใช้งานจำนวน 34 คน แบ่งการประเมินออกเป็น 3 ด้าน ได้แก่ 1. ด้านการทำงานได้ตามฟังก์ชันของระบบ (Functional Test) 2. ด้านความง่ายต่อการใช้งาน (Usability Test) และ 3. ด้านความปลอดภัยของระบบ (Security Test) เว็บบระบบตรวจสอบข่าวปลอมสามารถใช้งานได้ง่าย เมื่อประชาชนต้องการตรวจสอบข่าว ทำได้โดยการพิมพ์ข้อความหรือก๊อปปี้ข่าวลงในกล่องข้อมูลแล้วกดตรวจสอบข่าว ระบบจะทำการประมวลผลแล้วแสดงผลการวิเคราะห์ข่าว โดยแสดงเป็นผลลัพธ์เป็น “ข่าวจริง” “ข่าวปลอม” หรือ “ข่าวน่าสงสัย” พร้อมมีลิงค์ข่าวที่เกี่ยวข้องให้ผู้ตรวจข่าวมีข้อมูลเพิ่มเติมเกี่ยวกับข่าว นอกจากการวิเคราะห์ข่าว ระบบทำการรวบรวมข่าวปลอม ข่าวที่กำลังมีผู้สนใจ และข่าวภาครัฐ ระบบเก็บรวบรวมลิงค์จากเว็บไซต์ต่าง ๆ ไว้โดยอัตโนมัติ เพื่อความสะดวกกับผู้ใช้จะได้ตรวจสอบผ่านเว็บที่น่าเชื่อถือ ด้านความปลอดภัยของเว็บแอปพลิเคชัน ระบบมีฟังก์ชันการล็อกอิน/ลงทะเบียน สำหรับผู้ใช้ที่ประสงค์จะป้อนกลับข้อมูลข่าว ต้องมีการลงทะเบียนไว้ด้วยเพื่อระบุตัวตน ให้มีความมั่นใจในการให้ข้อคิดเห็น และระบบมีการป้องกันการบุกรุกโดยใช้ซอฟต์แวร์แบบเปิด ได้แก่ Fail2ban ป้องกันเซิร์ฟเวอร์ มีไฟร์วอลล์แบบ Uncomplicated Firewall (ufw) กำหนดกฎให้ป้องกัน Hackers และ Bad Bots ตัวอย่างหน้าจอแรกของเว็บตรวจข่าวปลอมแสดงดังรูปที่ 4

หลังจากเปิดการใช้งานผู้วิจัยติดตามผลการใช้งาน พบว่ามีการสืบค้นเฉลี่ยมากกว่าวันละ 1,000 ข้อความที่ได้รับการสืบค้น จากการวิเคราะห์พบว่าจากการสุ่มถามผู้ใช้งานระบบตรวจข่าวปลอม จำนวน 34 คน มีความคิดเห็นเกี่ยวกับในระบบที่ระดับความคิดเห็นในทุกด้านทุกรายการประเมินในระดับ ดีมาก ด้วยค่าเฉลี่ยทุกด้านทุกรายการประเมิน มีค่าเท่ากับ 4.53 และส่วนเบี่ยงเบนมาตรฐาน 0.63 ผลประเมินด้านความปลอดภัยของระบบอยู่ในระดับต่ำสุด มีค่าเฉลี่ยเท่ากับ 4.46 ค่าเบี่ยงเบนมาตรฐานเท่ากับ 0.70 อยู่ในระดับดี เมื่อพิจารณาผลการประเมินด้านการทำงานได้ตามฟังก์ชันของระบบ มีค่าเฉลี่ย 4.52 และส่วนเบี่ยงเบนมาตรฐานเท่ากับ 0.54 อยู่ในระดับดีมาก ผลการประเมินด้านความง่ายต่อการใช้งาน มีค่าเฉลี่ยเท่ากับ 4.59 ส่วนเบี่ยงเบนมาตรฐาน 0.55 อยู่ในระดับดีมาก ผู้ประเมินซึ่งส่วนใหญ่จบการศึกษาสาขาที่เกี่ยวข้องกับคอมพิวเตอร์และเทคโนโลยีสารสนเทศ และมีประสบการณ์การทำงานด้านคอมพิวเตอร์และเทคโนโลยีสารสนเทศโดยตรง ให้ความคิดเห็นในภาพรวมว่าระบบที่ผู้วิจัยพัฒนาขึ้นนี้มีประสิทธิภาพดีมาก มีการจัดรูปแบบได้สวยงาม มีการใช้ฟอนต์ขนาดพอดี และเป็นแอปพลิเคชันที่มีประโยชน์และสามารถใช้งานได้จริง

สรุป

จากปัญหาข่าวปลอมที่มีจำนวนมากในสื่อออนไลน์ ทำให้มีผู้หลงเชื่อข่าวปลอมว่าเป็นข่าวจริง และนำไปแชร์ต่อทำให้เกิดปัญหาอื่นๆ ตามมา ปัจจุบันยังไม่มีเครื่องมือที่จะช่วยตรวจสอบข่าวปลอมภาษาไทย ที่จะช่วยชะลอการแพร่กระจายข่าว การตรวจจับข่าวปลอมเป็นงานที่ยากเนื่องจากข่าวมีความเคลื่อนไหวอย่างมาก งานวิจัยนี้เสนอวิธีการใหม่ที่มีประสิทธิภาพในการจัดการกับข่าวปลอมหรือข้อมูลที่ผิด ผู้วิจัยจึงได้นำเสนอระบบตรวจสอบข่าว

ปลอมภาษาไทยที่ทำงานบนพื้นฐานเทคนิคทางปัญญาประดิษฐ์ โดยเครื่องประมวลผลตรวจสอบข่าวบนคลาวด์และส่งผลการวิเคราะห์ข่าวให้ผู้ใช้งานทราบ

ผู้วิจัยใช้เว็บเบราว์เซอร์รวบรวมข้อมูลสำหรับตัวอย่าง 41,448 ตัวอย่าง และทำการจัดประเภทไว้ล่วงหน้าเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย จำนวนตัวอย่างข้อมูลในแต่ละกลุ่มมีความสมดุลโดยสุ่มเพิ่มเป็น 53,220 ตัวอย่าง ผู้วิจัยใช้วิธีการสอนแบบไขว้ทดสอบจำนวน 10 ชุดข้อมูล (10-Fold Cross Validation) แบบจำลองการเรียนรู้ของเครื่องซึ่งเป็นศาสตร์ด้านปัญญาประดิษฐ์ที่ใช้ในการศึกษา ได้แก่ การถดถอยโลจิสติกส์ (LR) เพื่อนบ้านใกล้เคียง (KNN) นาอิฟเบย์ (NB) โครงข่ายประสาทเทียมเพอร์เซ็ปตรอนหลายชั้น (MLP) ซัพพอร์ตเวกเตอร์ (SVM) ป่าสุ่ม (RF) และโครงข่ายประสาทเทียมชนิดหน่วยความจำระยะสั้นแบบยาว (LSTM) ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยของการเรียนรู้ของเครื่องชนิดต่าง ๆ กลุ่มโมเดลที่พบว่าดีที่สุดในงานวิจัยนี้ได้แก่ DT, RF, SVM, LSTM และ MLP ซึ่งเหมาะสำหรับการเรียนรู้ของเครื่องที่จะนำไปใช้เป็นระบบตรวจข่าวปลอม

ระบบตรวจสอบข่าวปลอมบนพื้นฐานด้านปัญญาประดิษฐ์มีการทำงานประกอบด้วย 3 องค์ประกอบหลัก ได้แก่ ส่วนการสืบค้นข้อมูลข่าวบนเว็บ ส่วนการประมวลผลภาษาธรรมชาติ และส่วนการเรียนรู้ของเครื่อง ส่วนแรกที่เป็นสืบค้นข่าวบนเว็บทำหน้าที่สืบค้นข่าวที่ตรงกับข่าวสืบค้นซึ่งป้อนเข้าระบบโดยผู้ใช้งานที่ประสงค์จะตรวจข่าว ส่วนที่สองจะรับข้อมูลข่าวที่สืบค้นได้จากส่วนแรกและนำข้อมูลข่าวไปประมวลผล โดยทำการตัดคำและสกัดคำที่เป็นลักษณะเด่นของข่าวจริงข่าวปลอม และส่วนที่สามเป็นการเรียนรู้ของเครื่องทำหน้าที่จำแนกข้อมูลที่ได้รับจากส่วนที่สองและแสดงผลเป็น 3 ประเภท ได้แก่ ข่าวจริง ข่าวปลอม และข่าวน่าสงสัย

ระบบจะตรวจข่าวจากการสืบค้นข่าวเผยแพร่ในอินเทอร์เน็ตตามสกัดคำฟิเจอร์หรือลักษณะเด่นของข่าวก่อนส่งไปตัดสินใจ หากข้อความข่าวไม่มีการแชร์หรือเผยแพร่มากนักทำให้ผลการสกัดคำฟิเจอร์ไปตกอยู่ที่กลุ่มน่าสงสัย แต่หากข่าวมีการเผยแพร่จำนวนมากและมีกลุ่มคำที่มักปรากฏในข่าวปลอม เมื่อสกัดคำฟิเจอร์ส่งผลการตัดสินใจได้เป็นข่าวปลอม และมีการแชร์และเผยแพร่เป็นจำนวนมากและมีกลุ่มคำที่มักปรากฏในข่าวจริง เมื่อสกัดคำฟิเจอร์ส่งผลการตัดสินใจว่าเป็นข่าวจริง

ข้อจำกัดของระบบคือระบบยังไม่เข้าใจเนื้อหาข่าวทั้งหมดแบบกว้างอย่างที่มีมนุษย์เข้าใจ ในการวิจัยต่อไปอนาคต อาจจะต้องประยุกต์ใช้การเรียนรู้เชิงลึกประเภท BERT (Bidirectional Encoder Representations from Transformers) เพื่อให้มีความสามารถในการเข้าใจภาษาธรรมชาติ (Natural Language Understanding) ทำอย่างไรให้เครื่องเข้าใจข่าวมากขึ้นเหมือนมนุษย์ เครื่องสามารถทำนายประเภทของข่าวและอธิบายว่าทำไมจึงให้คำตอบหรือการตอบสนองดังกล่าว นอกจากนี้ หากเนื้อหาข่าวประกอบด้วยข้อความ เสียง และวิดีโอ เครื่องต้องวิเคราะห์และตอบสนองอย่างถูกต้องและมีการเสนอข่าวใกล้เคียงความสอดคล้องกันกับข่าวสืบค้นด้วยซ้ำ

กิตติกรรมประกาศ

งานวิจัยนี้อยู่ภายใต้ความร่วมมือระหว่างหน่วยปฏิบัติการวิจัยและพัฒนาทรัพยากรมนุษย์ด้านความรู้ทางดิจิทัลและการรู้เท่าทันสื่อ (DIRU) คณะนิเทศศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย และคณะเทคโนโลยีสารสนเทศและนวัตกรรมดิจิทัล มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ (KMUTNB) บทความนี้เป็นส่วนหนึ่งของโครงการวิจัยเรื่อง “การพัฒนาการตรวจจับข่าวปลอมโดยการเรียนรู้ของเครื่องและการตรวจสอบข้อเท็จจริงของประชาชน” ได้รับการสนับสนุนเงินทุนวิจัยจากกองทุนพัฒนาสื่อปลอดภัยและสร้างสรรค์ ประจำปีงบประมาณ 2562 ได้รับรหัสโครงการ: 2562-T3/0064

เอกสารอ้างอิง

- Ayoub, J., Yang, X. J., & Zhou, F. (2021). Combat COVID-19 infodemic using explainable natural language processing models. *Information Processing and Management*, 58(4). <https://doi.org/10.1016/j.ipm.2021.102569>
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273–297. <https://doi.org/10.1007/BF00994018>
- Gori, N. (2018). *Machine Learning: A constraint-based approach*. Morgan Kaufmann.
- Guo, G., Wang, H., Bell, D., Bi, Y., & Greer, K. (2003). KNN Model-Based Approach in Classification. In R. Meersman, R., Tari, Z., & Schmidt, D. C. (Eds.), *On The Move to Meaningful Internet Systems 2003: CoopIS, DOA, and ODBASE* (Vol. 2888, pp. 986–996). Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-540-39964-3_62
- Hagan, M. T., Demuth, H. B., & Beale, M. H. (2014). *Neural network design*. Martin Hagan.
- Hochreiter, S., & Schmidhuber, J. (1997). Long Short-Term Memory. *Neural Computation*, 9(8), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- Kleinbaum, D. G., & Klein, M. (2010). *Logistic Regression: A Self-Learning Text* (3rd ed.). Springer.
- Krallinger, M., Rabal, O., Lourenço, A., Oyarzabal, J., & Valencia, A. (2017). Information Retrieval and Text Mining Technologies for Chemistry. *Chemical Reviews*, 117(12), 7673–7761. <https://doi.org/10.1021/acs.chemrev.6b00851>
- Myles, A. J., Feudale, R. N., Liu, Y., Woody, N. A., & Brown, S. D. (2004). An introduction to decision tree modeling. *Journal of Chemometrics*, 18(6), 275–285. <https://doi.org/10.1002/cem.873>

- Phatthiyaphaibun, W., Chaovavanich, K., Polpanumas, C., Suriyawongkul, A., Lowphansirikul, L., Chormai, P., smeeklai, Tanruangporn, P. P., Charoenchainetr, P., Udomcharoenchaikit, C., Supaseth, Janthong, A., Chaovavanich, K., Korkeat W., Jitchiranant, N., Ninyawee, N., boomsquared, Dubois, Y., nyamakawa, Somsontorn, C., Cody, Prakalphakul, K., Yuankrathok, P., Badger, C., fossabot, hopedataannotations, & pontakornth. (2016). *PyThaiNLP: Thai Natural Language Processing in Python*. <https://doi.org/10.5281/zenodo.4662045>
- Svetnik, V., Liaw, A., Tong, C., Culberson, J. C., Sheridan, R. P., & Feuston, B. P. (2003). Random Forest: A Classification and Regression Tool for Compound Classification and QSAR Modeling. *Journal of Chemical Information and Computer Sciences*, 43(6), 1947–1958. <https://doi.org/10.1021/ci034160g>
- Yager, R. R. (2006). An extension of the naive Bayesian classifier. *Information Sciences*, 176(5), 577–588. <https://doi.org/10.1016/j.ins.2004.12.006>

Research Article

การประเมินแบบจำลองความสูงของต้นข้าวด้วยภาพถ่ายจากอากาศยานไร้คนขับ (UAV)

Rice canopy height model assessment using unmanned aerial vehicle (UAV) images

ปริญัต จุนถาวร¹ และ กฤษณ อินทรรัตน์^{1*}

Parinchat Juntaworn¹ and Kritchayan Intarat^{1*}

¹คณะศิลปศาสตร์ สาขาวิชาภูมิศาสตร์ มหาวิทยาลัยธรรมศาสตร์ อำเภอกองหลวง จังหวัดปทุมธานี 12121 ประเทศไทย

¹Faculty of Liberal Arts, Department of Geography, Thammasat University, Khlong Luang, Pathumthani 12121 Thailand

*E-mail: intaratt@tu.ac.th

Received: 10/06/2021; Revised: 20/08/2021; Accepted: 23/09/2021

บทคัดย่อ

การเกษตรแม่นยำถูกนำมาใช้ด้านการเกษตรมากขึ้นในประเทศไทยโดยเฉพาะการปลูกข้าว ข้อมูลที่ต้องแม่นยำเป็นปัจจัยสำคัญในการจัดการนาข้าวอย่างมีประสิทธิภาพ นอกจากข้อมูลดัชนีพืชพรรณแล้ว ข้อมูลความสูงของพืชมีความสำคัญอย่างมากในการตรวจสอบการเจริญเติบโตและประเมินผลผลิต อย่างไรก็ตามการวัดความสูงในแปลงโดยตรงอาจส่งผลกระทบต่อพืชและใช้เวลานาน การทดลองนี้ได้ใช้ความสามารถของอากาศยานไร้คนขับเพื่อสร้างแบบจำลองความสูงเรือนยอด (CHM) ดำเนินการเก็บตัวอย่างจำนวน 108 ตัวอย่างในระยะข้าวแตกกอและระยะข้าวตั้งท้องซึ่งเป็นระยะที่แสดงการเจริญเติบโตทางกายภาพได้ชัดเจน เช่น การเจริญเติบโตทางลำต้นและการเจริญเติบโตด้านการสืบพันธุ์ตามลำดับ CHM ของแต่ละพื้นที่ตัวอย่างถูกกำหนดด้วยค่าเปอร์เซ็นต์ไทล์ (Pr) ที่ 91 ถึง 99 ผลการทดลองพบว่า CHM และความสูงจากแปลงตัวอย่างแสดงความสัมพันธ์เชิงเส้นสูงที่สุดในทั้งสองระยะการเจริญเติบโตที่ Pr₉₉ ($r^2 = 0.87$ และ 0.86 ตามลำดับ) ค่า RMSE ของความสูงพืชมีค่า 1.83 และ 1.71 เซนติเมตรตามลำดับที่นัยสำคัญ $p < 0.01$ ซึ่งผลการทดลองสรุปได้ว่า UAV สามารถสร้างความสูงได้อย่างแม่นยำโดยไม่ต้องลงเก็บข้อมูลภาคสนาม ดังนั้นวิธีนี้จึงเป็นวิธีการที่ดีในการได้มาซึ่งข้อมูลทางกายภาพของต้นข้าวและช่วยป้องกันเรื่องความเสียหายของพืช

คำสำคัญ: แบบจำลองความสูง, อากาศยานไร้คนขับ, นาข้าว

Abstract

Precision agriculture implementation is currently increasing in Thailand, especially in rice cultivation. Accurate information is an essential factor in efficient rice field management. Apart from the vegetation indices, plant height (PH) is critical for crop monitoring and yield prediction. However, field height measurements could cause crop damage and take much time. This study utilizes the capability of an unmanned aerial vehicle (UAV) to construct the canopy height model (CHM) remotely. The samples are collected in the tillering and booting stage for 108 plots, representing the physical extension (e.g., rice elongation and reproduction, respectively). The CHM of each plot is determined at 91 to 99 percentiles. The result revealed that the best linear relationship between the CHM and field PH in both stages was found at the 99th percentile ($r^2 = 0.87$ and 0.86 , respectively). The RMS error of plant height was reported 1.83 and 1.71 cm from the prediction ($p < 0.01$). This experiment authenticates that the UAV could accurately obtain the PH without invading the field. Therefore, it is a good practice for achieving rice physical information and preventing crop destruction.

Keywords: canopy height model, UAV, paddy field

บทนำ

การศึกษาตัวแปรที่มีอิทธิพลต่อการเจริญเติบโตของพืชทางด้านเกษตรกรรมเป็นเป้าหมายหลักเพื่อให้ได้มาซึ่งผลผลิตในปริมาณสูงภายใต้การจัดการอย่างมีประสิทธิภาพ (De Souza et al., 2017) ปัญหาการทำการเกษตรกรรมส่วนใหญ่จะเกิดขึ้นในระหว่างขั้นตอนการปลูก (รัฐพงษ์ ม่วงประโคน และคณะ, 2564) จึงจำเป็นต้องมีการจัดการและประเมินผลผลิตพืชตั้งแต่เริ่มระยะการเจริญเติบโต (Grenzdörffer, 2014; Suphan et al., 2019) เพื่อแก้ไขปัญหาให้ทันทั่วทั้งที่และถูกต้องแม่นยำ การติดตามการเจริญเติบโตของพืชทำได้จากการตรวจสอบลักษณะชีพลักษณะที่สามารถบ่งบอกลักษณะการเติบโตของพืชได้ดีที่สุด (Großkinsky et al., 2015) แต่ปัจจุบันยังขาดวิธีการติดตามข้อมูลชีพลักษณะในแปลงผลผลิตที่ดีและแม่นยำ ด้วยข้อจำกัดด้านต้นทุนและความสามารถในการเก็บข้อมูลที่ส่งผลให้การประเมินผลผลิตมีประสิทธิภาพน้อยลง (รัฐพงษ์ ม่วงประโคน และคณะ, 2564; พรหมชัย สุพรรณและคณะ, 2561)

การวัดความสูงของพืชเป็นตัวแปรหนึ่งที่สามารถสะท้อนถึงลักษณะชีพลักษณะและสามารถใช้ติดตามการเจริญเติบโต การประมาณค่าชีวมวลเหนือพื้นดิน และประเมินผลผลิตตามกระบวนการเกษตรแม่นยำ (Bendig et al., 2015) ที่ได้รับความนิยมอย่างแพร่หลาย (Feng A. et al., 2019; Wang et al., 2018; Yang et al., 2020; Zhang et al., 2018) โดยทั่วไปการวัดความสูงจะส่มวัดด้วยไม้วัดโดยเลือกจากพืชบางต้นต่อแปลงที่วางแผนการเก็บอย่างเป็นรูปแบบสำหรับใช้เป็นตัวแทนความสูงของแปลง (Kawamura et al., 2020) การทำงานด้านเกษตรกรรมส่วนมากมัก

ใช้พื้นที่ดำเนินงานขนาดใหญ่ การเก็บข้อมูลโดยใช้แรงงานตามที่กล่าวไปอาจกระทบต่อประสิทธิภาพและก่อให้เกิดความเสียหายต่อพืชภายในแปลงเป็นผลให้ประสิทธิภาพการประเมินผลผลิตลดน้อยลง (Yang et al., 2020; รัฐพงษ์ ม่วงประโคน และคณะ, 2564; พรหมชัย สุพรรณ และคณะ, 2561)

ในปัจจุบันการวัดความสูงเรือนยอดของพืชได้รับการพัฒนาให้สามารถวัดค่าได้ด้วยเทคนิคการสแกนด้วยเลเซอร์ ได้แก่ เทคโนโลยีไลดาร์หรือเครื่องสแกนเลเซอร์ภาคพื้นดิน (Bareth et al., 2016; Grenzdörffer, 2014; Zhang et al., 2018; Zhou et al., 2020) อย่างไรก็ตามวิธีดังกล่าวใช้ต้นทุนค่อนข้างสูง การสำรวจระยะไกลด้วยอากาศยานไร้คนขับ (UAV) ในปัจจุบันถูกพัฒนาให้สามารถวิเคราะห์และสร้างแบบจำลองสามมิติ (3D model) ของพื้นที่วิธีการดังกล่าวสามารถนำไปใช้หาความสูงของพืช (Begashaw, 2018) และเป็นการวัดความสูงที่ไม่ก่อให้เกิดความเสียหายแก่พืชและมีประสิทธิภาพสำหรับการสำรวจพื้นที่ขนาดใหญ่ (Kawamura et al., 2020) ส่งผลให้การประเมินผลผลิตทางการเกษตรมีประสิทธิภาพเพิ่มขึ้นและช่วยลดเวลาในการทำงานให้น้อยลง

ข้อมูลจาก UAV ถูกนำมาประยุกต์กับงานด้านเกษตรกรรมได้หลากหลายรูปแบบทั้งการประมาณค่าชีวมวลของข้าว การตรวจวัดความสูง การทำนายหรือประมาณค่าผลผลิต และการติดตามการเจริญเติบโตของข้าวทั้งในประเทศ (พลพจน์ เอี่ยมสอาด และคณะ, 2563; รัฐพงษ์ ม่วงประโคน และคณะ, 2564; Suphan et al., 2019) และต่างประเทศ (Araus & Cairns, 2014; Bareth et al., 2016; Bendig et al., 2014; Bendig et al., 2015; Feng A. et al., 2019; Grenzdörffer, 2014; Iersel et al., 2017; Iizuka et al., 2017; Jones et al., 2020; Kawamura et al., 2020; Lu et al., 2021; Matese et al., 2016; De Souza et al., 2017; Tirado et al., 2019; Wang et al., 2018; Xie & Yang, 2020; Yang et al., 2020; Zhang et al., 2018; Zhou et al., 2020) Iersel et al. (2017) ได้ทดลองตรวจสอบความสูงและความเขียวของพืชสัมฤทธิ์ด้วยอนุกรมเวลาโดยใช้ข้อมูลจาก UAV งานทดลองดังกล่าวได้ศึกษาความสูงและค่าความเขียวของพืช จากนั้นนำมาสร้างแผนที่การเปลี่ยนแปลงความสูงของพืชสัมฤทธิ์ให้มีความถี่ของการเปลี่ยนแปลงมากขึ้น ผลการศึกษาพบว่าความสูงที่วัดได้แต่ละช่วงเวลามีความแม่นยำที่สุดในฤดูหนาวเนื่องจากสภาพแวดล้อมที่มีความผันแปรต่ำ Kawamura et al. (2020) ได้วิเคราะห์แบบจำลองความสูงของพืชโดยเปรียบเทียบความสูงพืชที่ได้จากภาพถ่าย UAV และความสูงที่ได้จากการวัดค่าในพื้นที่ งานดังกล่าวได้ข้อสรุปถึงประสิทธิภาพในการนำข้อมูลความสูงจาก UAV มาใช้ติดตามการเจริญเติบโตได้ จึงเป็นที่น่าสนใจในการสร้างข้อมูลความสูงจาก UAV เพื่อใช้ในการติดตามการเจริญเติบโตของข้าวแทนการลงสำรวจภาคสนาม

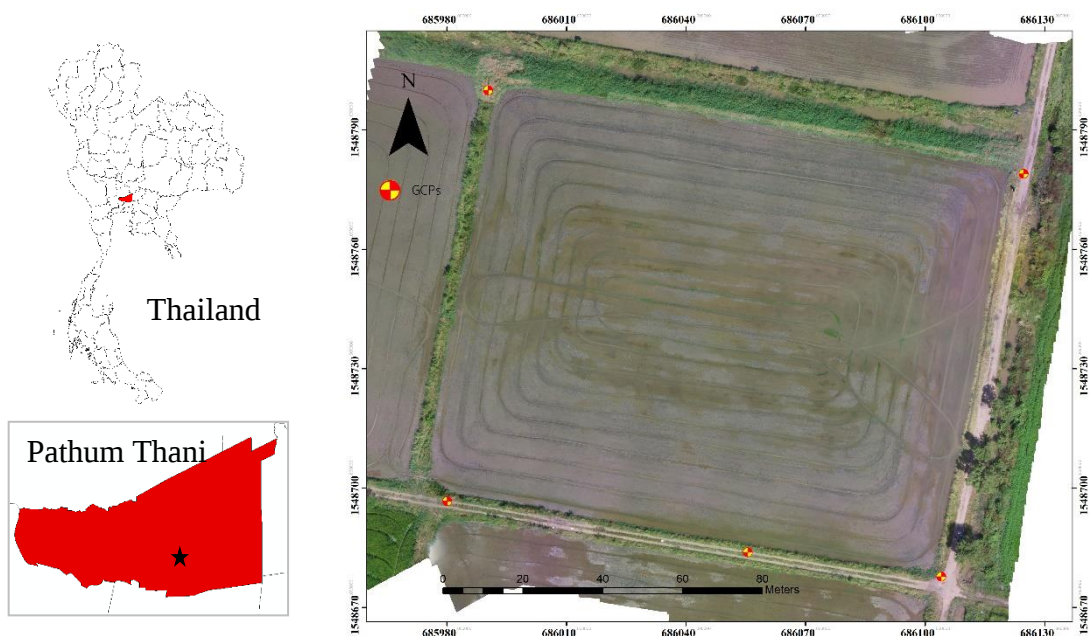
การทดลองนี้มุ่งเน้นไปที่การประเมินประสิทธิภาพของ UAV สำหรับการประมาณค่าความสูงต้นข้าวด้วยแบบจำลองความสูงเรือนยอด (CHM) บนสมมติฐานที่ว่า UAV เกรดตามท้องตลาด (commercial-grade UAV) สามารถนำมาใช้ตรวจสอบความสูงของต้นข้าวได้ การทดลองแบ่งการเก็บข้อมูลในแปลงปลูกข้าวออกเป็น 2 ระยะการเจริญเติบโต ได้แก่ ระยะแตกกอและระยะตั้งท้องเนื่องจากระยะการเจริญเติบโตทั้งสองระยะนี้เป็นช่วงที่สำคัญต่อปริมาณผลผลิตของข้าว (Kawamura et al., 2020) การวิเคราะห์ความสูงใช้ CHM โดยคาดหวังว่าการทดลองนี้จะสามารถให้ข้อสรุปถึงความสัมพันธ์ระหว่างความสูงของข้าวที่ได้จาก UAV และการเก็บข้อมูลในแปลง นอกจากนี้

การทดลองยังช่วยให้สามารถประเมินประสิทธิภาพของ UAV ในการเก็บข้อมูลความสูงของดินข้าวเพื่อเป็นแนวทางในการพัฒนาระบบการทำงานร่วมกับวิธีเกษตรแม่นยำ

วิธีการดำเนินงานวิจัย

1. พื้นที่ศึกษา

ข้าวเป็นพืชที่มีความสำคัญอันดับ 1 ใน 3 ของโลก (Yang et al., 2020) และเป็นพืชเศรษฐกิจที่สำคัญอันดับ 2 ของประเทศไทย การทดลองนี้ทำในพื้นที่ศูนย์วิจัยข้าวปทุมธานี ตำบลรังสิต อำเภอรัญบุรี จังหวัดปทุมธานี (ละติจูดที่ $13^{\circ}59'50.2''$ เหนือ และลองจิจูดที่ $100^{\circ}40'00.2''$ ตะวันออก) ทำการทดลองในแปลงนาข้าวขนาด 8 ไร่ปลูกข้าวเจ้าพันธุ์ปทุมธานี 1 (Pathum Thani 1) ดำเนินการปลูกแบบนาหว่าน เริ่มปลูกวันที่ 16 พฤศจิกายน พ.ศ. 2563 และเก็บเกี่ยววันที่ 13 มีนาคม พ.ศ. 2564 ได้แสดงพื้นที่ศึกษาในรูปที่ 1

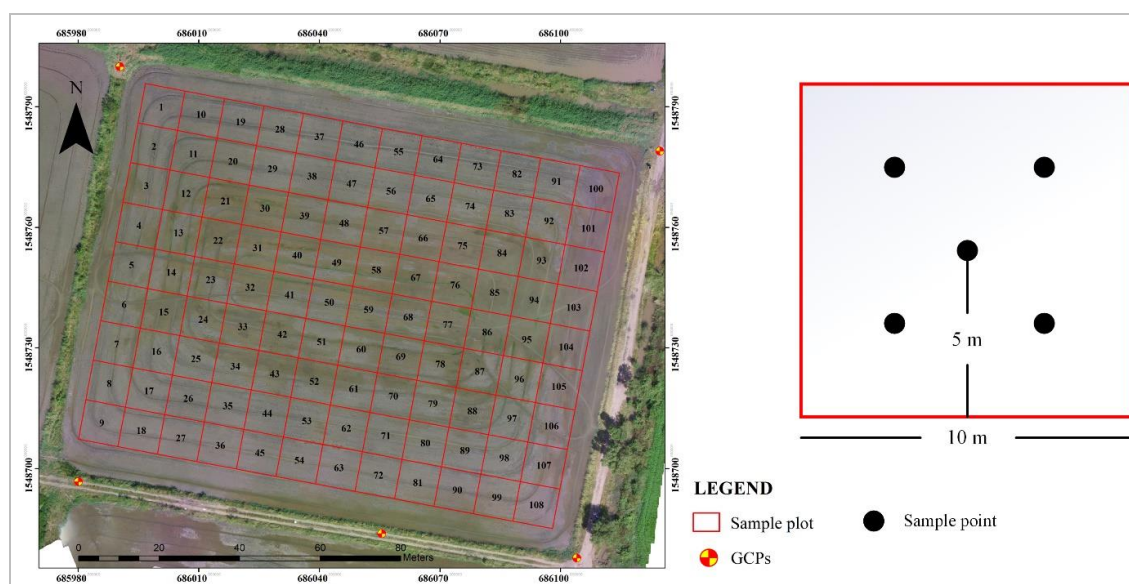


รูปที่ 1 แปลงนาทดลองประเมินความสูงดินข้าวดังอยู่ที่ศูนย์วิจัยข้าวจังหวัดปทุมธานี

2. การเก็บข้อมูลความสูงดินข้าวในแปลงทดลอง

ดำเนินการเก็บข้อมูลตัวอย่างภาคสนามโดยเลือกเก็บความสูงของข้าวที่การเจริญเติบโตสองระยะ ได้แก่ ระยะแตกกอ เก็บข้อมูลวันที่ 6 มกราคม พ.ศ. 2564 และระยะตั้งท้อง เก็บข้อมูลวันที่ 4 กุมภาพันธ์ พ.ศ. 2564 ข้อมูลถูกเก็บทั้งหมด 108 ตัวอย่างด้วยไม้วัดเป็นอุปกรณ์ในการเก็บข้อมูลมีหน่วยวัดเป็นเซนติเมตร ระยะห่างของแต่ละ

พื้นที่ตัวอย่างเป็นระยะทาง 10 เมตร (ทั้งด้านกว้างและด้านยาว) กำหนดตำแหน่งกลางพื้นที่ตัวอย่างด้วยเทปวัดระยะ และทำเครื่องหมายกำกับตำแหน่งไว้กับเชือกที่วางอยู่รอบขอบแปลง การเก็บข้อมูลจะเก็บภายในพื้นที่ตัวอย่างโดย วัดค่าความสูงบริเวณกึ่งกลางของพื้นที่ตัวอย่าง จากนั้นวัดค่าโดยการสุ่มตัวอย่างภายในพื้นที่จำนวนอีก 4 ตัวอย่าง รวมเป็น 5 ตัวอย่างต่อพื้นที่โดยสุ่มด้วยวิธี quadrat ซึ่งเป็นวิธีการสุ่มตัวอย่างข้อมูลในรูปแบบสี่เหลี่ยมจัตุรัส (รูปที่ 2) ช่วยให้การเก็บข้อมูลกระจายทั่วพื้นที่ตัวอย่างและมีรูปแบบในการเก็บง่ายขึ้น นอกจากนี้ยังเหมาะสำหรับการศึกษา ความสูงและชีวมวลในพืช (Zhang et al., 2018)



รูปที่ 2 การออกแบบวิธีการเก็บข้อมูลตัวอย่างภายในแปลงทดลอง

3. การเก็บข้อมูลความสูงต้นข้าวด้วยภาพถ่ายจาก UAV

ข้อจำกัดด้านการเก็บข้อมูลด้วย UAV สำหรับการเกษตรในประเทศไทยส่วนใหญ่มาจากต้นทุนในการทำงานที่ค่อนข้างสูงเนื่องจากต้องใช้อุปกรณ์เฉพาะทาง เครื่อง UAV ที่สามารถหาได้ตามท้องตลาดและราคาไม่สูงมาก มีกล้องที่ใช้เซนเซอร์ RGB ติดมากับตัวเครื่องจึงเป็นอุปกรณ์ที่ถูกนำมาทดลองเพื่อเก็บข้อมูลค่าความสูงของต้นข้าว งานวิจัยนี้ใช้ UAV รุ่น Phantom 3 Professional พร้อมกล้องดิจิทัลที่ภาพเซนเซอร์ RGB มีความละเอียด 12 ล้านจุดภาพ ประกอบด้วย 4 โบพัด บินได้ประมาณ 15 ถึง 20 นาทีต่อแบตเตอรี่หนึ่งก้อน แผนการบินถูกกำหนดด้วยแอปพลิเคชัน Pix4D Capture บินในรูปแบบ double grid (รูปที่ 3) ระหว่างช่วงเวลา 10.00 น. ถึง 14.00 น. ซึ่งเป็นช่วงเวลาที่มีความเหมาะสมสำหรับการเก็บข้อมูล (Starý et al., 2020; Wan et al., 2020; รัฐพงษ์ ม่วงประโคน และคณะ, 2564) กำหนดความสูงของการบินที่ 20 เมตร ให้ความละเอียดของจุดภาพต่อพื้นที่ (GSD) 1 เซนติเมตร

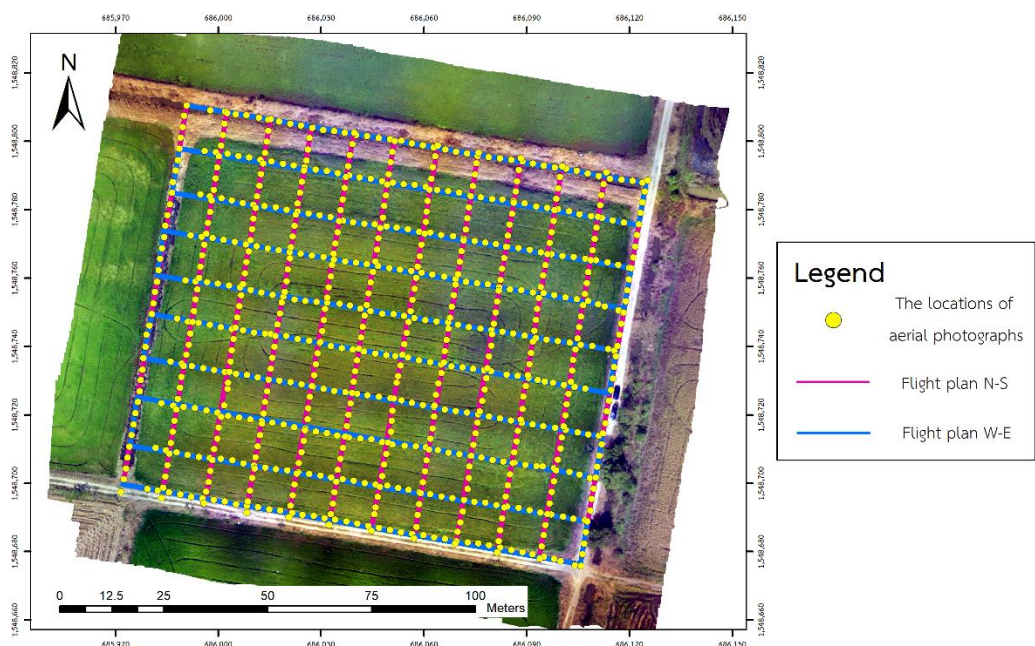
ต่อจุดภาพ กำหนดส่วนซ้อนทับด้านหน้า (over-lap) ร้อยละ 90 และด้านข้าง (side-lap) ร้อยละ 70 ทำมุมกล้อง 80 องศาขณะบินถ่าย ใช้เวลาในการบินประมาณ 63 นาทีต่อการเก็บข้อมูลหนึ่งครั้ง กำหนดจุดควบคุมภาคพื้นดิน จำนวน 5 จุด (รูปที่ 1) วางอยู่รอบตำแหน่งแปลงทดลอง ใช้ชุดเครื่องรับสัญญาณดาวเทียม GNSS ยี่ห้อ CHC รุ่น i80 และเครื่องควบคุมรุ่น LT500 ตรวจจับด้วยวิธีการตรวจจับด้วยดาวเทียมแบบจลน์ (RTK-GNSS) เพื่อเป็นข้อมูลสำหรับ ตรึงค่าพิกัดของแปลงทดลอง

ตารางที่ 1 การเก็บข้อมูลของระยะการเจริญเติบโตของข้าว

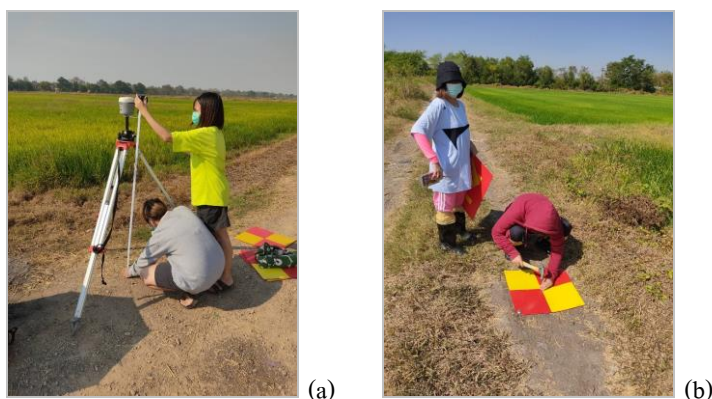
ระยะ	ระยะเวลา (นับตั้งแต่เริ่มปลูก)	วันถ่ายภาพ	รูปภาพแปลงในแต่ละระยะการเจริญเติบโต
ก่อนปลูก	-	24 พ.ย. 2563	
ระยะ แตกกอ	43 วัน	6 ม.ค. 2564	
ระยะ ตั้งท้อง	72 วัน	4 ก.พ. 2564	

เก็บข้อมูลความสูงในระยะการเจริญเติบโตของข้าวทั้งหมด 3 ระยะ ได้แก่ ระยะก่อนปลูก เก็บข้อมูลวันที่ 30 พฤศจิกายน พ.ศ. 2563 ระยะแตกกอ เก็บข้อมูลวันที่ 6 มกราคม พ.ศ. 2564 (ข้าวมีอายุ 52 วันนับจากวันปลูก) และระยะตั้งท้อง เก็บข้อมูลวันที่ 4 กุมภาพันธ์ พ.ศ. 2564 (ข้าวมีอายุ 81 วันนับจากวันปลูก) ดังแสดงในตารางที่ 1 โดยงานวิจัยของ Kawamura et al. (2020) ได้เสนอช่วงเวลาการเก็บข้อมูลความสูงต้นข้าวที่สองระยะนี้เนื่องจากการวัดความสูงของต้นข้าวจะมีความแม่นยำบางช่วงของระยะการเจริญเติบโต โดยระยะแรกที่ต้นข้าวมีความสูงไม่มากอาจเกิดข้อผิดพลาดในเรื่องความสูงของพืชกับพื้นดินปะปนกัน นอกจากนี้ช่วงการเจริญเติบโตทางลำต้นซึ่งแบ่งได้ 4

ระยะข้อยมีความแตกต่างของความสูงเพียง 1 ถึง 2 เซนติเมตรเท่านั้น (Prathumchai et al., 2018) เช่นเดียวกับระยะการเจริญเติบโตทางสืบพันธุ์ ซึ่ง UAV สามารถถ่ายภาพที่มีข้อผิดพลาดระดับเซนติเมตร เป็นเหตุผลการทำงานทดลองนี้เลือกข้าวระยะแตกกอซึ่งอยู่ในการเจริญเติบโตทางลำต้นและระยะตั้งท้องอยู่ในการเจริญเติบโตทางสืบพันธุ์เพื่อหลีกเลี่ยงการเกิดข้อผิดพลาดดังกล่าว



รูปที่ 3 แนวบินถ่ายภาพข้อมูลแปลงข้าว จุดสีเหลืองในภาพแสดงตำแหน่งจุดกึ่งกลางภาพแต่ละภาพ เส้นสีชมพูและสีฟ้าแสดงแนวบินของ UAV



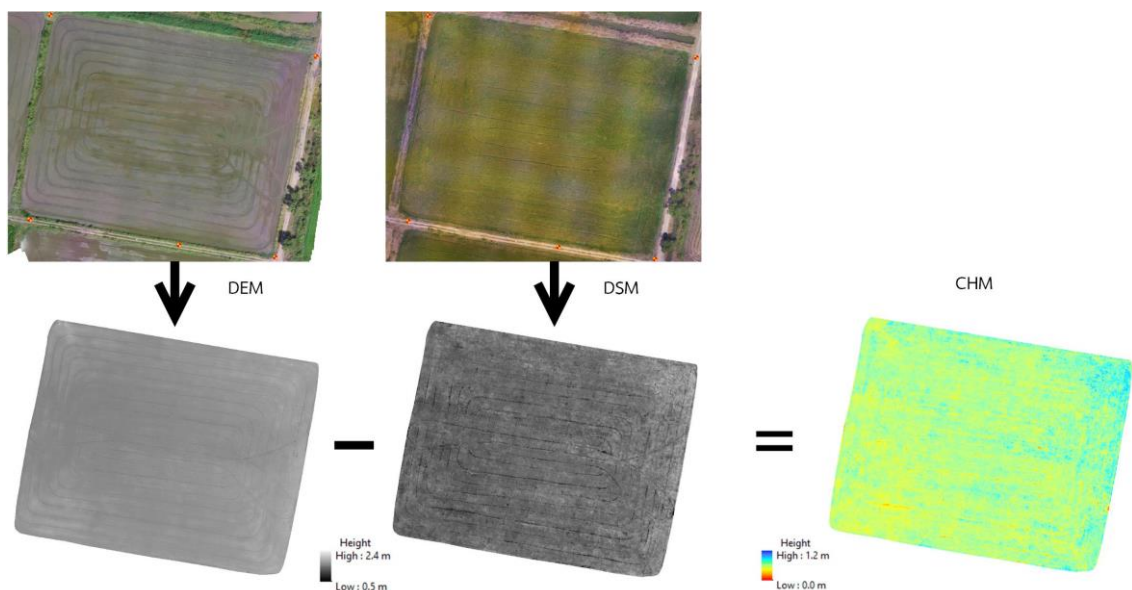
รูปที่ 4 (a) การเก็บข้อมูลจุดอ้างอิงสำหรับตรึงพิกัดภาพถ่ายจาก UAV (b) การติดตั้งเป้าอ้างอิงก่อนการบินเก็บข้อมูล

4. การประมวลผลจุดพิกัดสามมิติและแบบจำลองความสูงเรื่อนยอด

ภาพถ่ายจาก UAV ถูกนำมาสร้างแบบจำลองสามมิติ (3D model) เพื่อวิเคราะห์แบบจำลองความสูงพื้นผิวเชิงเลข (DSM) และแบบจำลองความสูงภูมิประเทศเชิงเลข (DEM) ของแปลงทดลอง ค่าความสูงที่ได้จากการวิเคราะห์จะถูกจัดเก็บร่วมกับค่าพิกัดทำให้มีความถูกต้อง (Brocks & Bareth, 2018) และมีความต่อเนื่องครอบคลุมทั้งพื้นที่ทดลอง (Kawamura et al., 2020) จากนั้นหาค่าความสูงของต้นข้าวด้วย CHM จากค่าความต่างของ DEM และ DSM ดังรูปที่ 5 ได้แสดงรายละเอียดการตั้งค่าพารามิเตอร์สำหรับคำนวณ CHM ในตารางที่ 2

5. การสร้างแบบจำลองความสัมพันธ์เชิงเส้นของความสูงจาก UAV และจากแปลงปลูกข้าว

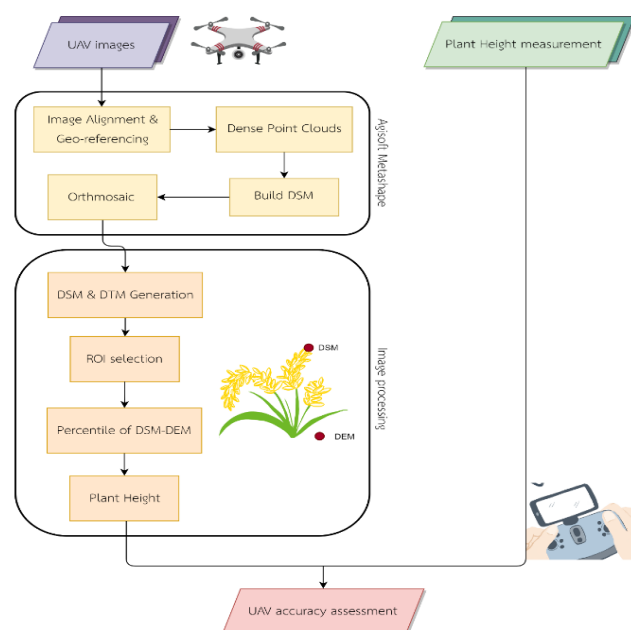
ข้อมูลความสูงที่วัดจากแปลงปลูกข้าวและจาก UAV (Pr_1 ถึง Pr_{99}) ถูกนำมาสร้างแบบจำลองความสัมพันธ์การถดถอยเชิงเส้น (linear regression model) เพื่อทดลองข้อมูลความสูงในแต่ละ Pr ที่ให้ค่าสัมประสิทธิ์การตัดสินใจ (r^2) สูงที่สุด จากนั้นคำนวณค่ารากของค่าคลาดเคลื่อนกำลังสองเฉลี่ย (root mean square error; RMSE) เพื่อตรวจสอบค่าคลาดเคลื่อนจาก CHM เปรียบเทียบกับค่าความสูงที่วัดจากแปลงปลูกข้าว รูปที่ 6 แสดงกระบวนการทดลองเพื่อหาความสูงของต้นข้าวด้วยข้อมูลจาก UAV



รูปที่ 5 ตัวอย่างการสร้าง CHM จากข้อมูล DEM และ DSM ที่ได้จากภาพถ่ายจาก UAV

ตารางที่ 2 การตั้งค่าพารามิเตอร์สำหรับการประมวลผลภาพถ่ายจาก UAV

ขั้นตอน	ตัวแปร	การตั้งค่า
Align Photos	Accuracy	Highest
	Adaptive camera model fitting	Yes
Build Dense Cloud	Quality	Ultra-high
	Depth filtering	Mild
Build texture	Mapping mode	Adaptive ortho-photo
	Blending mode	Mosaic
Build Tiled Model	Source data	Dense cloud
	Tile size	2048
Build DEM	Source data	Dense cloud
	Interpolation	Enabled
Build Orthomosaic	Blending mode	Mosaic
	Surface	DSM



รูปที่ 6 แผนผังแสดงขั้นตอนการทดลองประมาณค่าความสูงต้นข้าวด้วย UAV

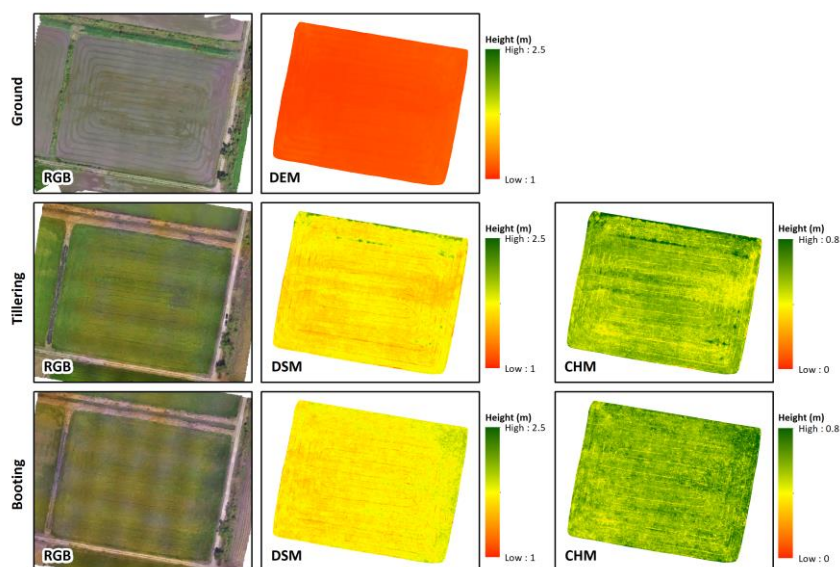
ผลและวิจารณ์ผลการทดลอง

ผลการเก็บข้อมูลความสูงจากแปลงทดลองจำนวน 108 ตัวอย่าง โดยสุ่มตัวอย่างละ 5 ค่า (ตามรูปที่ 2) ได้ค่าต่ำสุด (min) ค่าสูงสุด (max) ค่ามัธยฐาน (median) ค่าเฉลี่ยความสูง (mean) ส่วนเบี่ยงเบนมาตรฐาน (standard deviation) และค่าความแปรปรวนร่วมเกี่ยว (covariance) ดังแสดงในตารางที่ 3

ตารางที่ 3 ค่าสูงสุด ค่าต่ำสุด ค่ามัธยฐาน ค่าเฉลี่ยความสูง และส่วนเบี่ยงเบนมาตรฐานของความสูงต้นข้าวในแปลง

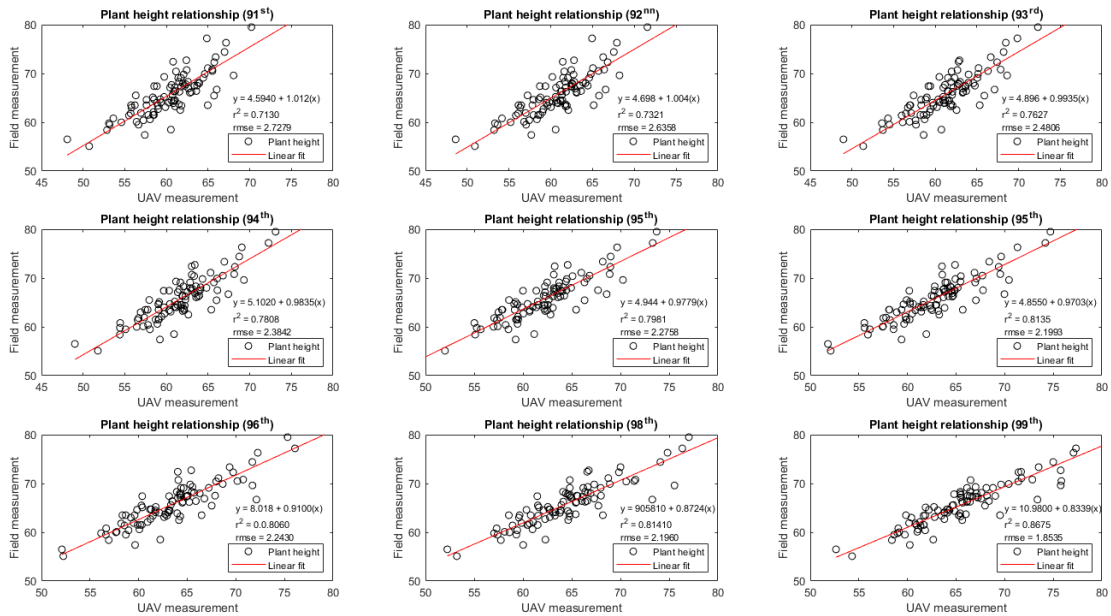
ชุดข้อมูล	จำนวน (จุด)	ค่าต่ำสุด (ซม.)	ค่าสูงสุด (ซม.)	ค่ามัธยฐาน (ซม.)	ค่าเฉลี่ย (ซม.)	ส่วน เบี่ยงเบน มาตรฐาน	ความ แปรปรวน ร่วมเกี่ยว
ระยะข้าว แตกกอ	108	45.55	102.58	67.03	67.93	5.93	8.73
ระยะข้าว ตั้งท้อง	108	47.58	124.14	82.90	83.54	7.19	8.61

จากตารางที่ 3 ค่าความสูงของข้าวระยะแตกกอและระยะตั้งท้องที่ได้จากการวัดในแปลงทดลองมีช่วงกว้าง (ระหว่าง 45.5 – 102.5 เซนติเมตร และ 47.5 – 124.1 เซนติเมตร) มีค่าเฉลี่ย±ส่วนเบี่ยงเบนมาตรฐานที่ 67.9 ± 5.9 และ 83.5 ± 7.2 เซนติเมตร โดยค่าความแปรปรวนร่วมเกี่ยวอยู่ที่ 8.7 และ 8.6 ตามลำดับ ความสูงต้นข้าวที่ได้จากข้อมูล UAV ถูกนำมาสร้าง CHM ของข้าวระยะแตกกอและระยะตั้งท้องดังแสดงในรูปที่ 7

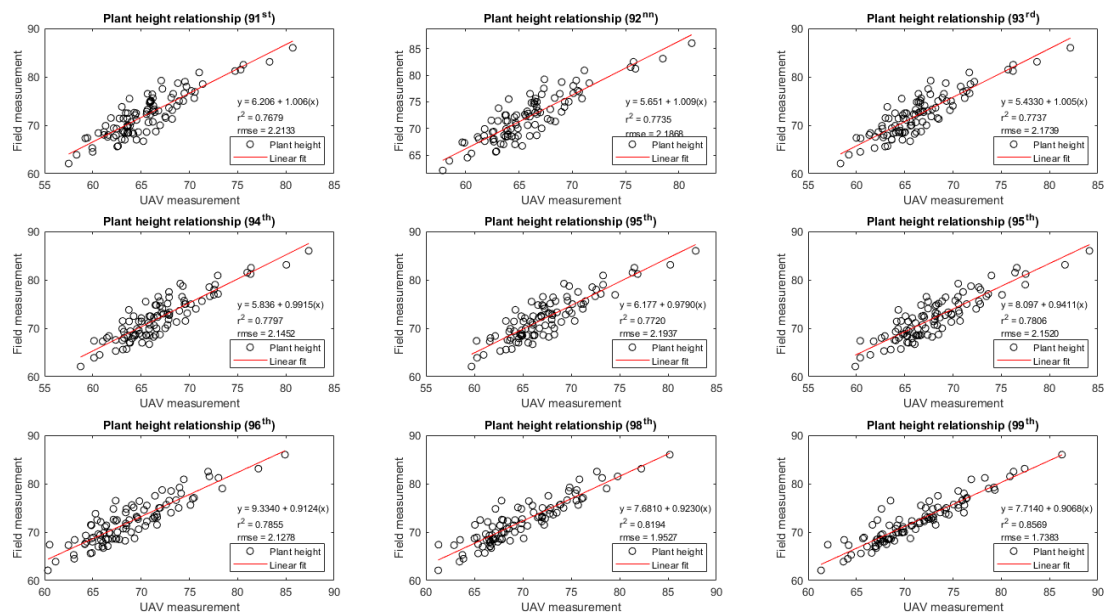


รูปที่ 7 CHM ของต้นข้าวในระยะแตกกอ (tillering) และระยะตั้งท้อง (booting)

นำความสูงต้นข้าวจาก UAV ในแต่ละระยะการเจริญเติบโตลำดับที่ Pr_{91} ถึง Pr_{99} ไปเปรียบเทียบกับความสูงต้นข้าวที่วัดจากแปลงทดลองพบว่าเฉลี่ยของความสูงให้ความสัมพันธ์ที่ดีที่สุดในการนำมาหาความสัมพันธ์กับความสูงที่ได้จาก UAV แสดงความสัมพันธ์ในรูปที่ 8 และ 9



รูปที่ 8 ความสัมพันธ์เชิงเส้นระหว่างข้อมูลความสูงจาก UAV และแปลงข้าวทดลองในระยะข้าวแตกกอ



รูปที่ 9 ความสัมพันธ์เชิงเส้นระหว่างข้อมูลความสูงจาก UAV และแปลงข้าวทดลองในระยะข้าวตั้งท้อง

จากรูปที่ 8 และ 9 ความสัมพันธ์ระหว่างความสูงของต้นข้าวในแปลงทดลองและ CHM ที่ระดับ Pr_{91} ถึง Pr_{99} แสดงความสัมพันธ์อย่างมีนัยสำคัญ ($p < 0.01$) CHM ให้ความสัมพันธ์มากที่สุดอยู่ใน Pr_{99} โดยในระยะต้นข้าวแตกกอให้ค่า r^2 ที่ 0.87 และในต้นข้าวระยะตั้งท้องให้ r^2 ที่ 0.86 ค่า RMSE ของความสูงต้นข้าวมีค่า 1.83 และ 1.71 เซนติเมตรตามลำดับ ความสัมพันธ์มีค่าลดน้อยลงตามค่าระดับ Pr การประเมินค่าความสูงด้วย UAV อาจให้ค่าความสูงของต้นข้าวต่ำกว่าความเป็นจริง ส่วนหนึ่งเป็นผลมาจากการเลือกข้อมูลตัวแทนความสูงจาก UAV (Holman et al., 2016) การเลือกข้อมูลโดยใช้ลำดับของ Pr จากชุดข้อมูลทั้งหมดช่วยลดข้อผิดพลาดที่อาจเกิดขึ้นจากกระบวนการดังกล่าวได้ Lu et al. (2021) นำความสูงที่ Pr_{10} และ Pr_{95} ถึง Pr_{99} เป็นตัวแทนค่า DEM และ DSM ที่ได้จาก UAV สำหรับวิเคราะห์ CHM เช่นเดียวกับ Iersel et al. (2017) ได้เปรียบเทียบความสูงจาก UAV โดยใช้ค่าความสูงที่ Pr_{95} เพียงค่าเดียวส่งผลให้ข้อมูลที่ใช้วิเคราะห์ความสัมพันธ์ของความสูงอาจเกิดอคติ งานวิจัยนี้ใช้ค่า Pr_{91} ถึง Pr_{99} จากข้อมูล UAV เพื่อทดสอบค่าตัวอย่างที่ให้แบบจำลองความสัมพันธ์ความสูงของต้นข้าวมีค่าถูกต้องมากที่สุด พื้นที่สุ่มแต่ละพื้นที่มีขนาด 100 ตารางเมตรส่งผลให้มีจุดภาพในแต่ละพื้นที่สุ่มจำนวน 1,000,000 จุดภาพ การใช้ค่าความสูงที่เลือกจาก Pr จึงเป็นค่าที่ช่วยให้การสร้าง CHM มีประสิทธิภาพเนื่องจากมีความเหมาะสมสำหรับการคำนวณข้อมูลความสูงที่มีจำนวนจุดภาพมากกว่า 20,000 จุดภาพขึ้นไป อีกทั้งยังให้ค่าความสูงที่โดดเด่นคล้ายกับค่าความสูงที่วัดได้จากแปลงทดลองมากที่สุดและช่วยลดค่าคลาดเคลื่อนที่เกิดจากการสะท้อนของวัสดุอื่นที่ไม่ใช่ต้นข้าว (Holman et al., 2016) ในงานของ Kawamura et al. (2020) ได้ใช้ค่า Pr_{90} ถึง Pr_{99} ในการแสดงค่าความสูงที่ได้จาก UAV เพื่อหลีกเลี่ยงค่าจุดภาพที่มีค่าความสูงมากเกินไปหรือความสูงของพืชที่มีความผิดปกติ สอดคล้องกับ Malambo et al. (2018) ที่ได้สรุปว่า Pr_{90} , Pr_{95} และ Pr_{99} ของแบบจำลองสามมิติมีความสัมพันธ์กับความสูงจากภาคสนามมากกว่าค่าสูงสุด แต่ไม่พบนัยสำคัญทางสถิติจากความแตกต่างของค่า Pr ดังกล่าว

ในการทดลองมีโอกาสดพบค่าคลาดเคลื่อนระหว่างจุดพิกัดสามมิติซึ่งเป็นผลมาจากความสูงบิน Iersel et al. (2017) ใช้ความสูงบินที่ 150 เมตรในการเก็บข้อมูลทำให้ค่าที่ได้จากการประมาณความสูงมีความผิดพลาด นอกจากนี้ส่วนเบี่ยงเบนมาตรฐานและช่วงข้อมูลมีผลต่อความถูกต้องของแบบจำลองเชิงเส้น (Grenzdörffer, 2014; Iersel et al., 2017) ส่งผลต่อค่าคลาดเคลื่อนที่เกิดขึ้นในการประมาณค่าความสูงของต้นข้าวด้วย UAV เช่นกัน (Feng A. et al., 2019; Holman et al., 2016; Kawamura et al., 2020; Wang et al., 2018) งานวิจัยนี้ได้ใช้ความสูงบินต่ำ (20 เมตร) เพื่อลดค่าคลาดเคลื่อนดังกล่าวที่เกิดขึ้นในขั้นตอนการสร้างจุดพิกัดสามมิติ (Kawamura, et al., 2020) และทดสอบความสัมพันธ์ของความสูงโดยเลือกจากค่า Pr เพื่อให้ได้มาซึ่งแบบจำลองเชิงเส้นที่มีค่าความสัมพันธ์ดีที่สุดเพื่อแก้ปัญหาค่าความแปรปรวนของข้อมูลที่เกิดขึ้นจากการเก็บข้อมูล

อย่างไรก็ตาม การวัดข้อมูลความสูงดังกล่าวยังไม่สามารถแยกชนิดของพืชได้ทำให้ข้อมูลความสูงอาจได้รับการปะปนจากพืชชนิดอื่นที่ไม่ใช่ข้าว เช่น วัชพืช และอาจส่งผลกระทบต่อการวิเคราะห์จุดพิกัดสามมิติซึ่งเป็นปัจจัยที่ไม่สามารถควบคุมได้ ยิ่งไปกว่านั้นรูปแบบของต้นข้าวมีลักษณะที่แตกต่างกันไปตามปัจจัยผลิต เช่น พันธุ์ข้าว การดูแล และการบริหารจัดการน้ำรวมทั้งการเกิดโรคในพืชอาจส่งผลให้เกิดค่าคลาดเคลื่อน หากเป็นไปได้ควร

มีการทดสอบความสูงของข้าวในทุกระยะการเจริญเติบโต (Bareth et al., 2016; Kawamura et al., 2020) เพื่อติดตามการเปลี่ยนแปลงความสูงในการเจริญเติบโตของข้าวได้ละเอียดและแม่นยำมากขึ้น

สรุปผลการทดลอง

การทดลองนี้ได้นำภาพถ่ายจาก UAV ช่วงคลื่น RGB มาใช้เพื่อสร้างจุดพิกัดสามมิติสำหรับวิเคราะห์ความสูงของต้นข้าวโดยสร้าง CHM ของข้าวในระยะแตกกอและตั้งท้องซึ่งเป็นช่วงที่มีนัยสำคัญต่อการเจริญเติบโตของต้นข้าว ค่าความสูงที่ระดับ $Pr_{0.1}$ ถึง $Pr_{0.9}$ ถูกนำมาทดสอบเปรียบเทียบหาความสัมพันธ์กับข้อมูลความสูงต้นข้าวที่วัดได้จากแปลงทดลองด้วยแบบจำลองการถดถอยเชิงเส้น เมื่อพิจารณา r^2 และ RMSE พบว่าความสูงที่ $Pr_{0.9}$ แสดงความสัมพันธ์กับค่าความสูงที่วัดจากแปลงทดลองได้สูงที่สุด ค่าความสัมพันธ์ระหว่างข้อมูลจาก UAV และข้อมูลที่ได้จากแปลงทดลองจะลดลงตามระดับ Pr การใช้ UAV เพื่อประมาณค่าความสูงของต้นข้าวสามารถดำเนินการได้รวดเร็ว สะดวก และมีต้นทุนต่ำ สามารถนำไปประยุกต์กับงานด้านการติดตามการเจริญเติบโตของต้นข้าวและเป็นอีกหนึ่งทางเลือกที่มีประสิทธิภาพ

กิตติกรรมประกาศ

คณะผู้วิจัยขอขอบคุณศูนย์วิจัยข้าวจังหวัดปทุมธานีที่อนุญาตให้ดำเนินการทดลองงานวิจัยในพื้นที่ปลูกข้าวของศูนย์วิจัยฯ และขอขอบคุณภาควิชาวิศวกรรมสำรวจ คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัยที่ให้ความอนุเคราะห์เครื่องมือ GNSS สำหรับบันทึกพิกัดอ้างอิง

เอกสารอ้างอิง

- Araus, J. L., & Cairns, J. E. (2014). Field high-throughput phenotyping: the new crop breeding frontier. *Trends in plant science*, 19(1), 52-61.
- Bareth, G., Bendig, J., Tilly, N., Hoffmeister, D., Aasen, H., & Bolten, A. (2016). A comparison of UAV-and TLS-derived plant height for crop monitoring: using polygon grids for the analysis of crop surface models (CSMs). *Photogramm. Fernerkund. Geoinf*, 2016, 85-94.
- Berhanu, S. B. (2018). *Accuracy of DTM derived from UAV imagery and its effect on canopy height model compared to Airborne LiDAR in part of tropical rain forests of Berklah, Malaysia* [Unpublished Master's thesis]. University of Twente.
- Bendig, J., Bolten, A., Bennertz, S., Broscheit, J., Eichfuss, S., & Bareth, G. (2014). Estimating biomass of barley using crop surface models (CSMs) derived from UAV-based RGB imaging. *Remote sensing*, 6(11), 10395-10412.

- Bendig, J., Willkomm, M., Tilly, N., Gnyp, M. L., Bennertz, S., Lenz-Wiedemann, V. I., Qiang, C., Miao, Y., & Bareth, G. (2015). Very high resolution crop surface models (CSMs) from UAV-based stereo images for rice growth monitoring in Northeast China. *The international Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, 40, 45-50.
- Brocks, S., & Bareth, G. (2018). Estimating Barley Biomass with Crop Surface Models from Oblique RGB Imagery. *Remote Sensing* 2018, 10(2). <https://doi.org/10.3390/rs10020268>
- Feng, A., Zhang, M., Sudduth, K. A., Vories, E. D., & Zhou, J. (2019). Cotton Yield Estimation from UAV-Based Plant Height. *American Society of Agricultural and Biological Engineers*, 62(2). <https://doi.org/10.13031/trans.13067>
- Grenzdörffer, G. J. (2014). Crop height determination with UAS point clouds. *The International Archives of Photogrammetry, Remote Sensing and Spatial Information Sciences*, 40(1), 135.
- Großkinsky, D. K., Svendsgaard, J., Christensen, S., & Roitsch, T. (2015). Plant phenomics and the need for physiological phenotyping across scales to narrow the genotype-to-phenotype knowledge gap. *Journal of Experimental Botany*, 66(18). <https://doi.org/10.1093/jxb/erv345>
- Holman, F. H., Riche, A. B., Michalski, A., Castle, M., Wooster, M. J., & Hawkesford, M. J. (2016). High throughput field phenotyping of wheat plant height and growth rate in field plot trials using UAV based remote sensing. *Remote Sensing*, 8(12), 1031.
- Iersel, W. V., Straatsma, M., Addink, E., & Middelkoop, H. (2018). Monitoring height and greenness of non-woody floodplain vegetation with UAV time series. *ISPRS journal of photogrammetry and remote sensing*, 141, 112-123.
- Iizuka, K., Yonehara, T., Itoh, M., & Kosugi, Y. (2018). Estimating tree height and diameter at breast height (DBH) from digital surface models and orthophotos obtained with an unmanned aerial system for a Japanese cypress (*Chamaecyparis obtusa*) forest. *Remote Sensing*, 10(1), 13.
- Jones, A. R., Raja Segaran, R., Clarke, K. D., Waycott, M., Goh, W. S., & Gillanders, B. M. (2020). Estimating Mangrove Tree Biomass and Carbon Content: A Comparison of Forest Inventory Techniques and Drone Imagery. *Frontiers in Marine Science*, 6, 784.
- Kawamura, K., Asai, H., Yasuda, T., Khanthavong, P., Soisouvanh, P., & Phongchanmixay, S. (2020). Field phenotyping of plant height in an upland rice field in Laos using low-cost small unmanned aerial vehicles (UAVs). *Plant Production Science*, 23(4), 452-465.

- Lu, J., Cheng, D., Geng, C., Zhang, Z., Xiang, Y., & Hu, T. (2021). Combining plant height, canopy coverage and vegetation index from UAV-based RGB images to estimate leaf nitrogen concentration of summer maize. *Biosystems Engineering*, 202, 42-54.
- Matese, A., Di Gennaro, S. F., & Berton, A. (2017). Assessment of a canopy height model (CHM) in a vineyard using UAV-based multispectral imaging. *International Journal of Remote Sensing*, 38(8-10), 2150-2160.
- Prathumchai, K., Nagai, M., Tripathi, N. K., & Sasaki, N. (2018). Forecasting transplanted rice yield at the farm scale using moderate-resolution satellite imagery and the AquaCrop model: A case study of a rice seed production community in thailand. *ISPRS International Journal of Geo-Information*, 7(2), 73.
- De Souza, C. H. W., Lamparelli, R. A. C., Rocha, J. V., & Magalhães, P. S. G. (2017). Height estimation of sugarcane using an unmanned aerial system (UAS) based on structure from motion (SfM) point clouds. *International journal of remote sensing*, 38(8-10), 2218-2230.
- Starý, K., Jelínek, Z., Kumhálová, J., Chyba, J., & Balážová, K. (2020). Comparing RGB-based vegetation indices from UAV imageries to estimate hops canopy area. *Agronomy Research*, 18(4), 2592-2601.
- Suphan, P., Kaewplang, S., & Sa-Ngiamvibool, W. (2019). Monitoring of Rice Growth with UAV-derived aerial imagery. *Maharakham International Journal of Engineering Technology*, 5(1), 28-32.
- Tirado, S. B., Hirsch, C. N., & Springer, N. M. (2020). UAV-based imaging platform for monitoring maize growth throughout development. *Plant Direct*, 4(6), e00230.
- Wan, L., Cen, H., Zhu, J., Zhang, J., Zhu, Y., Sun, D., & He, Y. (2020). Grain yield prediction of rice using multi-temporal UAV-based RGB and multispectral images and model transfer—a case study of small farmlands in the South of China. *Agricultural and Forest Meteorology*, 291, 108096.
- Wang, X., Zhang, R., Song, W., Han, L., Liu, X., Sun, X., & Zhao, J. (2019). Dynamic plant height QTL revealed in maize through remote sensing phenotyping using a high-throughput unmanned aerial vehicle (UAV). *Scientific reports*, 9(1), 1-10.
- Xie, C., & Yang, C. (2020). A review on plant high-throughput phenotyping traits using UAV-based sensors. *Computers and Electronics in Agriculture*, 178, 105731.
- Yang, C. Y., Yang, M. D., Tseng, W. C., Hsu, Y. C., Li, G. S., Lai, M. H., & Lu, H. Y. (2020). Assessment of Rice Developmental Stage Using Time Series UAV Imagery for Variable Irrigation Management. *Sensors*, 20(18), 5354.

- Zhang, H., Sun, Y., Chang, L., Qin, Y., Chen, J., Qin, Y., & Wang, Y. (2018). Estimation of grassland canopy height and aboveground biomass at the quadrat scale using unmanned aerial vehicle. *Remote sensing*, 10(6), 851.
- Zhou, L., Gu, X., Cheng, S., Yang, G., Shu, M., & Sun, Q. (2020). Analysis of plant height changes of lodged maize using UAV-LiDAR data. *Agriculture*, 10(5), 146.
- Suphan, P., Aekkuerkul, T., & Srihanu, N. (2018). Monitoring of rice growth with UAV-derived aerial imagery. In Proceedings of 56th Kasetsart University Annual Conference: Science and Genetic Engineering, Architecture and Engineering, Agro-Industry, Natural Resources and Environment (pp. 548-554). Kasetsart University, Bangkok. (in thai)
- Iamsaard, P., Apiyarat, V., Chaengchen, S., Chulert, P., & Aobpaet, A. (2020). Creating 3D models to estimate the volume of trees using UAV, a case study of Foxtail palm. In *The 25th National Convention on Civil Engineering*. SGI07-SGI07. (in thai)
- Muangprakhon, R., Khaengkhan, P., & Kaewpland, S. (2021). An evaluation of UAV-derived aerial imagery for estimating the yield of rice. *Khon kaen agriculture*, 1, 314-322. (in thai)

Research Article

ปัจจัยที่มีผลกระทบต่อระยะทางการเดินหยิบสินค้าในคลังสินค้ารูป ใบไม้

Factors affecting the picking travel distance in the leaf warehouse

มะกุสี มะแซ^{1*}, วิชชุดา บุญเรือง¹, ภักธิยา หมาดหลู¹ และ ภาณุพงศ์ วิจิตรคุณากร²

Makusee Masae^{1*}, Witchuda Boonreung¹, Phakthiya Hmadhloo¹ and Panupong Vichitkunakorn²

¹หน่วยวิจัยสถิติและการประยุกต์ สาขาวิทยาศาสตร์การคำนวณ คณะวิทยาศาสตร์ มหาวิทยาลัยสงขลานครินทร์ ประเทศไทย

²หน่วยวิจัยคณิตวิเคราะห์เชิงประยุกต์ สาขาวิทยาศาสตร์การคำนวณ คณะวิทยาศาสตร์ มหาวิทยาลัยสงขลานครินทร์ ประเทศไทย

¹Statistics and Applications Research Unit, Division of Computational Science, Faculty of Science, Prince of Songkla University, Thailand

²Applied Analysis Research Unit, Division of Computational Science, Faculty of Science, Prince of Songkla University, Thailand

*E-mail: makusee.ma@psu.ac.th

Received: 27/03/2021; Revised: 16/10/2021; Accepted: 16/11/2021

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อประเมินปัจจัยที่สำคัญที่ปัจจัยในการจัดการคลังสินค้า ได้แก่ วิธีการเดินหยิบสินค้า วิธีการเก็บสินค้า ขนาดของคลังสินค้า และขนาดของใบสั่งซื้อ ที่มีผลต่อระยะทางรวมที่ใช้ในการเดินหยิบสินค้าในคลังสินค้าที่มีลักษณะเป็นรูปใบไม้ โดยผู้วิจัยได้จำลองสถานการณ์และใช้ผลลัพธ์จากการจำลองไปศึกษาอิทธิพลหลักและอิทธิพลร่วมของปัจจัยต่าง ๆ ผลการศึกษาพบว่าปัจจัยทั้งสี่มีผลต่อระยะทางการเดินของพนักงานหยิบสินค้าและพบว่าการเดินหยิบสินค้าแบบฮิวริสติกส์ที่ต่างกันมีผลต่อระยะทางการเดินอย่างมีนัยสำคัญทางสถิติ และเมื่อเปรียบเทียบกับวิธีเดินหยิบสินค้าที่สั้นที่สุดแบบแม่นยำตรง (Exact) พบว่าวิธีฮิวริสติกส์แบบ Leaf S-shape มีประสิทธิภาพสูงกว่าวิธีฮิวริสติกส์แบบ Leaf largest gap โดยระยะทางการเดินหยิบสินค้าเฉลี่ยของวิธีฮิวริสติกส์แบบ Leaf S-shape และ Leaf largest gap แตกต่างจากวิธีเดินหยิบสินค้าที่สั้นที่สุดแบบแม่นยำตรงอยู่ในช่วง 18-28% และ 31-37% ตามลำดับ

คำสำคัญ: คลังสินค้า, การหยิบสินค้า, การเดินหยิบสินค้า, การเก็บสินค้า

Abstract

The research aims to evaluate the effects of four main factors, including order picker routing policies, storage assignment policies, warehouse sizes, and pick-list sizes, on the picking travel distance in leaf warehouses. The simulations of all possible combinations of levels were created which the outputs from them were, then, used as the inputs for a statistical study assessing the main and interaction effects from different factors. The result indicated that the four-factor interaction had a statistically significant effect on the picking travel distance. The two heuristics showed statistically significant different average travel distances. Comparing to the exact shortest routing, the Leaf S-shape outperformed the Leaf largest-gap with optimal gaps of 18-28% and 31-37%, respectively.

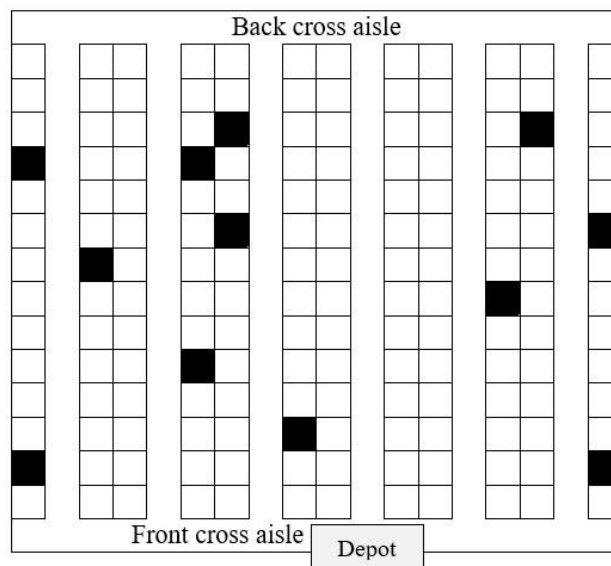
Keywords: warehousing, order picking, order picker routing, storage assignment

บทนำ

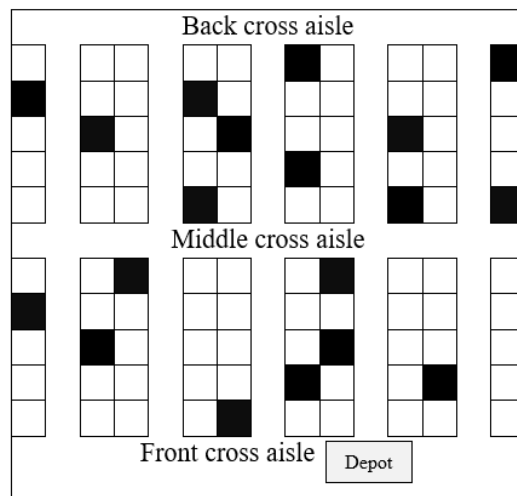
คลังสินค้ามีบทบาทสำคัญต่อระบบโลจิสติกส์เพราะทำหน้าที่เก็บสินค้าก่อนส่งสินค้าไปยังลูกค้า การจัดการคลังสินค้าอย่างมีประสิทธิภาพจะช่วยลดค่าใช้จ่ายที่ไม่จำเป็นและช่วยให้สามารถส่งสินค้าให้กับลูกค้าได้ทันเวลาที่กำหนด นักวิจัยและผู้ประกอบการได้มีการศึกษา ปรับปรุง พัฒนาการดำเนินการในคลังสินค้าให้มีความรวดเร็วและมีประสิทธิภาพ กระบวนการในคลังสินค้าประกอบด้วยขั้นตอนหลัก ๆ คือ การรับสินค้าจากผู้ผลิต การเก็บสินค้าเข้าคลังสินค้า การหยิบสินค้าจากที่เก็บสินค้า และการส่งสินค้าไปยังลูกค้า (De Koster et al., 2007; Çelik & Süral, 2014; Grosse et al., 2017) กระบวนการที่กล่าวมาทั้งหมดพบว่าขั้นตอนการหยิบสินค้า (order picking) จากที่เก็บสินค้าในคลังสินค้าใช้เวลาและแรงงานมาก ก่อให้เกิดค่าใช้จ่ายที่สูงและคิดเป็นสัดส่วนที่มากเมื่อเทียบกับค่าใช้จ่ายทั้งหมดที่เกิดขึ้นในคลังสินค้า (Won & Olafsson, 2005; Bukchin et al., 2012; Pansart et al., 2018) การหยิบสินค้าเริ่มจากพนักงานรับใบสั่งซื้อที่จุดเริ่มต้นที่เรียกว่าจุด depot จากนั้นพนักงานจะเดินไปยังตำแหน่งที่เก็บสินค้าแต่ละตำแหน่งในคลังสินค้า และเมื่อหยิบสินค้าครบตามใบสั่งซื้อแล้วพนักงานก็จะเดินกลับมายังจุด depot อีกครั้ง จากรูปที่ 1-3 แสดงการหยิบสินค้าในคลังสินค้าที่มีลักษณะเป็นแบบหนึ่งบล็อก สองบล็อก และรูปใบไม้ ตามลำดับ โดยกล่องสี่เหลี่ยมที่แรเงาด้วยสีดำแทนตำแหน่งของสินค้าที่ปรากฏในใบสั่งซื้อที่ผู้หยิบสินค้าต้องเดินไปหยิบ การลดเวลาที่ไม่ก่อให้เกิดประโยชน์และไม่เพิ่มมูลค่าจึงเป็นปัจจัยสำคัญในการลดต้นทุนการดำเนินงานของคลังสินค้า และเพิ่มประสิทธิภาพของการหยิบสินค้า เนื่องจากมีหลายปัจจัยที่คาดว่าจะส่งผลกระทบต่อค่าใช้จ่ายในการหยิบสินค้า งานวิจัยในอดีตได้ศึกษาเกี่ยวกับอิทธิพลของปัจจัยต่าง ๆ ที่มีต่อระยะทางการเดินหยิบสินค้าและพบว่าหลายปัจจัยที่ส่งผลกระทบต่อระยะทางการเดินหยิบสินค้าในคลังสินค้า เช่น วิธีการเดินหยิบสินค้า การเก็บสินค้าในคลังสินค้า ขนาดของคลังสินค้า และขนาดของใบสั่งซื้อ เป็นต้น อย่างไรก็ตามงานวิจัยส่วนใหญ่จะมุ่งเน้นไปยังคลังสินค้าที่มี

ลักษณะเป็นแบบหนึ่งบล็อกดังรูปที่ 1 หรือสองบล็อกดังรูปที่ 2 แต่ยังไม่มีการศึกษาในคลังสินค้าที่มีลักษณะเป็นรูป

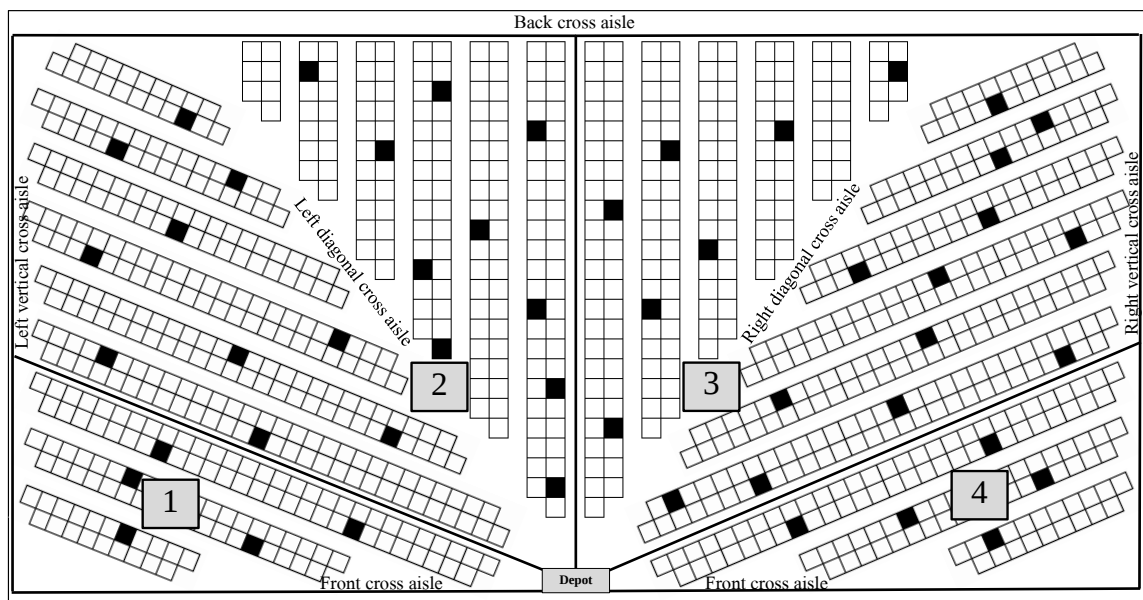
โอบีไม่มีดังรูปที่ 3 โดยคลังสินค้านี้มีเพียงการศึกษาเกี่ยวกับรูปแบบผังของคลังสินค้า (Öztürkoglu et al., 2012) และการเดินหยิบสินค้า (Masae et al., 2021) นอกจากนี้งานวิจัยของ Pohl et al. (2009) และ Çelik & Süral (2014) พบว่าการเดินหยิบสินค้าในคลังสินค้าที่ไม่ใช่แบบบล็อกสามารถลดระยะทางการเดินหยิบสินค้าในบางกรณีได้ ผู้วิจัยจึงได้เล็งเห็นความสำคัญและประโยชน์ของคลังสินค้าที่ไม่ใช่แบบบล็อก ดังนั้นงานวิจัยนี้มุ่งเน้นที่จะศึกษาปัจจัยที่มีผลกระทบต่อระยะทางการเดินหยิบสินค้าในคลังสินค้านี้ โดยผู้วิจัยได้ทำการจำลองสถานการณ์การดำเนินการในคลังสินค้าผ่านการเขียนโปรแกรมคอมพิวเตอร์ และนำผลที่ได้จากการจำลองไปประเมินอิทธิพลของปัจจัยต่าง ๆ โดยใช้เครื่องมือทางสถิติ ซึ่งผลการวิจัยจะเป็นประโยชน์ต่อนักวิจัยและผู้ประกอบการในการจัดการการหยิบสินค้าในคลังสินค้านี้ต่อไป



รูปที่ 1 คลังสินค้าแบบหนึ่งบล็อก



รูปที่ 2 คลังสินค้าแบบสองบลิ๊อค



รูปที่ 3 คลังสินค้ารูปใบไม้ (Masae et al., 2021)

วิธีการทดลอง

1. ลักษณะทั่วไปของคลังสินค้า

คลังสินค้าที่ศึกษาเป็นคลังสินค้ารูปใบไม้ถูกนำเสนอครั้งแรกโดย Öztürkoglu et al. (2012) ดังแสดงในรูปที่ 3 จะเห็นว่าคลังสินค้าประกอบด้วยทางเดินตามขวางตามแนวนอนด้านหน้าและหลัง (front and back cross aisles) ทางเดินตามแนวตั้งด้านซ้ายและขวา (left and right vertical cross aisles) และทางเดินตามเส้นทแยงมุมด้านซ้ายและขวา (left and right diagonal cross aisles) คลังสินค้ารูปใบไม้แบ่งออกได้เป็น 4 ส่วน (ตามตัวเลขที่ปรากฏในรูปที่ 3) โดยส่วนที่ 1 กับ 4 และ 2 กับ 3 มีจำนวนทางเดินหยิบสินค้าเท่ากัน คลังสินค้ามีจุด depot แต่ละตำแหน่ง ตั้งอยู่ตรงจุดกึ่งกลางของทางเดินตามแนวนอนด้านหน้า กระบวนการหยิบสินค้าเริ่มต้นจากพนักงานหยิบสินค้ารับใบสั่งซื้อที่จุด depot โดยใบสั่งซื้อนี้จะบอกจำนวนและตำแหน่งของสินค้าที่ต้องไปหยิบ (ตำแหน่งที่แทนด้วยกล่องสีดำในรูปที่ 3) เมื่อพนักงานหยิบสินค้าทุกชิ้นที่ปรากฏในใบสั่งซื้อเสร็จ ก็จะนำสินค้าที่หยิบมาทั้งหมดกลับมายังตำแหน่ง depot โดยคลังสินค้าที่ศึกษานี้เป็นคลังสินค้าที่มีทางเดินหยิบสินค้าที่แคบซึ่งจะทำให้พนักงานหยิบสินค้าสามารถหยิบสินค้าจากทั้งสองฝั่งของทางเดินได้โดยไม่มีการเพิ่มระยะทางและเวลาในการเดินส่วนนี้ นอกจากนี้คลังสินค้าประกอบด้วยชั้นวางที่มีความสูงระดับที่ผู้หยิบสามารถหยิบสินค้าได้โดยตรงโดยไม่ต้องใช้อุปกรณ์เคลื่อนที่ในแนวตั้ง

2. การออกแบบการทดลอง

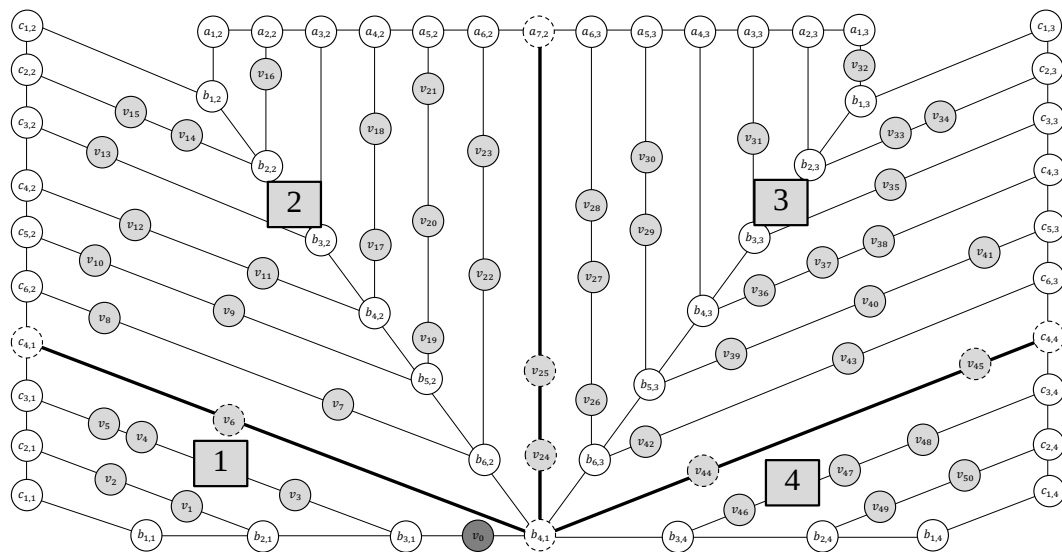
ผู้วิจัยประยุกต์แผนแบบแฟกทอเรียลเต็มรูปแบบ (full factorial design) ในการศึกษาอิทธิพลหลัก (main effect) ซึ่งเป็นอิทธิพลของปัจจัยหนึ่งที่มีต่อระยะทางการเดินหยิบสินค้าและอิทธิพลร่วม (interaction effect) ซึ่งเป็นอิทธิพลของปัจจัยตั้งแต่สองปัจจัยขึ้นไปซึ่งมีความเกี่ยวพันกันที่มีต่อระยะทางการเดินหยิบสินค้า โดยปัจจัยที่ศึกษามีดังต่อไปนี้

2.1 วิธีการเดินหยิบสินค้า (Routing) ผู้วิจัยพิจารณา 3 ระดับ (level) คือ วิธีแม่นยำตรง (Exact) วิธีอิวิริสติกส์แบบ Leaf S-shape และวิธีอิวิริสติกส์แบบ Leaf largest gap ซึ่งถูกเสนอโดย Masae et al. (2021) โดยสามารถสรุปวิธีการเดินหยิบสินค้าแต่ละวิธีได้ดังนี้

วิธีแม่นยำตรง (Exact-E) Masae et al. (2021) ประยุกต์วิธีการเดินหยิบสินค้าแบบแม่นยำตรงที่ใช้ในคลังสินค้าแบบหนึ่งบล็อกและสองบล็อกที่เสนอโดย Ratliff & Rosenthal (1983) Roodbergen & De Koster (2001) ตามลำดับ โดยปรับใช้กับคลังสินค้ารูปใบไม้ สมมติมีสินค้าทั้งหมดที่ต้องไปหยิบ m ชิ้น Masae et al. (2021) ได้แปลงการหาเส้นทางเดินหยิบสินค้า m ชิ้นในคลังสินค้าให้เป็นปัญหาการหาเส้นทางที่สั้นที่สุดบนกราฟตัวแทนของคลังสินค้า (graph representation) เริ่มต้น Masae et al. (2021) ได้สร้างกราฟตัวแทนของปัญหาดังรูปที่ 4 โดยจุด (vertex) แต่ละจุดในกราฟจะแทนตำแหน่งของสินค้าในใบสั่งซื้อ จุดตัดระหว่างทางเดิน และจุด depot ส่วนเส้น

เชื่อม (edge) แต่ละเส้นบนกราฟจะเชื่อมจุดสองจุดที่สามารถเดินทางถึงกันได้ โดยความยาวของเส้นเชื่อมมีค่าเท่ากับระยะห่างระหว่างจุดปลายทางทั้งสอง หรือระยะทางการเดินทางจริงในคลังสินค้า

เนื่องจากพนักงานหยิบสินค้าต้องเดินผ่านทุกจุดที่แทนตำแหน่งของสินค้าและต้องเดินกลับมายังจุด depot ดังนั้นปัญหาการหาเส้นทางที่สั้นที่สุดบนกราฟจึงถูกแปลงมาเป็นปัญหาการหากราฟย่อยเชื่อมโยงที่สั้นที่สุดที่บรรจุจุด depot (จุด v_0 ในรูปที่ 4) และทุกจุดที่แทนตำแหน่งของสินค้า (จุด v_j ในรูปที่ 4) โดยไม่จำเป็นต้องเดินผ่านจุดตัดระหว่างทางเดิน (จุด $a_{i,j}$, $b_{i,j}$ และ $c_{i,j}$ ในรูปที่ 4) ทุกจุด จากทฤษฎีของกราฟออยเลอร์ (Eulerian graph theorem) ทำให้ทราบว่าทุกจุดในกราฟย่อยจะมีดีกรีเป็นจำนวนคู่ ซึ่งช่วยลดกรณีในการพิจารณาไปได้จำนวนหนึ่ง จากนั้น Masae et al. (2021) ได้ประยุกต์ใช้กำหนดการพลศาสตร์ (dynamic programming) เพื่อหากราฟย่อยที่ต้องการ

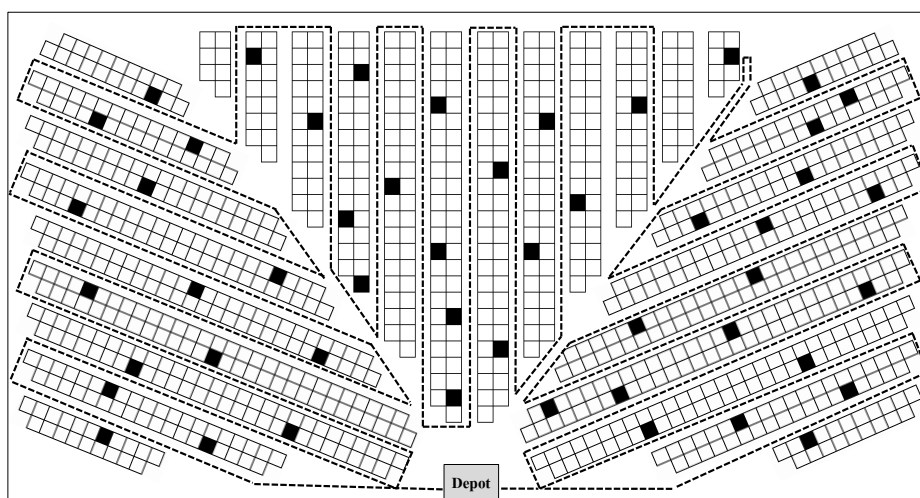


รูปที่ 4 กราฟตัวแทน (graph representation) ของปัญหาในคลังสินค้านี้แบบรูปใบไม้ (Masae et al., 2021)

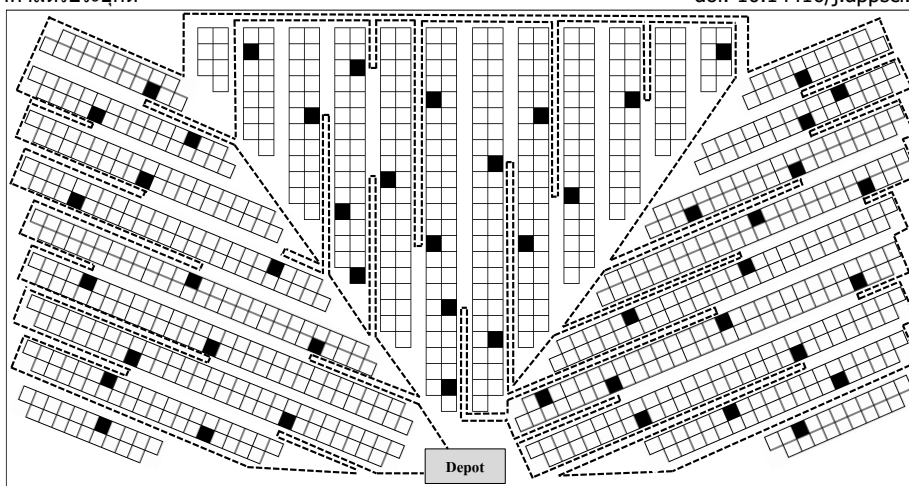
วิธีการเดินหยิบสินค้าแบบฮิวริสติกส์ได้รับความนิยมเป็นอย่างมากในทางอุตสาหกรรมเพราะมีเส้นทางการเดินที่ง่ายต่อการจดจำของผู้หยิบสินค้า (Masae et al., 2020a) ทำให้สามารถหยิบสินค้าได้อย่างรวดเร็วและลดความผิดพลาดในการหยิบสินค้าได้

วิธี Leaf-S-shape (LS) Masae et al. (2021) แบ่งคลังสินค้าออกเป็น 3 ส่วน คือ ส่วนซ้าย ส่วนกลาง และส่วนขวา พนักงานหยิบสินค้าเริ่มต้นโดยการรับใบสั่งซื้อที่จุด depot จากนั้นพนักงานหยิบสินค้าจะหยิบสินค้าเริ่มจากส่วนซ้าย ส่วนกลาง และส่วนขวา ตามลำดับ โดยจะเดินเก็บสินค้าเป็นรูปตัว S เมื่อหยิบสินค้าครบทุกชิ้น (ผ่านจุดค้าทุกจุด) พนักงานจะเดินกลับมายังจุด depot ดังแสดงในรูปที่ 5

วิธี *Leaf largest gap (LL)* การเดินโดยวิธีนี้สามารถแสดงได้ในรูปที่ 6 วิธีนี้จะแบ่งทางเดินหีบสินค้าออกเป็นสองส่วนคือ ส่วนบนและส่วนล่างโดยใช้ระยะห่างที่มากที่สุด (largest gap) ในทางเดินเป็นตัวแบ่ง โดยระยะห่างที่มากที่สุดนี้ อาจจะเป็นระยะห่างระหว่างตำแหน่งของสินค้าสองตำแหน่งที่ต้องหีบภายในทางเดินนั้น หรือเป็นระยะห่างระหว่างตำแหน่งของสินค้าที่ต้องหีบกับทางออกของทางเดิน โดยการเดินหีบสินค้าในแต่ละทางเดินที่มีสินค้าจะไม่ผ่านระยะห่างที่มากที่สุดนี้



รูปที่ 5 เส้นทางการเดินของวิธีการเดินหีบสินค้าแบบ LS (Masae et al., 2021)



รูปที่ 6 เส้นทางเดินของวิธีการเดินหิบบินค้ำแบบ LL (Masae et al., 2021)

2.2 วิธีการเก็บสินค้าในคลังสินค้า (Storage) ซึ่งจะพิจารณาด้วยกัน 4 ระดับ (level) คือ การเก็บสินค้าแบบสุ่ม (random storage; R) การเก็บตามการหมุนเวียนที่มีความเบ้ของความต้องการแบบ 20/40, 20/60 และ 20/80 (turnover-based with 20/40, 20/60, and 20/80 demand skewness; TB1, TB2, TB3)

การเก็บสินค้าแบบสุ่ม เป็นการเก็บสินค้า ณ ตำแหน่งที่ว่างในคลังสินค้า ณ เวลานั้น ซึ่งชนิดสินค้าที่จะเข้ามาเก็บในคลังสินค้า ณ ตำแหน่งใด ๆ ไม่จำเป็นต้องเป็นสินค้าที่เคยอยู่ที่ตำแหน่งนั้นมาก่อน หมายความว่าไม่มีการเจาะจงพื้นที่เก็บของสินค้าแต่ละชนิด การเก็บสินค้าวิธีนี้เป็นการเก็บสินค้าที่ง่ายเพราะไม่ต้องอาศัยการจัดระเบียบที่ซับซ้อน

การเก็บตามการหมุนเวียนที่มีความเบ้ของความต้องการ เป็นการเก็บสินค้าที่มีความต้องการสูงให้อยู่ใกล้กับตำแหน่ง depot นั่นคือระยะทางจาก depot ของสินค้าชนิดหนึ่ง ๆ จะแปรผันตรงกับความน่าจะเป็นที่สินค้าชนิดนั้นจะปรากฏอยู่ในใบสั่งซื้อ ในการศึกษานี้ ผู้วิจัยได้ปรับวิธีการจำลองการเก็บสินค้าตามการหมุนเวียนที่นำเสนอโดย Pohl et al. (2009) และ Çelik & Süral (2014) ซึ่งสามารถสรุปได้ดังนี้

กำหนดให้คลังสินค้ามีตำแหน่งเก็บสินค้าทั้งหมด N ตำแหน่ง เรียงตามระยะทางจาก depot จากน้อยไปมาก นั่นคือ ตำแหน่งที่ 1 จะอยู่ใกล้กับ depot ที่สุด และตำแหน่งที่ N จะอยู่ไกลที่สุด

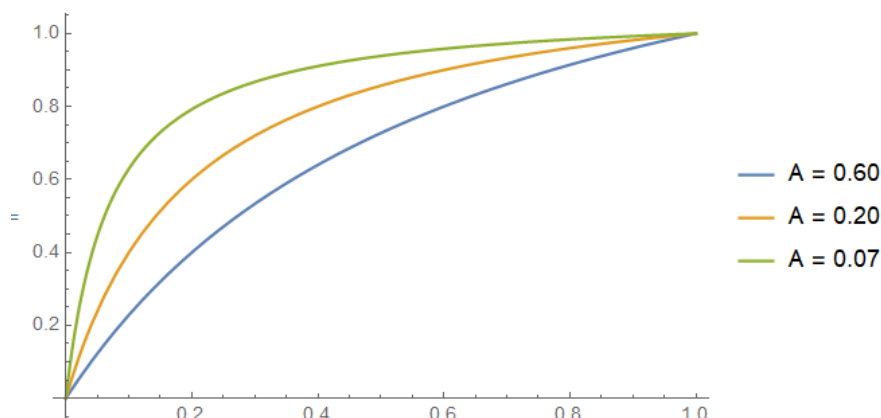
ให้ p_i แทนความต้องการของสินค้า (demand) (ความน่าจะเป็นที่สินค้าจะปรากฏในใบสั่งซื้อ) ในตำแหน่งที่ i สำหรับ $i = 1, 2, 3, \dots, N$ จากนั้นจะจำลองค่า p_i ที่สินค้าในตำแหน่งที่ i จะปรากฏในใบสั่งซื้อสำหรับทุก $i = 1, 2, 3, \dots, N$ โดยที่

$$p_i = F\left(\frac{i}{N}\right) - F\left(\frac{i-1}{N}\right) \quad (1)$$

เมื่อ $F(x)$ เป็นฟังก์ชันเพิ่มที่มีกราฟเว้าลง (concave down) (Bender, 1981) โดยที่

$$F(x) = \frac{(1+A)x}{A+x} \quad (2)$$

สำหรับ $x \in [0,1]$ และ $F(x)$ มีพารามิเตอร์ A ที่ทำให้ความเว้าของกราฟของฟังก์ชันเปลี่ยนไป ดังตัวอย่างแสดงในรูปที่ 7



รูปที่ 7 กราฟของ $F(x)$ ที่ $A = 0.60, A = 0.20$ และ $A = 0.07$

เนื่องจาก $F(x)$ เป็นฟังก์ชันเพิ่มที่มีกราฟเว้าลง จะได้ $p_1 \geq p_2 \geq \dots \geq p_N$ ดังเช่นตัวอย่างในรูปที่ 8 สำหรับกรณี $A = 0.2$ และ $N = 10$ โดยเราจะได้

$$\begin{aligned} p_1 &= F(0.1) - F(0) = 0.4 - 0 = 0.4 \\ p_2 &= F(0.2) - F(0.1) = 0.6 - 0.4 = 0.2 \\ p_3 &= F(0.3) - F(0.2) = 0.72 - 0.6 = 0.12 \\ p_4 &= F(0.4) - F(0.3) = 0.8 - 0.72 = 0.08 \\ &\vdots \\ p_{10} &= F(1) - F(0.9) = 1 - 0.98 = 0.01 \end{aligned}$$

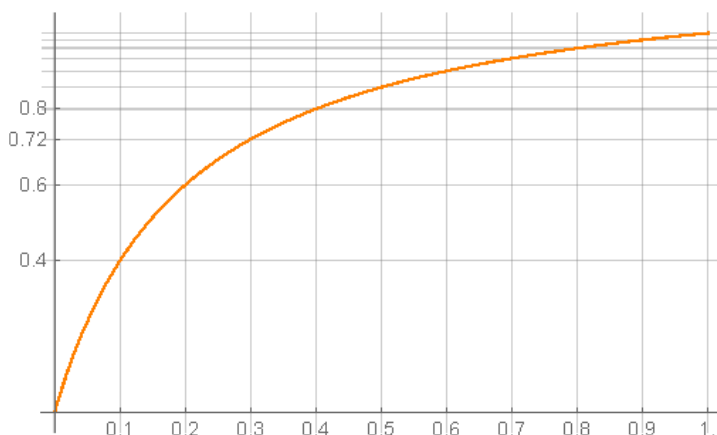
ในรูปที่ 7 จะพบว่า ถ้าค่า A ยิ่งน้อย กราฟของ $F(x)$ จะยิ่งมีค่าความเว้ามากขึ้น โดยค่าความเว้าของ $F(x)$ ที่เพิ่มขึ้น จะทำให้ผลต่างระหว่าง p_1 และ p_N มีค่าสูงขึ้น และทำให้ความแตกต่างระหว่างความต้องการสินค้าแต่ละชั้นมีค่าสูงขึ้น เราเรียกความแตกต่างดังกล่าวว่าความเบ้ของความต้องการสินค้า (demand skewness) (ความเบ้นี้แตกต่างจากความเบ้ทางสถิติที่มีค่าเท่ากับ $\frac{\text{mean}-\text{median}}{SD}$) โดยในงานวิจัยนี้เราจะกล่าวถึงความเบ้ดังกล่าวด้วยค่า $F(0.2)$ ซึ่งมีค่าเท่ากับความต้องการรวมของสินค้าที่มีความต้องการสูงสุด 20% แรก เช่น

สำหรับ $A = 0.60$ จะได้ $F(0.2) = 0.4 = 40\%$ เราใช้สัญลักษณ์ 20/40

สำหรับ $A = 0.20$ จะได้ $F(0.2) = 0.6 = 60\%$ เราใช้สัญลักษณ์ 20/60

สำหรับ $A = 0.07$ จะได้ $F(0.2) \approx 0.8 = 80\%$ เราใช้สัญลักษณ์ 20/80

ในงานวิจัยนี้ จะพิจารณาการเก็บสินค้าตามการหมุนเวียนที่มีความเบ้ของความต้องการแบบ 20/40, 20/60 และ 20/80 โดยจะแทนด้วย TB1, TB2 และ TB3 ตามลำดับ



รูปที่ 8 กราฟของ $F(x)$ ที่ $A = 0.20$ และ $N = 10$

2.3 ขนาดของคลังสินค้า ซึ่งจะพิจารณา 3 ระดับ (level) ได้แก่ คลังสินค้าที่มีจำนวนทางเดินเท่ากับ 21, 27 และ 33 ทางเดิน และให้แทนด้วย WZ1, WZ2, และ WZ3 ตามลำดับ

2.4 ขนาดของใบสั่งซื้อ โดยพิจารณา 3 ระดับ (level) ได้แก่ ขนาด 10, 20 และ 30 รายการ

ปัจจัยการทดลองและระดับที่ใช้ในการทดลองสามารถสรุปได้ดังตารางที่ 1 จากนั้นนำปัจจัยทั้ง 4 ปัจจัยมาทำการออกแบบการทดลองแบบ $3 \times 4 \times 3 \times 3$ แฟกทอเรียลเต็มรูปแบบ (full factorial design) โดยแต่ละตัวแบบจะทำการจำลองข้อมูลในโปรแกรมจาวา (Java) เพื่อที่จะประมาณค่าเฉลี่ยของระยะทางการเดินหยิบสินค้าในคลังสินค้า

ตารางที่ 1 ปัจจัยการทดลองและระดับที่ใช้ในการทดลอง

ปัจจัย (factor)	จำนวนระดับ (number of levels)	ระดับ (levels)
วิธีการเดินหีบสินค้า (routing)	3	Exact, LS, LL
วิธีการเก็บสินค้า (storage)	4	R, TB1, TB2, TB3
ขนาดคลังสินค้า (warehouse sizes)	3	WZ1, WZ2, WZ3
ขนาดของใบสั่งซื้อ (pick-list sizes)	3	10, 20, 30

ตัวแบบสถิติสำหรับการทดลองแบบแฟกทอเรียลเต็มรูปที่มี 4 ปัจจัยเขียนได้ดังนี้

$$y_{ijkl} = \mu + \alpha_i + \beta_j + \gamma_k + \lambda_l + (\alpha\beta)_{ij} + (\alpha\gamma)_{ik} + (\alpha\lambda)_{il} + (\beta\gamma)_{jk} + (\beta\lambda)_{jl} + (\gamma\lambda)_{kl} \\ + (\alpha\beta\gamma)_{ijk} + (\alpha\beta\lambda)_{ijl} + (\alpha\gamma\lambda)_{ikl} + (\beta\gamma\lambda)_{jkl} + (\alpha\beta\gamma\lambda)_{ijkl} + e_{ijkl}m$$

$i = 1,2,3; j = 1,2,3,4; k = 1,2,3; l = 1,2,3; m = 1,2,3, \dots, 30$

โดยที่

y_{ijkl}	แทนระยะทางการเดินที่ได้รับทริตเมนต์ $ijkl$ ตัวที่ m
μ	แทนค่าเฉลี่ยระยะทางการเดินทั้งหมด
α_i	แทนอิทธิพลของปัจจัยวิธีการเดินหีบสินค้าที่มี 3 ระดับ
β_j	แทนอิทธิพลของปัจจัยวิธีการเก็บสินค้าที่มี 4 ระดับ
γ_k	แทนอิทธิพลของปัจจัยขนาดคลังสินค้าที่มี 3 ระดับ
λ_l	แทนอิทธิพลของปัจจัยขนาดของใบสั่งซื้อที่มี 3 ระดับ
$(\alpha\beta)_{ij}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเดินหีบสินค้าและวิธีการเก็บสินค้า
$(\alpha\gamma)_{ik}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเดินหีบสินค้าและขนาดคลังสินค้า
$(\alpha\lambda)_{il}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเดินหีบสินค้าและขนาดของใบสั่งซื้อ
$(\beta\gamma)_{jk}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเก็บสินค้าและขนาดคลังสินค้า
$(\beta\lambda)_{jl}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเก็บสินค้าและขนาดของใบสั่งซื้อ
$(\gamma\lambda)_{kl}$	แทนอิทธิพลร่วมของปัจจัยขนาดคลังสินค้าและขนาดของใบสั่งซื้อ
$(\alpha\beta\gamma)_{ijk}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเดินหีบสินค้า วิธีการเก็บสินค้า และขนาดคลังสินค้า
$(\alpha\beta\lambda)_{ijl}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเดินหีบสินค้า วิธีการเก็บสินค้า และขนาดของใบสั่งซื้อ
$(\alpha\gamma\lambda)_{ikl}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเดินหีบสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ
$(\beta\gamma\lambda)_{jkl}$	แทนอิทธิพลร่วมของปัจจัยวิธีการเก็บสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ

$(\alpha\beta\gamma\lambda)_{ijkl}$ แทนอิทธิพลร่วมของปัจจัยวิธีการเดินหีบสินค้า วิธีการเก็บสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ

e_{ijklm} แทนความคลาดเคลื่อนสุ่มของการทดลอง

โดยต้องทดสอบสมมติฐานทางสถิติต่อไปนี้

1. การทดสอบอิทธิพลหลัก

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้า

$$H_0: \alpha_1 = \alpha_2 = \alpha_3 = 0, H_1: \alpha_i \neq 0 \text{ อย่างน้อย 1 ค่า เมื่อ } i = 1, 2, 3$$

การทดสอบสมมติฐานของปัจจัยวิธีการเก็บสินค้า

$$H_0: \beta_1 = \beta_2 = \beta_3 = \beta_4 = 0, H_1: \beta_j \neq 0 \text{ อย่างน้อย 1 ค่า เมื่อ } j = 1, 2, 3, 4$$

การทดสอบสมมติฐานของปัจจัยขนาดคลังสินค้า

$$H_0: \gamma_1 = \gamma_2 = \gamma_3 = 0, H_1: \gamma_k \neq 0 \text{ อย่างน้อย 1 ค่า เมื่อ } k = 1, 2, 3$$

การทดสอบสมมติฐานของปัจจัยขนาดของใบสั่งซื้อ

$$H_0: \lambda_1 = \lambda_2 = \lambda_3 = 0, H_1: \lambda_l \neq 0 \text{ อย่างน้อย 1 ค่า เมื่อ } l = 1, 2, 3$$

2. การทดสอบอิทธิพลร่วม

2.1 อิทธิพลร่วม 2 ปัจจัย

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้าและวิธีการเก็บสินค้า

$$H_0: (\alpha\beta)_{ij} = 0, H_1: (\alpha\beta)_{ij} \neq 0 \text{ อย่างน้อย 1 ค่า}$$

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้าและขนาดคลังสินค้า

$$H_0: (\alpha\gamma)_{ik} = 0, H_1: (\alpha\gamma)_{ik} \neq 0 \text{ อย่างน้อย 1 ค่า}$$

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้าและขนาดของใบสั่งซื้อ

$$H_0: (\alpha\lambda)_{il} = 0, H_1: (\alpha\lambda)_{il} \neq 0 \text{ อย่างน้อย 1 ค่า}$$

การทดสอบสมมติฐานของปัจจัยวิธีการเก็บสินค้าและขนาดคลังสินค้า

$$H_0: (\beta\gamma)_{jk} = 0, H_1: (\beta\gamma)_{jk} \neq 0 \text{ อย่างน้อย 1 ค่า}$$

การทดสอบสมมติฐานของปัจจัยวิธีการเก็บสินค้าและขนาดของใบสั่งซื้อ

$$H_0: (\beta\lambda)_{jl} = 0, H_1: (\beta\lambda)_{jl} \neq 0 \text{ อย่างน้อย 1 ค่า}$$

การทดสอบสมมติฐานขนาดคลังสินค้าและขนาดของใบสั่งซื้อ

$$H_0: (\gamma\lambda)_{kl} = 0, H_1: (\gamma\lambda)_{kl} \neq 0 \text{ อย่างน้อย 1 ค่า}$$

2.2 อิทธิพลร่วม 3 ปัจจัย

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้า วิธีการเก็บสินค้า และขนาดคลังสินค้า

$H_0: (\alpha\beta\gamma)_{ijk} = 0, H_1: (\alpha\beta\gamma)_{ijk} \neq 0$ อย่างน้อย 1 ค่า

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้า วิธีการเก็บสินค้า และขนาดของใบสั่งซื้อ

$H_0: (\alpha\beta\lambda)_{ijl} = 0, H_1: (\alpha\beta\lambda)_{ijl} \neq 0$ อย่างน้อย 1 ค่า

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ

$H_0: (\alpha\gamma\lambda)_{ikl} = 0, H_1: (\alpha\gamma\lambda)_{ikl} \neq 0$ อย่างน้อย 1 ค่า

การทดสอบสมมติฐานของปัจจัยวิธีการเก็บสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ

$H_0: (\beta\gamma\lambda)_{jkl} = 0, H_1: (\beta\gamma\lambda)_{jkl} \neq 0$ อย่างน้อย 1 ค่า

2.3 อิทธิพลร่วม 4 ปัจจัย

การทดสอบสมมติฐานของปัจจัยวิธีการเดินหีบสินค้า วิธีการเก็บสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ

$H_0: (\alpha\beta\gamma\lambda)_{ijkl} = 0, H_1: (\alpha\beta\gamma\lambda)_{ijkl} \neq 0$ อย่างน้อย 1 ค่า

2.5 การจำลองสถานการณ์ สำหรับแต่ละระดับของขนาดคลังสินค้าและขนาดใบสั่งซื้อ สินค้าแต่ละชิ้นจะถูกแทนด้วยตำแหน่งในคลังสินค้า ตำแหน่งในคลังสินค้าทั้งหมดจะถูกเรียงโดยพิจารณาจากระยะทางไปยังจุด depot

โดยเรียงจากระยะทางน้อยไปมาก นั่นคือ ตำแหน่ง $i = 1$ จะเป็นตำแหน่งที่อยู่ใกล้กับจุด depot มากที่สุด สำหรับตำแหน่งที่มีระยะทางเท่ากัน ลำดับของตำแหน่งเหล่านั้นจะถูกเลือกอย่างสุ่ม

สำหรับในกรณีของการเก็บสินค้าแบบสุ่ม ใบสั่งซื้อจะถูกสร้างขึ้นโดยการสุ่มสินค้า (ตำแหน่งของสินค้า) ทีละชิ้นไปจนกระทั่งได้จำนวนสินค้าเท่ากับขนาดของใบสั่งซื้อ โดยสุ่มจำนวนเต็มในช่วง 1 ถึง N เมื่อ N คือจำนวนตำแหน่งทั้งหมดในคลังสินค้า

สำหรับในกรณีของการเก็บตามการหมุนเวียนที่มีความเบ้ของความต้องการ เราพิจารณาสามระดับของการเก็บสินค้าโดยใช้ฟังก์ชัน $F(x)$ ที่มีค่า A เท่ากับ 0.60, 0.20 และ 0.07 ตามลำดับ ในแต่ละระดับดังกล่าว ใบสั่งซื้อจะถูกสร้างขึ้นโดยการสุ่มสินค้า (ตำแหน่งของสินค้า) ทีละชิ้นไปจนกระทั่งได้จำนวนสินค้าเท่ากับขนาดของใบสั่งซื้อ โดยการสุ่มสินค้าแต่ละชิ้นเกิดจากการสุ่มจำนวนจริง $x \in [0,1]$ จากนั้นหาค่า i ที่ทำให้ $F\left(\frac{i-1}{N}\right) < x \leq F\left(\frac{i}{N}\right)$ ที่แปลความหมายได้ว่า สินค้าในตำแหน่งที่ i จะปรากฏในใบสั่งซื้อ

ในแต่ละระดับของขนาดคลังสินค้า ขนาดใบสั่งซื้อ และการเก็บสินค้า เราสร้างใบสั่งซื้อทั้งหมด 30 ใบ จากนั้นพิจารณาการหีบสินค้าทั้ง 3 แบบกับแต่ละใบสั่งซื้อ แล้วจึงหาค่าเฉลี่ยของระยะการเดินหีบสินค้าจากใบสั่งซื้อทั้ง 30 ใบดังกล่าว

ผลการทดลอง

จากการวิเคราะห์ความแปรปรวนข้อมูล (Analysis of variance: ANOVA) (ตารางที่ 2) ที่ระดับนัยสำคัญ 0.05 พบว่าปัจจัยหลักทุก ๆ ปัจจัยมีอิทธิพลต่อระยะทางการเดินหีบสินค้าในคลังสินค้าอย่างมีนัยสำคัญทางสถิติ โดยจะเห็นได้จากผลการทดสอบความแตกต่างของวิธีการเดินหีบสินค้าทั้ง 3 วิธีได้ค่าสถิติ F เท่ากับ 1613932.433 และค่า $p\text{-value} < 0.001$ แสดงว่าวิธีการเดินหีบสินค้ามีผลต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติ และผลการทดสอบการเปรียบเทียบพหุคูณ (Multiple comparison) ที่ระดับนัยสำคัญ 0.05 พบว่าระยะทางการเดินหีบสินค้าเฉลี่ยของทั้งสามวิธีแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ผลการทดสอบความแตกต่างของวิธีการเก็บสินค้าทั้ง 4 วิธีได้ค่าสถิติ F เท่ากับ 907588.470 และค่า $p\text{-value} < 0.001$ แสดงว่าวิธีการเก็บสินค้ามีผลต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติ ผลการทดสอบความแตกต่างของขนาดคลังสินค้าได้ค่าสถิติ F เท่ากับ 4437129.574 และค่า $p\text{-value} < 0.001$ แสดงว่าขนาดคลังสินค้ามีผลต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติ ผลการทดสอบความแตกต่างของขนาดใบสั่งซื้อต่างกัน ได้ค่าสถิติ F เท่ากับ 5487089.837 และค่า $p\text{-value} < 0.001$ แสดงว่าขนาดใบสั่งซื้อมีผลต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติ ตารางที่ 2 ยังแสดงให้เห็นว่าขนาดของใบสั่งซื้อและขนาดของคลังสินค้ามีผลกระทบอย่างมากต่อระยะทางการเดินหีบสินค้า โดยเฉลี่ยโดยสังเกตจากค่า Sequential sum of square (Seq. SS) ที่มีค่าสูง ผลลัพธ์นี้เป็นไปตามที่คาดไว้และสมเหตุสมผลกับความเป็นจริงที่ว่าระยะทางการเดินหีบสินค้ามีแนวโน้มเพิ่มขึ้นตามขนาดของใบสั่งซื้อและขนาดของคลังสินค้า นอกจากนี้พบว่าวิธีการเก็บสินค้าในคลังสินค้ามีผลกระทบต่อระยะทางการเดินเฉลี่ยน้อยที่สุด ผล

การทดสอบผลกระทบร่วมระหว่าง 2 ปัจจัยได้แก่ 1) วิธีการเดินหีบสินค้าและวิธีการเก็บสินค้า 2) วิธีการเดินหีบสินค้าและขนาดคลังสินค้า 3) วิธีการเดินหีบสินค้าและขนาดของใบสั่งซื้อ 4) วิธีการเก็บสินค้าและขนาดคลังสินค้า 5) วิธีการเก็บสินค้าและขนาดของใบสั่งซื้อ และ 6) ขนาดคลังสินค้าและขนาดของใบสั่งซื้อ ได้ค่าสถิติ F เท่ากับ 1) 4170.581, 2) 44262.040, 3) 29658.037, 4) 19623.296, 5) 15674.445 และ 6) 163565.069 ตามลำดับ และค่า $p\text{-value} < 0.001$ สำหรับทุกการทดสอบ แสดงว่าปัจจัยร่วม 2 ปัจจัยข้างต้นมีผลต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติ ผลการทดสอบผลกระทบร่วมระหว่าง 3 ปัจจัยได้แก่ 1) วิธีการเดินหีบสินค้า วิธีการเก็บสินค้า และขนาดคลังสินค้า 2) วิธีการเดินหีบสินค้า วิธีการเก็บสินค้า และขนาดของใบสั่งซื้อ 3) วิธีการเดินหีบสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ และ 4) วิธีการเก็บสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ ได้ค่าสถิติ F เท่ากับ 1) 61.022, 2) 1183.454, 3) 2272.331, และ 4) 1138.894 ตามลำดับ และค่า $p\text{-value} < 0.001$ สำหรับทุกการทดสอบ แสดงว่าปัจจัยร่วม 3 ปัจจัยข้างต้นมีผลต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติ ผลการทดสอบผลกระทบร่วมระหว่างปัจจัยทั้ง 4 ซึ่งประกอบด้วย วิธีการเดินหีบสินค้า วิธีการเก็บสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ ได้ค่าสถิติ F เท่ากับ 21.928 และค่า $p\text{-value} < 0.001$ แสดงว่าปัจจัยทั้ง 4 มีผลกระทบร่วมกันต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติ เมื่อพิจารณาค่าสัมประสิทธิ์การตัดสินใจ (R Squared) ของตัวแบบมีค่าเท่ากับ 0.9999 หมายความว่าตัวแบบนี้สามารถอธิบายความผันแปรของข้อมูลได้

99.99% ซึ่งเป็นตัวเลขที่สูง แสดงให้เห็นว่าตัวแบบมีความเหมาะสมในการประมาณค่าระยะทางการเดินเฉลี่ยของพนักงานหยิบสินค้าในคลังสินค้า

ตารางที่ 2 ผลการวิเคราะห์ของการทดสอบ ANOVA สำหรับระยะทางการเดินเฉลี่ย

source	df	Seq. SS	Adj. MS	F	p-value
Corrected Model	107	31936921.634 ^a	298475.903	252646.028	<0.001
Intercept	1	273126002.713	273126002.713	231188511.329	<0.001
routing	2	3813398.092	1906699.046	1613932.433	< 0.001
storage	3	3216673.823	1072224.608	907588.470	< 0.001
warehouse sizes	2	10484045.743	5242022.872	4437129.574	< 0.001
pick-list sizes	2	12964890.903	6482445.452	5487089.837	< 0.001
routing * storage	6	29562.739	4927.123	4170.581	< 0.001
routing * warehouse sizes	4	209164.615	52291.154	44262.040	< 0.001
routing * pick-list sizes	4	140151.968	35037.992	29658.037	< 0.001
storage * warehouse sizes	6	139097.718	23182.953	19623.296	< 0.001
storage * pick-list sizes	6	111106.696	18517.783	15674.445	< 0.001
warehouse sizes * pick-list sizes	4	772942.795	193235.699	163565.069	< 0.001
routing * storage * warehouse sizes	12	865.096	72.091	61.022	< 0.001
routing * storage * pick-list sizes	12	16777.591	1398.133	1183.454	< 0.001
routing * warehouse sizes * pick-list sizes	8	21476.243	2684.530	2272.331	< 0.001
storage * warehouse sizes * pick-list sizes	12	16145.866	1345.489	1138.894	< 0.001
routing * storage * warehouse sizes * pick-list sizes	24	621.744	25.906	21.928	< 0.001
Error	3132	3700.143	1.181		
Total	3240	305066624.490			
Corrected Total	3239	31940621.777			

a. R Squared = 0.9999 (Adjusted R Squared = 0.9999)

เมื่อพิจารณาวิธีการเดินหิบบินค้ำแบบ E, LS และ LL ในคลังสินค้าขนาด WZ1, WZ2 และ WZ3 จากตารางที่ 3 พบว่าวิธีการเดินหิบบินค้ำแบบ E เป็นวิธีที่ให้ระยะทางการเดินเฉลี่ยน้อยที่สุดในทุก ๆ ขนาดคลังสินค้า โดยสามารถลดระยะทางการเดินเฉลี่ยได้ประมาณ 18-36% เมื่อเทียบกับวิธีการเดินหิบบินค้ำแบบ LS และ LL อย่างไรก็ตามวิธีการเดินหิบบินค้ำแบบฮิวริสติกส์ได้รับความนิยมในทางอุตสาหกรรมมากกว่าวิธีการเดินหิบบินค้ำแบบ E เนื่องจากความง่ายต่อการนำไปใช้ จากผลการทดลองพบว่า ร้อยละความแตกต่างของระยะทางของวิธีการเดินหิบบินค้ำแบบ LS เมื่อเทียบกับวิธีการเดินหิบบินค้ำแบบ E มีค่าเพิ่มขึ้นตามขนาดของคลังสินค้า เพราะความยาวของทางเดินเพิ่มขึ้นประกอบกับลักษณะการเดินหิบบินค้ำแบบ LS ที่เมื่อเดินเข้าทางใดทางหนึ่งแล้วต้องเดินไปให้สุดทางเดินถึงจะเปลี่ยนไปทางเดินอื่นได้ ต่างจากการเดินหิบบินค้ำแบบ LL ที่ค่าร้อยละความแตกต่างดังกล่าวมีค่าลดลงเมื่อคลังสินค้ามีขนาดใหญ่ขึ้น เพราะการเดินหิบบินค้ำแบบ LL ไม่จำเป็นต้องเดินให้สุดทางเดินหนึ่ง ๆ แต่สามารถเดินไปจนถึงตำแหน่งที่เก็บสินค้าแล้วเดินออกมาตามทางเข้าเดิมทำให้สามารถลดระยะทางในการเดินได้ ดังนั้นการเดินหิบบินค้ำแบบ LS จึงมีประสิทธิภาพสูงกว่าการเดินหิบบินค้ำแบบ LL ในคลังสินค้าขนาดเล็ก ส่วนการเดินหิบบินค้ำแบบ LL มีแนวโน้มที่จะมีประสิทธิภาพสูงกว่าการเดินหิบบินค้ำแบบ LS ในคลังสินค้าขนาดใหญ่

ตารางที่ 3 ระยะทางการเดินเฉลี่ย(เมตร)เมื่อเดินหิบบินค้ำในคลังสินค้าขนาดต่าง ๆ และร้อยละที่เพิ่มของระยะทางการเดินเฉลี่ยเมื่อเทียบกับวิธีการเดินหิบบินค้ำแบบ E

วิธีการเดินหิบบินค้ำ	WZ1		WZ2		WZ3	
	ร้อยละที่เพิ่ม		ร้อยละที่เพิ่ม		ร้อยละที่เพิ่ม	
	ค่าเฉลี่ย	จากวิธีการเดิน	ค่าเฉลี่ย	จากวิธีการเดิน	ค่าเฉลี่ย	จากวิธีการเดิน
		หิบบินค้ำ แบบ E		หิบบินค้ำ แบบ E		หิบบินค้ำ แบบ E
E	187.159	-	243.805	-	301.141	-
LS	220.847	18.00	299.303	22.76	382.714	27.09
LL	255.152	36.33	325.655	33.57	397.297	31.93

เมื่อพิจารณาวิธีการเดินหิบบินค้ำแบบ E, LS และ LL เมื่อขนาดใบสั่งซื้อเป็น 10, 20 และ 30 ในตารางที่ 4 พบว่าวิธีการเดินหิบบินค้ำแบบ E เป็นวิธีที่ให้ระยะทางการเดินเฉลี่ยน้อยที่สุดและมีระยะทางการเดินเฉลี่ยน้อยกว่าการเดินแบบฮิวริสติกส์อยู่ประมาณ 22-36% ในทุก ๆ ขนาดใบสั่งซื้อ ส่วนวิธีการเดินแบบฮิวริสติกส์พบว่าวิธีการเดินหิบบินค้ำแบบ LS และ LL มีค่าของร้อยละความแตกต่างจากวิธีการเดินหิบบินค้ำแบบ E ประมาณ 22-25%

และ 31-36% ตามลำดับ เห็นได้ว่าการเดินหิบบินค้ำแบบ LS มีประสิทธิภาพสูงกว่าการเดินหิบบินค้ำแบบ LL ในทุกขนาดของใบสั่งซื้อ

ตารางที่ 4 ระยะทางการเดินเฉลี่ย(เมตร)เมื่อใบสั่งซื้อมีขนาดต่าง ๆ และร้อยละที่เพิ่มของระยะทางการเดินเฉลี่ยเมื่อเทียบกับวิธีการเดินหิบบินค้ำแบบ E

ขนาดใบสั่งซื้อ	10		20		30	
วิธีการเดิน หิบบินค้ำ	ค่าเฉลี่ย	ร้อยละที่เพิ่ม	ค่าเฉลี่ย	ร้อยละที่เพิ่ม	ค่าเฉลี่ย	ร้อยละที่เพิ่ม
		จากวิธีการเดิน หิบบินค้ำ		จากวิธีการเดิน หิบบินค้ำ		จากวิธีการเดิน หิบบินค้ำ
		แบบ E		แบบ E		แบบ E
E	173.878	-	252.627	-	305.600	-
LS	212.437	22.18	315.488	24.88	374.939	22.69
LL	236.331	35.92	339.053	34.21	402.719	31.78

ตารางที่ 5 ระยะทางการเดินเฉลี่ย(เมตร)ภายใต้วิธีการเก็บสินค้าที่แตกต่างกัน

วิธีการเดินหิบบินค้ำ					
E		LS		LL	
วิธีการเก็บสินค้า	ค่าเฉลี่ย	วิธีการเก็บสินค้า	ค่าเฉลี่ย	วิธีการเก็บสินค้า	ค่าเฉลี่ย
TB3	199.244	TB3	253.381	TB3	274.088
TB2	238.479	TB2	290.896	TB2	320.819
TB1	264.213	TB1	318.754	TB1	350.070
R	274.204	R	340.788	R	359.161

จากตารางที่ 2 พบว่าวิธีการเดินหิบบินค้ำกับวิธีการเก็บสินค้ามีผลกระทบร่วมกันต่อระยะทางการเดินหิบบินค้ำอย่างมีนัยสำคัญทางสถิติ และเมื่อพิจารณาจากตารางที่ 5 แสดงให้เห็นว่าการเก็บสินค้าแบบ TB1-TB3 สามารถลดระยะทางการเดินในคลังสินค้าได้เมื่อเทียบกับการเก็บสินค้าแบบ R ไม่ว่าจะเลือกวิธีการเดินหิบบินค้ำแบบ E, LS หรือ LL โดยการเดินหิบบินค้ำแบบ E กับการเก็บสินค้าแบบ TB3 ให้ผลลัพธ์ระยะทางการเดินเฉลี่ยน้อยที่สุด กรณีการเดินหิบบินค้ำแบบฮิวริสติกส์พบว่าการเดินหิบบินค้ำแบบ LS ให้ระยะทางการเดินหิบบินค้ำน้อยกว่าการเดินหิบบินค้ำแบบ LL ไม่เพียงแต่การเก็บสินค้าแบบ TB3 เท่านั้น แต่รวมถึงการเก็บสินค้าแบบ TB2, TB1 และ R ด้วย

วิจารณ์ผลการทดลอง

ผลการวิจัยพบว่าวิธีการเดินหีบสินค้า วิธีการเก็บสินค้า ขนาดของคลังสินค้า และขนาดของใบสั่งซื้อ มีผลกระทบร่วมกันต่อระยะทางการเดินหีบสินค้า อย่างมีนัยสำคัญทางสถิติ โดยการประยุกต์ใช้วิธีการเดินหีบสินค้าแบบวิธีแมนตรงและการเก็บตามการหมุนเวียนที่มีความเบ้ของความต้องการแบบ 20/80 เป็นวิธีที่ให้ระยะทางการเดินสั้นที่สุด อย่างไรก็ตาม การเดินหีบสินค้าแบบวิธีแมนตรงมีรูปแบบการเดินที่ไม่แน่นอนและซับซ้อน ดังนั้นคลังสินค้าที่จะประยุกต์ใช้วิธีนี้ควรเป็นคลังสินค้าที่มีการนำเทคโนโลยีที่ทันสมัย มาช่วยแสดงเส้นทางการเดินให้กับพนักงานหีบสินค้า โดยอาจจะเป็นการแสดงเส้นทางเดินบนแท็บเล็ตหรือเครื่องมือช่วยหีบสินค้าอื่น ๆ สำหรับคลังสินค้าที่จะประยุกต์ใช้การเดินหีบสินค้าแบบอีวีรสถิตส์ คลังสินค้าควรเลือกการเดินหีบสินค้าแบบ Leaf S-shape เพราะจากผลการวิจัยพบว่าวิธีการเดินหีบสินค้าแบบ Leaf S-shape มีประสิทธิภาพสูงกว่า Leaf Largest gap ประกอบกับรูปแบบการเดินหีบสินค้าแบบ Leaf S-shape ง่ายต่อการจดจำของพนักงานหีบสินค้า

สรุปผลการทดลอง

งานวิจัยนี้ศึกษาผลกระทบของปัจจัยต่าง ๆ ซึ่งประกอบด้วย วิธีการเดินหีบสินค้า วิธีการเก็บสินค้า ขนาดของคลังสินค้า และขนาดใบสั่งซื้อ ต่อระยะทางการเดินหีบสินค้าของพนักงานหีบสินค้าในคลังสินค้านำไปไม่ โดยวิธีการเดินหีบสินค้าผู้วิจัยพิจารณา 3 ระดับ ได้แก่ วิธีแมนตรง วิธีอีวีรสถิตส์แบบ Leaf S-shape และ วิธีอีวีรสถิตส์แบบ Leaf Largest gap วิธีการเก็บสินค้าในคลังสินค้าพิจารณา 4 ระดับ ได้แก่ การเก็บแบบสุ่ม การเก็บตามการหมุนเวียนที่มีความเบ้ของความต้องการแบบ 20/40, 20/60 และ 20/80 ขนาดคลังสินค้าประกอบด้วย 3 ระดับ ได้แก่ คลังสินค้านำไปไม่ที่มีจำนวนทางเดินหีบสินค้าเท่ากับ 21, 27 และ 33 ทางเดิน และสุดท้ายคือขนาดของใบสั่งซื้อซึ่งพิจารณา 3 ระดับ ได้แก่ ขนาด 10, 20 และ 30 รายการ ผลจาก ANOVA สรุปได้ว่าปัจจัยหลักทุก ๆ ปัจจัย ซึ่งประกอบด้วย วิธีการเดินหีบสินค้า วิธีการเก็บสินค้า ขนาดคลังสินค้า และขนาดของใบสั่งซื้อ มีผลกระทบต่อระยะทางการเดินหีบสินค้าในคลังสินค้าอย่างมีนัยสำคัญทางสถิติที่ระดับนัยสำคัญ 0.05 นอกจากนี้ ปัจจัยร่วม 2 3 และ 4 ปัจจัยก็มีผลกระทบร่วมกันต่อระยะทางการเดินหีบสินค้าอย่างมีนัยสำคัญทางสถิติเช่นกัน งานวิจัยนี้สามารถต่อยอดโดยการพิจารณาปัจจัยอื่น ๆ เพิ่มเติม เช่น การจัดกลุ่มใบสั่งซื้อ (order batching) หรือการพิจารณารูปแบบวิธีการเดินหีบสินค้าและการเก็บสินค้าแบบอื่น ๆ นอกจากนี้ยังสามารถขยายงานวิจัยนี้ไปยังคลังสินค้านำไปแบบที่แตกต่างออกไป เช่น คลังสินค้าแบบเชฟรอน (chevron warehouse) และแบบก้างปลา (fishbone warehouse) เป็นต้น

- Bender, P. S. (1981). Mathematical modeling of the 20/80 rule: theory and practice. *Journal of Business Logistics*, 2(2), 139-157.
- Bukchin, Y., Khmelnitsky, E., & Yakuel, P. (2012). Optimizing a dynamic order-picking process. *European Journal of Operational Research*, 219(2), 335-346.
- Çelik, M., & Süral, H. (2014). Order picking under random and turnover-based storage policies in fishbone aisle warehouses. *IIE Transactions*, 46(3), 283-300.
- Chan, F. T., & Chan, H. K. (2011). Improving the productivity of order picking of a manual-pick and multi-level rack distribution warehouse through the implementation of class-based storage. *Expert systems with applications*, 38(3), 2686-2700.
- De Koster, R., Le-Duc, T., & Roodbergen, K. J. (2007). Design and control of warehouse order picking: A literature review. *European Journal of Operational Research*, 182(2), 481-501.
- Grosse, E. H., Glock, C. H., & Neumann, W. P. (2017). Human factors in order picking: a content analysis of the literature. *International Journal of Production Research*, 55(5), 1260-1276.
- Manzini, R., Gamberi, M., Persona, A., & Regattieri, A. (2007). Design of a class based storage picker to product order picking system. *The International Journal of Advanced Manufacturing Technology*, 32(7-8), 811-821.
- Masae, M., Glock, C. H., & Grosse, E. H. (2020a). Order picker routing in warehouses: A systematic literature review. *International Journal of Production Economics*, 224, 107564.
- Masae, M., Glock, C. H., & Vichitkunakorn, P. (2020b). Optimal picker routing in a conventional warehouse with two blocks and arbitrary starting and ending points of a tour. *International Journal of Production Research*, 58(17), 5337-5358.
- Masae, M., Glock, C. H., & Vichitkunakorn, P. (2020c). Optimal order picker routing in the chevron warehouse. *IIE Transactions*, 52(6), 665-687.
- Masae, M., Glock, C. H., & Vichitkunakorn, P. (2021). A method for efficiently routing order pickers in the leaf Warehouse. *International Journal of Production Economics*. <https://doi.org/10.1016/j.ijpe.2021.108069>
- Öztürkoglu, Ö., Gue, K. R., & Meller, R. D. (2012). Optimal unit-load warehouse designs for single-command operations. *IIE Transactions*, 44(6), 459-475.
- Pansart, L., Catusse, N., & Cambazard, H. (2018). Exact algorithms for the order picking problem. *Computers & Operations Research*, 100, 117-127.

- Petersen, C. G., Aase, G. R., & Heiser, D. R. (2004). Improving order-picking performance through the implementation of class-based storage. *International Journal of Physical Distribution & Logistics Management*, 34(7), 534-544.
- Pohl, L. M., Meller, R. D., & Gue, K. R. (2009). Optimizing fishbone aisles for dual-command operations in a warehouse. *Naval Research Logistics (NRL)*, 56(5), 389-403.
- Ratliff, H. D., & Rosenthal, A. S. (1983). Order-picking in a rectangular warehouse: a solvable case of the traveling salesman problem. *Operations Research*, 31(3), 507-521.
- Roodbergen, K. J., & De Koster, R. (2001). Routing order pickers in a warehouse with a middle aisle. *European Journal of Operational Research*, 133(1), 32-43.
- Theys, C., Bräysy, O., Dullaert, W., & Raa, B. (2010). Using a TSP heuristic for routing order pickers in warehouses. *European Journal of Operational Research*, 200(3), 755-763.
- Won, J., & Olafsson, S. (2005). Joint order batching and order picking in warehouse operations. *International Journal of Production Research*, 43(7), 1427-1442.

Research Article

การเปรียบเทียบประสิทธิภาพของการทดสอบทีและการทดสอบ ผลบวกลำดับที่ของวิลค็อกซันสำหรับประชากรสองกลุ่มที่เป็น อิสระกันเมื่อข้อมูลเป็นจำนวนนับ

A comparing on performance of t-test and Wilcoxon rank sum test for two independent populations using counting data

เขมิกา อูระวงศ์¹, จุฬารัตน์ ชุมناول^{1,2*}, จิรนนท์ เผ่าจรูญ¹, อัสลยา สิงห์บำรุง¹

Khemika Urawong¹, Jularat Chumnaul^{1,2*}, Chiranan Phaochamrun¹, Asalaya Singhabumrung¹

¹สาขาวิทยาศาสตร์การคำนวณ คณะวิทยาศาสตร์ มหาวิทยาลัยสงขลานครินทร์ หาดใหญ่สงขลา 90112 ประเทศไทย

¹Division of Computational Science, Faculty of Science, Prince of Songkla University, Hat Yai, Songkhla 90112, Thailand

²หน่วยวิจัยสถิติและการประยุกต์ คณะวิทยาศาสตร์ มหาวิทยาลัยสงขลานครินทร์ หาดใหญ่สงขลา 90112 ประเทศไทย

²Statistics and Applications Research Unit, Faculty of Science, Prince of Songkla University, Hat Yai, Songkhla 90112, Thailand

*E-mail: jularat.c@psu.ac.th

Received: 04/10/2021; Revised: 27/03/2022; Accepted: 06/05/2022

บทคัดย่อ

การวิจัยในครั้งนี้มีวัตถุประสงค์เพื่อเปรียบเทียบประสิทธิภาพของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับทดสอบความแตกต่างระหว่างค่ากลางของประชากรสองกลุ่มที่เป็นอิสระกันเมื่อข้อมูลเป็นจำนวนนับ โดยการแจกแจงที่นำมาศึกษา คือ การแจกแจงทวินาม และการแจกแจงปัวซอง การเปรียบเทียบประสิทธิภาพของการทดสอบพิจารณาจากความสามารถในการควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 และกำลังการทดสอบ ผลการศึกษาพบว่า การทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสามารถควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 ได้ทุกกรณีที่ศึกษา สำหรับกำลังการทดสอบพบว่า ในกรณีที่ขนาดตัวอย่างทั้งสองกลุ่มเท่ากัน กำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงทวินามและการแจกแจงปัวซองมีแนวโน้มเพิ่มขึ้นเมื่อขนาดตัวอย่างทั้งสองกลุ่มและขนาดอิทธิพลเพิ่มขึ้น และเมื่อพิจารณาอิทธิพลของอัตราส่วนของขนาด

ตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ที่มีต่อการทดสอบภายใต้ขนาดตัวอย่างรวมที่เท่ากันพบว่า การทดสอบของ การทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันจะมีค่าสูงสุดเมื่ออัตราส่วนของขนาดตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ($n_1 : n_2$) เท่ากับ 1 : 1 ยกเว้นกรณีที่ n_1 และ n_2 เท่ากับ 10 และจะมีค่าน้อยที่สุดเมื่ออัตราส่วนของขนาด ตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ($n_1 : n_2$) เท่ากับ 1 : 5 (หรือ 5 : 1) ทั้งในกรณีที่ข้อมูลมีการแจกแจงทวินามและกรณีที่ ข้อมูลมีการแจกแจงปัวซอง นอกจากนี้ยังพบว่า การทดสอบทีซึ่งเป็นการทดสอบแบบอิงพารามิเตอร์ให้ค่าการ ทดสอบสูงกว่าการทดสอบผลบวกลำดับที่ของวิลค็อกซันทุกกรณี

คำสำคัญ: การทดสอบที, การทดสอบผลบวกลำดับที่ของวิลค็อกซัน, ค่าการทดสอบ, ความผิดพลาดแบบที่ 1

Abstract

This research aimed to compare the performance of the t-test and Wilcoxon rank sum test for testing the difference between the mean values of two independent populations using counting data. The distribution considered in this study were binomial distribution and Poisson distribution. The performance of these two tests was compared considering the ability to control the probability of type I error and power. The results showed that the t-test and Wilcoxon rank sum test could control the probability of type I error for all situations. In the case of power, when sample sizes of two groups were equal, the powers of t-test and Wilcoxon rank sum test for testing the difference between the mean values of two independent populations using counting data tended to increase when sample sizes of two groups, effect size, and the value of parameter n increased. Considering the effect of the ratios of sample sizes between groups toward powers, the results showed that the powers of the t-test and Wilcoxon rank sum test became maximum when the ratio of sample sizes between groups ($n_1 : n_2$) was 1 : 1 except for n_1 and n_2 equal 10 and became minimum when the ratio of sample sizes between groups ($n_1 : n_2$) was 1 : 5 (or 5 : 1) for both binomial and Poisson data. Moreover, the t-test which is a parametric test yielded higher powers than the Wilcoxon rank sum test for all cases.

Keywords: t-test, Wilcoxon rank sum test, power of a test, type I error

บทนำ

วิธีทางสถิติแบบอิงพารามิเตอร์ในการทดสอบความแตกต่างระหว่างค่าเฉลี่ยของประชากรสองกลุ่มที่เป็นอิสระกันที่นิยมใช้กันอย่างแพร่หลาย คือ การทดสอบที (Two independent samples t-test) ซึ่งการทดสอบดังกล่าวมี ข้อตกลงเบื้องต้นที่สำคัญ คือ ข้อมูลตัวอย่างทั้งสองกลุ่มต้องมาจากประชากรที่มีการแจกแจงปกติ แต่ในทางปฏิบัติ ข้อมูลไม่ได้มีการแจกแจงปกติเสมอไปโดยเฉพาะข้อมูลในงานวิจัยเชิงประยุกต์ เช่น ข้อมูลจำนวนผู้เข้ารับบริการ

จำนวนสินค้าที่ชำรุด จำนวนผู้ป่วยติดเชื้อ COVID-19 รายใหม่ จำนวนครั้งการเข้ารับการรักษาในโรงพยาบาลในระยะเวลา 1 ปี จำนวนครั้งการพยายามฆ่าตัวตายของผู้ป่วยที่เป็นโรคซึมเศร้า เป็นต้น การใช้การทดสอบที่ซึ่งเป็นการทดสอบอิงพารามิเตอร์ (parametric test) ในการวิเคราะห์ข้อมูลดังกล่าวอาจส่งผลให้ข้อสรุปไม่ถูกต้อง ผลการวิเคราะห์ข้อมูลที่ได้ไม่น่าเชื่อถือและไม่สามารถนำไปอ้างอิงหรือไปอธิบายสิ่งที่ต้องการศึกษาได้ ทั้งนี้ เนื่องจากสถิติทดสอบที่ถูกสร้างขึ้นภายใต้เงื่อนไขที่ว่าข้อมูลต้องเป็นข้อมูลเชิงปริมาณแบบต่อเนื่องและมาจากประชากรที่มีการแจกแจงปกติที่ไม่ทราบค่าความแปรปรวนของประชากร ตัวสถิติทดสอบจึงจะมีการแจกแจงที่ (Student's t-distribution) ตามทฤษฎีทางสถิติและให้กำลังการทดสอบ (power of a test) สูง การนำการทดสอบที่มาใช้กับข้อมูลจำนวนนับซึ่งเป็นข้อมูลเชิงปริมาณแบบไม่ต่อเนื่องและไม่ได้มีการแจกแจงปกติ จึงไม่สามารถการันตีได้ว่า การทดสอบที่จะยังให้ประสิทธิภาพดีหรือให้กำลังการทดสอบที่สูงอย่างที่ควรจะเป็น ในบางชุดข้อมูล การทดสอบที่อาจยังให้กำลังการทดสอบสูง แต่ในบางชุดข้อมูล การทดสอบที่อาจให้กำลังการทดสอบต่ำ ดังนั้น เพื่อแก้ไขปัญหาดังกล่าว การทดสอบไม่อิงพารามิเตอร์ (non-parametric test) จึงเป็นอีกทางเลือกหนึ่งที่ถูกนำมาใช้ในกรณีที่ข้อมูลไม่เป็นไปตามเงื่อนไขหรือข้อตกลงเบื้องต้นของการทดสอบอิงพารามิเตอร์

สำหรับการทดสอบความแตกต่างระหว่างค่าเฉลี่ยของประชากร 2 กลุ่มที่เป็นอิสระกัน ในกรณีที่ข้อมูลไม่ผ่านข้อตกลงเบื้องต้นของการแจกแจงปกติ การทดสอบผลบวกลำดับที่ของวิลค็อกซัน (Wilcoxon rank sum test) มักจะถูกนำมาใช้แทนการทดสอบที (t-test) โดยมีงานวิจัยหลาย ๆ งานที่แสดงให้เห็นว่า การทดสอบผลบวกลำดับที่ของวิลค็อกซันให้กำลังการทดสอบ (power of the test) ที่สูงกว่าการทดสอบทีในหลาย ๆ สถานการณ์ ยกตัวอย่างเช่น Blair & Higgins (1980) และ Bridge & Sawilowsky (1999) ได้ศึกษากำลังการทดสอบของการทดสอบที (t-test) และการทดสอบผลบวกลำดับที่ของวิลค็อกซัน (Wilcoxon rank sum test) ในกรณีที่ข้อมูลมีขนาดเล็กและในกรณีที่ข้อมูลเป็นข้อมูลเชิงปริมาณแบบต่อเนื่องที่ไม่ได้มีการแจกแจงปกติ ซึ่งผลการวิจัยพบว่า การทดสอบผลบวกลำดับที่ของวิลค็อกซันมีกำลังการทดสอบสูงกว่าการทดสอบทีในกรณีที่ข้อมูลมีขนาดเล็กและมีความแปรปรวนสูง Sangthong (2013) ได้ศึกษาประสิทธิภาพของสถิติอิงพารามิเตอร์และสถิติไม่อิงพารามิเตอร์ในการทดสอบความแตกต่างของค่ากลางระหว่างประชากร 2 กลุ่มในกรณีที่ข้อมูลมีการแจกแจงปกติ ข้อมูลมีการแจกแจงเบ้ซ้ายและความโด่งสูงกว่าปกติ และข้อมูลมีการแจกแจงเบ้ขวาและความโด่งสูงกว่าปกติ ซึ่งผลการวิจัยพบว่า การทดสอบทีและการทดสอบ Van der Waerden มีกำลังการทดสอบสูงสุดและสามารถควบคุมความผิดพลาดแบบที่ 1 ได้เมื่อตัวอย่างมีขนาดเล็ก สำหรับทุกการแจกแจง ส่วนในกรณีที่ข้อมูลมีการแจกแจงเบ้ซ้ายและความโด่งสูงกว่าปกติและข้อมูลมีการแจกแจงเบ้ขวาและความโด่งสูงกว่าปกติ การทดสอบแมนน์-วิทนียู (Mann-Whitney U test) มีกำลังการทดสอบสูงสุดและสามารถควบคุมความผิดพลาดแบบที่ 1 ได้เมื่อตัวอย่างมีขนาดกลางและขนาดใหญ่ นอกจากนี้ Sangthong (2014) ยังได้ศึกษาความแกร่งและกำลังการทดสอบของสถิติอิงพารามิเตอร์และสถิติไม่อิงพารามิเตอร์ในการทดสอบความแตกต่างของค่ากลางระหว่างประชากรสองกลุ่มสำหรับข้อมูลแบบลิเคิร์ท 5 ระดับ ซึ่งผลการวิจัยพบว่า การทดสอบทีมีกำลังการทดสอบสูงสุดและสามารถควบคุมความผิดพลาดแบบที่ 1 ได้เมื่อตัวอย่างมีขนาดเล็กสำหรับทุกการ

แจกแจง ส่วนการทดสอบแมนน์-วิตนีย์มีกำลังการทดสอบสูงสุดและสามารถควบคุมความผิดพลาดแบบที่ 1 ได้ เมื่อข้อมูลมีการแจกแจงปกติและตัวอย่างมีขนาดใหญ่ และเมื่อข้อมูลมีการแจกแจงเบ้ซ้ายและความโด่งต่ำกว่า ปกติและตัวอย่างมีขนาดเล็ก Araveeporn et al. (2017) ได้ศึกษาประสิทธิภาพของการทดสอบซี (z-test) การทดสอบที (t-test) การทดสอบแบบสุ่ม (randomization test) และการทดสอบแมนน์-วิตนีย์ (Mann-Whitney U test) สำหรับทดสอบค่าเฉลี่ยหรือค่ากลางของประชากร 2 กลุ่มที่เป็นอิสระกัน โดยผลการวิจัยพบว่า การทดสอบซี การทดสอบที และการทดสอบแมนน์-วิตนีย์สามารถควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 ได้ทั้งในกรณี ตัวอย่างขนาดเล็กและตัวอย่างขนาดใหญ่ และเมื่อข้อมูลทั้ง 2 กลุ่มมีความแปรปรวนไม่เท่ากัน การทดสอบซีมีกำลังการทดสอบสูงที่สุด และ Kim & Park (2019) ได้ทำการศึกษาเกี่ยวกับข้อดคลงเบื้องต้นของการทดสอบทีสำหรับประชากร 2 กลุ่มที่เป็นอิสระกันและพบว่า ถ้าอัตราส่วนของขนาดตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ($n_1 : n_2$) เท่ากับ 1 : 1 จะให้กำลังการทดสอบที่สูงที่สุด

ดังนั้น ในการวิจัยนี้จึงต้องการศึกษาประสิทธิภาพของการทดสอบทีและการทดสอบผลบวกลำดับที่ของ วิลค็อกซันสำหรับข้อมูลเชิงปริมาณไม่ต่อเนื่องที่ได้จากการนับ (counting data) โดยการแจกแจงที่นำมาพิจารณา คือ การแจกแจงทวินาม (Binomial distribution) และการแจกแจงปัวซอง (Poisson distribution) ซึ่งเป็นการแจกแจงความน่าจะเป็นของตัวแปรสุ่มแบบไม่ต่อเนื่องที่สำคัญที่มีค่าของตัวแปรสุ่มเป็นจำนวนนับและเป็นการแจกแจงที่ใช้กัน อย่างแพร่หลายในงานวิจัยเชิงประยุกต์ ทั้งนี้ เพื่อเป็นแนวทางให้กับผู้วิเคราะห์ข้อมูลในการเลือกใช้สถิติทดสอบ สำหรับการทดสอบสมมุติฐานความแตกต่างระหว่างค่ากลางของประชากรสองกลุ่มที่เป็นอิสระกันสำหรับข้อมูล จำนวนนับได้อย่างเหมาะสม

วิธีการวิจัย

การวิจัยครั้งนี้ใช้การจำลองข้อมูล (simulation study) ในการตรวจสอบประสิทธิภาพของการทดสอบที และการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับประชากรสองกลุ่มที่เป็นอิสระกันเมื่อข้อมูลเป็นจำนวนนับ โดยผู้วิจัยได้กำหนดขอบเขตและขั้นตอนการดำเนินการวิจัยดังนี้

ขอบเขตการวิจัย

1. ข้อมูลที่ศึกษามีการแจกแจงทวินาม (Binomial distribution) และการแจกแจงปัวซอง (Poisson distribution) โดยมีพารามิเตอร์ของการแจกแจงดังนี้

การแจกแจงทวินาม $(n, p_1, p_2) = (10, 0.1, 0.1), (10, 0.5, 0.5), (10, 0.7, 0.7), (50, 0.1, 0.1), (50, 0.5, 0.5), (50, 0.7, 0.7), (100, 0.1, 0.1), (100, 0.5, 0.5), (100, 0.7, 0.7), (10, 0.1, 0.125), (10, 0.5, 0.525), (10, 0.7, 0.725), (50, 0.1, 0.125), (50, 0.5, 0.525), (50, 0.7, 0.725), (100, 0.1, 0.125), (100, 0.5, 0.525), (100, 0.7, 0.725)$

การแจกแจงปัวซอง $(\lambda_1, \lambda_2) = (1, 1), (5, 5), (10, 10), (1, 1.25), (1, 1.5), (1, 2), (5, 8)$

2. ขนาดตัวอย่างที่ศึกษาจำแนกเป็น 2 กรณี คือ กรณีขนาดตัวอย่างเท่ากัน และกรณีขนาดตัวอย่างไม่เท่ากัน

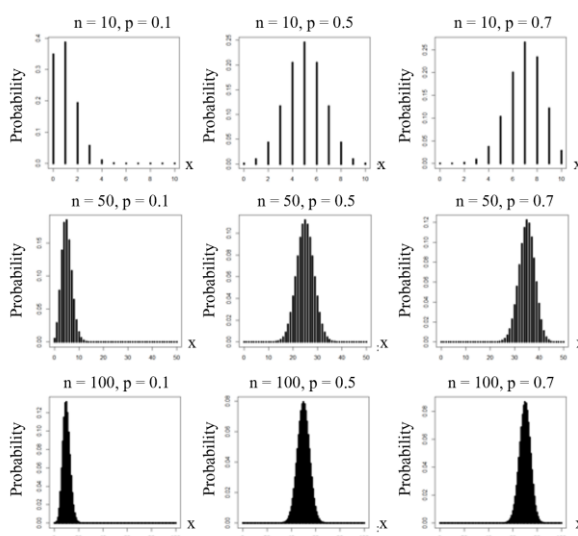
- 2.1 กรณีขนาดตัวอย่างเท่ากัน $(n_1, n_2) = (10, 10), (60, 60), (100, 100)$
- 2.2 กรณีขนาดตัวอย่างไม่เท่ากัน $(n_1, n_2) = (20, 100), (30, 90), (40, 80), (100, 20), (90, 30), (80, 40)$
3. การทดสอบที่ศึกษาเป็นการทดสอบแบบสองทาง (two-tailed test) โดยกำหนดระดับนัยสำคัญ $\alpha = 0.05$
4. ประสิทธิภาพของการทดสอบพิจารณาจากความสามารถในการควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 (type I error) และกำลังการทดสอบ (power of a test)
5. ทำการทดลองซ้ำ 10,000 ครั้งในแต่ละสถานการณ์ที่กำหนด

ขั้นตอนการดำเนินการวิจัย

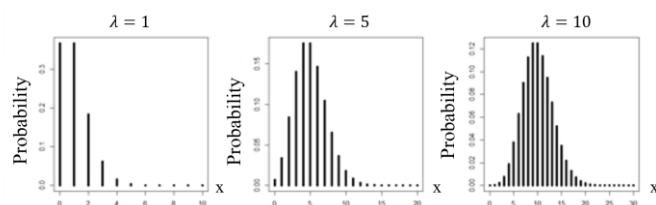
1. สร้างข้อมูลสำหรับการใช้ในการวิจัยตามสถานการณ์ต่าง ๆ ที่กำหนดไว้ในขอบเขตการวิจัย กล่าวคือ สำหรับการแจกแจงทวินาม ทำการสร้างตัวแปรสุ่มตัวที่ 1 ที่มีการแจกแจงทวินามด้วยพารามิเตอร์ n และ p_1 และสร้างตัวแปรสุ่มตัวที่ 2 ที่มีการแจกแจงทวินามด้วยพารามิเตอร์ n และ p_2 โดยมีขนาดตัวอย่างเท่ากับ n_1 และ n_2 ตามลำดับ สำหรับการแจกแจงปัวซอง ทำการสร้างตัวแปรสุ่มตัวที่ 1 ที่มีการแจกแจงปัวซองด้วยพารามิเตอร์ λ_1 และสร้างตัวแปรสุ่มตัวที่ 2 ที่มีการแจกแจงปัวซองด้วยพารามิเตอร์ λ_2 โดยมีขนาดตัวอย่างเท่ากับ n_1 และ n_2 ตามลำดับ โดยในการสร้างตัวแปรสุ่มทั้ง 2 ตัว กำหนดให้ค่าสัมประสิทธิ์สหสัมพันธ์ระหว่างตัวแปรสุ่มทั้ง 2 ตัวเท่ากับ 0 เพื่อให้ตัวแปรสุ่มทั้ง 2 ตัว เป็นอิสระกัน และในการศึกษาความน่าจะเป็นของความผิดพลาดแบบที่ 1 ทำการสร้างข้อมูลภายใต้สมมติฐานว่าง (null hypothesis, H_0) ส่วนการศึกษาการทดสอบ ทำการสร้างข้อมูลภายใต้สมมติฐานทางเลือก (alternative hypothesis, H_1) โดยสมมติฐานของการทดสอบ คือ

H_0 : ประชากรทั้งสองกลุ่มมีค่ากลาง (ค่าเฉลี่ย/มัธยฐาน) ไม่แตกต่างกัน

H_1 : ประชากรทั้งสองกลุ่มมีค่ากลาง (ค่าเฉลี่ย/มัธยฐาน) แตกต่างกัน



รูปที่ 1 ข้อมูลที่มีการแจกแจงทวินามภายใต้ขอบเขตการวิจัย



รูปที่ 2 ข้อมูลที่มีการแจกแจงปัวซองภายใต้ขอบเขตการวิจัย

2. จำนวนค่าสถิติทดสอบที่และสถิติทดสอบผลบวกลำดับที่ของวิลค็อกซัน

2.1 สถิติทดสอบที่ (Moore, McCabe & Craig, 2014)

2.1.1 ถ้า $\sigma_1^2 = \sigma_2^2$ ค่าสถิติทดสอบที่สามารถคำนวณได้ดังนี้

$$t^* = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{s_p^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}, \quad \nu = n_1 + n_2 - 2 \quad (1)$$

โดยที่ $\bar{X}_1$ คือ ค่าเฉลี่ยตัวอย่างกลุ่มที่ 1
 s_1^2 คือ ความแปรปรวนตัวอย่างกลุ่มที่ 1
 n_1 คือ ขนาดตัวอย่างกลุ่มที่ 1
 $\bar{X}_2$ คือ ค่าเฉลี่ยตัวอย่างกลุ่มที่ 2
 s_2^2 คือ ความแปรปรวนตัวอย่างกลุ่มที่ 2
 n_2 คือ ขนาดตัวอย่างกลุ่มที่ 2
 ν คือ องศาเสรี (degree of freedom)
 s_p^2 คือ ตัวประมาณความแปรปรวนรวม โดยที่

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2} \quad (2)$$

2.1.2 ถ้า $\sigma_1^2 \neq \sigma_2^2$ ค่าสถิติทดสอบที่สามารถคำนวณได้ดังนี้

$$t^* = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}, \quad \nu = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} \right)^2}{\left(\frac{1}{n_1 - 1} \right) \left(\frac{s_1^2}{n_1} \right)^2 + \left(\frac{1}{n_2 - 1} \right) \left(\frac{s_2^2}{n_2} \right)^2} \quad (3)$$

2.2 สถิติทดสอบผลบวกลำดับที่ของวิลค็อกซัน (Sinsomboonthong, 2020)

การทดสอบผลบวกลำดับที่ของวิลค็อกซันมีวิธีการคือ นำตัวอย่างทั้งสองกลุ่มมารวมเข้าด้วยกัน จากนั้นทำการเรียงลำดับตัวอย่างทั้งสองกลุ่มจากค่าน้อยที่สุดไปหาค่ามากที่สุดแล้วให้ค่าลำดับที่ โดยกำหนดให้ $y_1, y_2, \dots, y_{n_2}$ เป็นค่าที่ได้จากตัวอย่างกลุ่มที่มีขนาดน้อยกว่า (หรือเท่ากับ) ที่ถูกเรียงลำดับเรียบร้อยแล้ว และ $S_1, S_2, \dots, S_{n_2}$ คือค่าลำดับที่ของ $y_1, y_2, \dots, y_{n_2}$ ดังนั้น ค่าสถิติทดสอบผลบวกลำดับที่ของวิลค็อกซันสามารถคำนวณได้ดังนี้

$$w^* = \sum_{i=1}^{n_2} S_i \quad (4)$$

เมื่อ $n_2 \leq n_1$

3. เปรียบเทียบค่าสถิติทดสอบที่ได้ในข้อ 2 กับค่าวิกฤต (critical values) ที่ได้จากการคำนวณ โดยใช้โปรแกรมทางสถิติ จากนั้นทำการสรุปผลการทดสอบตามสมมุติฐานที่กำหนดไว้ในข้อ 1 ว่าปฏิเสธหรือยอมรับสมมุติฐานว่าง (H_0)

3.1 การทดสอบที่

ถ้า $t^* \geq t_{\alpha/2, v}$ หรือ $t^* \leq -t_{\alpha/2, v}$ ปฏิเสธสมมุติฐานว่าง (H_0)

3.2 การทดสอบผลบวกลำดับที่ของวิลค็อกซัน

ถ้า $w^* \geq w_{\alpha/2}$ หรือ $w^* \leq n_2(n_1 + n_2 + 1) - w_{\alpha/2}$ ปฏิเสธสมมุติฐานว่าง (H_0)

4. ดำเนินการตามข้อ 1-3 ซ้ำ จำนวน 10,000 ครั้งในแต่ละสถานการณ์

5. นับจำนวนครั้งในการปฏิเสธสมมุติฐานว่าง

6. คำนวณค่าประมาณความน่าจะเป็นของความผิดพลาดแบบที่ 1 และกำลังการทดสอบของทั้งสองการทดสอบ โดยค่าประมาณความน่าจะเป็นของความผิดพลาดแบบที่ 1 ($\hat{\alpha}$) คำนวณได้จาก

$$\hat{\alpha} = \frac{\text{จำนวนครั้งในการปฏิเสธสมมุติฐานว่างเมื่อสมมุติฐานว่างเป็นจริง}}{\text{จำนวนครั้งในการทำซ้ำ}} \quad (5)$$

และค่าประมาณกำลังการทดสอบ ($1 - \hat{\beta}$) คำนวณได้จาก

$$1 - \hat{\beta} = \frac{\text{จำนวนครั้งในการปฏิเสธสมมุติฐานว่างเมื่อสมมุติฐานทางเลือกเป็นจริง}}{\text{จำนวนครั้งในการทำซ้ำ}} \quad (6)$$

7. เปรียบเทียบค่าประมาณความน่าจะเป็นของความผิดพลาดแบบที่ 1 ของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันกับเกณฑ์ของ Cochran (Cochran, 1947) ที่ระดับนัยสำคัญ 0.05 ถ้าค่าประมาณความน่าจะเป็นของความผิดพลาดแบบที่ 1 ตกอยู่ในช่วง $[0.04, 0.06]$ จะสรุปว่า การทดสอบนั้นสามารถควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 ได้

8. เปรียบเทียบกำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซัน ถ้าการทดสอบใดสามารถควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 ได้ และให้กำลังการทดสอบที่มากกว่า จะสรุปว่า การทดสอบนั้นมีประสิทธิภาพมากกว่า

ผลการวิจัย

ผลการเปรียบเทียบประสิทธิภาพของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับประชากร 2 กลุ่มที่เป็นอิสระกันเมื่อข้อมูลเป็นจำนวนนับ แสดงดังตารางที่ 1-4

ความน่าจะเป็นของความผิดพลาดแบบที่ 1

จากตารางที่ 1 และตารางที่ 2 พบว่า การทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงทวินามและการแจกแจงปัวซองสามารถควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 ได้ทุกสถานการณ์ แม้ในกรณีที่ตัวอย่างทั้ง 2 กลุ่มมีขนาดเล็ก ทั้งในกรณีที่ขนาดตัวอย่างทั้งสองกลุ่มเท่ากันและไม่เท่ากัน นอกจากนี้ การทดสอบทีให้ค่าประมาณความน่าจะเป็นของความผิดพลาดแบบที่ 1 สูงกว่าการทดสอบผลบวกลำดับที่ของวิลค็อกซันเกือบทุกกรณี

ตารางที่ 1 ค่าประมาณความน่าจะเป็นของความผิดพลาดแบบที่ 1 ของการทดสอบที (t-test) และการทดสอบผลบวกลำดับที่ของวิลค็อกซัน (Wilcoxon rank sum test) สำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงทวินาม

พารามิเตอร์			การทดสอบ	ขนาดตัวอย่าง (n_1, n_2)			ขนาดตัวอย่างไม่เท่ากัน (n_1, n_2)		
n	p_1	p_2		(10, 10)	(60, 60)	(100, 100)	(20, 100)	(30, 90)	(40, 80)
10	0.1	0.1	t-test	0.0500	0.0487	0.0518	0.0526	0.0560	0.0511
			Wilcoxon rank sum test	0.0477	0.0509	0.0487	0.0503	0.0526	0.0506
	0.5	0.5	t-test	0.0503	0.0511	0.0527	0.0536	0.0528	0.0527
			Wilcoxon rank sum test	0.0453	0.0505	0.0511	0.0513	0.0502	0.0512
	0.7	0.7	t-test	0.0522	0.0505	0.0529	0.0530	0.0522	0.0519
			Wilcoxon rank sum test	0.0451	0.0478	0.0520	0.0511	0.0515	0.0507

ตารางที่ 1 (ต่อ)

พารามิเตอร์			การทดสอบ	ขนาดตัวอย่าง (n_1, n_2)			ขนาดตัวอย่างไม่เท่ากัน (n_1, n_2)		
n	p_1	p_2		(10, 10)	(60, 60)	(100, 100)	(20, 100)	(30, 90)	(40, 80)
50	0.1	0.1	t-test	0.0500	0.0500	0.0539	0.0533	0.0552	0.0513
			Wilcoxon rank sum test	0.0445	0.0503	0.0552	0.0534	0.0520	0.0488
	0.5	0.5	t-test	0.0517	0.0499	0.0517	0.0526	0.0551	0.0526
			Wilcoxon rank sum test	0.0540	0.0482	0.0540	0.0530	0.0531	0.0491
	0.7	0.7	t-test	0.0499	0.0501	0.0529	0.0523	0.0541	0.0512
			Wilcoxon rank sum test	0.0427	0.0496	0.0522	0.0521	0.0518	0.0494
100	0.1	0.1	t-test	0.0498	0.0509	0.0524	0.0519	0.0553	0.0520
			Wilcoxon rank sum test	0.0429	0.0507	0.0514	0.0526	0.0529	0.0495
	0.5	0.5	t-test	0.0490	0.0485	0.0519	0.0546	0.0529	0.0459
			Wilcoxon rank sum test	0.0420	0.0509	0.0533	0.0527	0.0503	0.0503
	0.7	0.7	t-test	0.0513	0.0465	0.0532	0.0518	0.0507	0.0540
			Wilcoxon rank sum test	0.0462	0.0491	0.0521	0.0528	0.0496	0.0525

ตารางที่ 2 ค่าประมาณความน่าจะเป็นของความผิดพลาดแบบที่ 1 ของการทดสอบที่ (t-test) และการทดสอบผลบวกลำดับที่ของวิลค็อกซัน (Wilcoxon rank sum test) สำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงปัวซอง

พารามิเตอร์		การทดสอบ	ขนาดตัวอย่าง (n_1, n_2)			ขนาดตัวอย่างไม่เท่ากัน (n_1, n_2)		
λ_1	λ_2		(10, 10)	(60, 60)	(100, 100)	(20, 100)	(30, 90)	(40, 80)
1	1	t-test	0.0501	0.0500	0.0524	0.0547	0.0560	0.0536
		Wilcoxon rank sum test	0.0484	0.0503	0.0509	0.0500	0.0528	0.0524
5	5	t-test	0.0486	0.0481	0.0533	0.0519	0.0544	0.0508
		Wilcoxon rank sum test	0.0431	0.0503	0.0541	0.0526	0.0523	0.0486
10	10	t-test	0.0480	0.0491	0.0499	0.0511	0.0541	0.0507
		Wilcoxon rank sum test	0.0426	0.0477	0.0494	0.0514	0.0530	0.0521

กำลังการทดสอบ

จากตารางที่ 3 และตารางที่ 4 พบว่า ในกรณีที่ขนาดตัวอย่างทั้งสองกลุ่มเท่ากัน ($n_1 = n_2$) กำลังการทดสอบของการทดสอบที่และการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงทวินามและการแจกแจงปัวซองมีแนวโน้มเพิ่มขึ้นเมื่อขนาดตัวอย่างทั้งสองกลุ่ม และขนาดอิทธิพลเพิ่มขึ้น ดังรูปภาพ 3-4 นอกจากนี้ การทดสอบที่ให้กำลังการทดสอบสูงกว่าการทดสอบผลบวกลำดับที่ของวิลค็อกซันทุกกรณี

ตารางที่ 3 กำลังการทดสอบของการทดสอบที่ (t-test) และการทดสอบผลบวกลำดับที่ของวิลค็อกซัน (Wilcoxon rank sum test) สำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงทวินาม

พารามิเตอร์			ขนาด อิทธิพล ^d	การทดสอบ	ขนาดตัวอย่างเท่ากัน (n_1, n_2)			ขนาดตัวอย่างไม่เท่ากัน (n_1, n_2)		
n	p_1	p_2			(10, 10)	(60, 60)	(100, 100)	(20, 100)	(30, 90)	(40, 80)
10	0.1	0.125	0.1771	t-test	0.0826	0.2720	0.4115	0.1717	0.2095	0.2476
								0.1746*	0.2215*	0.2512*
				Wilcoxon rank sum test	0.0763	0.2583	0.3939	0.1618	0.2022	0.2365
								0.1781*	0.2174*	0.2469*
	0.5	0.525	0.1119	t-test	0.0681	0.1432	0.2049	0.1018	0.1178	0.1292
								0.1007*	0.1163*	0.1299*
				Wilcoxon rank sum test	0.0612	0.1364	0.1951	0.0981	0.1115	0.1247
								0.0964*	0.1092*	0.1212*
	0.7	0.725	0.1236	t-test	0.0660	0.1658	0.2345	0.1158	0.1318	0.1537
								0.1100*	0.1288*	0.1471*
				Wilcoxon rank sum test	0.0582	0.1590	0.2285	0.1096	0.1256	0.1453
								0.1074*	0.1265*	0.1435*
50	0.1	0.125	0.3959	t-test	0.2190	0.8526	0.9729	0.6146	0.7443	0.8140
								0.6190*	0.7517*	0.8136*
				Wilcoxon rank sum test	0.1978	0.8358	0.9667	0.5964	0.7249	0.7923
								0.5893*	0.7265*	0.7948*
	0.5	0.525	0.2502	t-test	0.1219	0.4816	0.7065	0.2966	0.3822	0.4395
								0.2941*	0.3824*	0.4405*
				Wilcoxon rank sum test	0.1033	0.4652	0.6892	0.2826	0.3616	0.4174
								0.2831*	0.3645*	0.4192*

ตารางที่ 3 (ต่อ)

พารามิเตอร์			ขนาด อิทธิพล ^d	การทดสอบ	ขนาดตัวอย่างเท่ากัน (n_1, n_2)			ขนาดตัวอย่างไม่เท่ากัน (n_1, n_2)		
n	p_1	p_2			(10, 10)	(60, 60)	(100, 100)	(20, 100)	(30, 90)	(40, 80)
100	0.7	0.725	0.2763	t-test	0.1311	0.5698	0.7925	0.3570	0.4549	0.5137
								0.3516*	0.4482*	0.5198*
				Wilcoxon rank	0.1143	0.5508	0.7777	0.3359	0.4334	0.4937
	0.5	0.525	0.3538	sum test				0.3401*	0.4290*	0.4947*
				t-test	0.3879	0.9894	0.9998	0.8870	0.9596	0.9818
								0.8871*	0.9570*	0.9824*
100	0.7	0.725	0.3907	Wilcoxon rank	0.3517	0.9873	0.9995	0.8756	0.9507	0.9777
				sum test				0.8678*	0.9494*	0.9770*
				t-test	0.1854	0.7740	0.9449	0.5317	0.6536	0.7243
	0.5	0.525	0.3538					0.5266*	0.6529*	0.7218*
				Wilcoxon rank	0.1619	0.7507	0.9329	0.5078	0.6273	0.7044
				sum test				0.5039*	0.6280*	0.6983*
100	0.7	0.725	0.3907	t-test	0.2065	0.8491	0.9741	0.6052	0.7385	0.8034
								0.5994*	0.7299*	0.8050*
				Wilcoxon rank	0.1893	0.8303	0.9671	0.5796	0.7140	0.7840
	0.5	0.525	0.3538	sum test				0.5797*	0.7054*	0.7793*
				t-test	0.1854	0.7740	0.9449	0.5317	0.6536	0.7243
								0.5266*	0.6529*	0.7218*

* ค่าประมาณกำลังการทดสอบของขนาดตัวอย่าง (n_1, n_2) เท่ากับ (100, 20) (90, 30) และ (80, 40)

^d ขนาดอิทธิพล (Effect size) คำนวณจาก $\left| (np_1 - np_2) / \sqrt{np_1(1-p_1) + np_2(1-p_2)} \right|$

เมื่อพิจารณาอิทธิพลของอัตราส่วนของขนาดตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ที่มีต่อกำลังการทดสอบภายใต้ขนาดตัวอย่างรวมที่เท่ากัน พบว่า สำหรับข้อมูลที่มีการแจกแจงทวินาม โดยที่ $p_1 < p_2$ เมื่ออัตราส่วนของ $n_1 : n_2$ เท่ากับ 1 : 1 ยกเว้นกรณีที่ n_1 และ n_2 เท่ากับ 10 กำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันจะมีค่าสูงสุด และเมื่ออัตราส่วนของ $n_1 : n_2$ ลดลงเป็น 1 : 2, 1 : 3 และ 1 : 5 กำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันจะลดลงตามลำดับ ดังรูปภาพ 5-6 ในทำนองเดียวกัน ถ้า $p_1 < p_2$ และ $n_1 > n_2$ พบว่า เมื่ออัตราส่วนของ $n_1 : n_2$ เพิ่มขึ้นจาก 1 : 1 เป็น 2 : 1, 3 : 1 และ 5 : 1 กำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันก็จะลดลงตามลำดับ

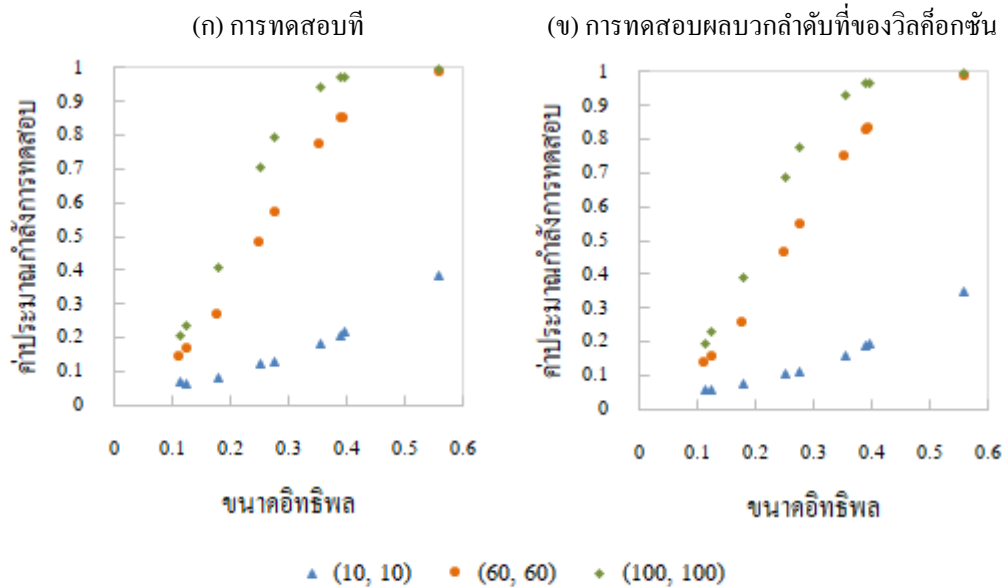
ตารางที่ 4 กำลังการทดสอบของการทดสอบที (t-test) และการทดสอบผลบวกลำดับที่ของวิลค็อกซัน (Wilcoxon rank sum test) สำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงปัวซง

พารามิเตอร์ λ_1	ขนาด λ_2 อิทธิพล ^d		การทดสอบ	ขนาดตัวอย่างเท่ากัน (n_1, n_2)			ขนาดตัวอย่างไม่เท่ากัน (n_1, n_2)		
				(10, 10)	(60, 60)	(100, 100)	(20, 100)	(30, 90)	(40, 80)
1	1.25	0.1667	t-test	0.0766	0.2418	0.3728	0.1624	0.1970	0.2284
							0.1538*	0.1914*	0.2270*
			Wilcoxon rank sum test	0.0740	0.2287	0.3488	0.1453	0.1803	0.2090
1	1.50	0.3162	t-test	0.1567	0.6765	0.8772	0.4283	0.5505	0.6329
							0.4171*	0.5448*	0.6206*
			Wilcoxon rank sum test	0.1454	0.6446	0.8500	0.4062	0.5153	0.5999
1	2	0.5774	t-test	0.4048	0.9937	0.9999	0.9104	0.9737	0.9893
							0.8817*	0.9630*	0.9861*
			Wilcoxon rank sum test	0.3780	0.9910	0.9998	0.9000	0.9656	0.9848
5	8	0.8321	t-test	0.7068	1.000	1.000	0.9982	0.9999	1.000
							0.8679*	0.9999*	1.000*
			Wilcoxon rank sum test	0.6707	1.000	1.000	0.9974	0.9997	1.000
5	8	0.8321	t-test	0.7068	1.000	1.000	0.9982	0.9999	1.000
							0.8679*	0.9999*	1.000*
			Wilcoxon rank sum test	0.6707	1.000	1.000	0.9974	0.9997	1.000

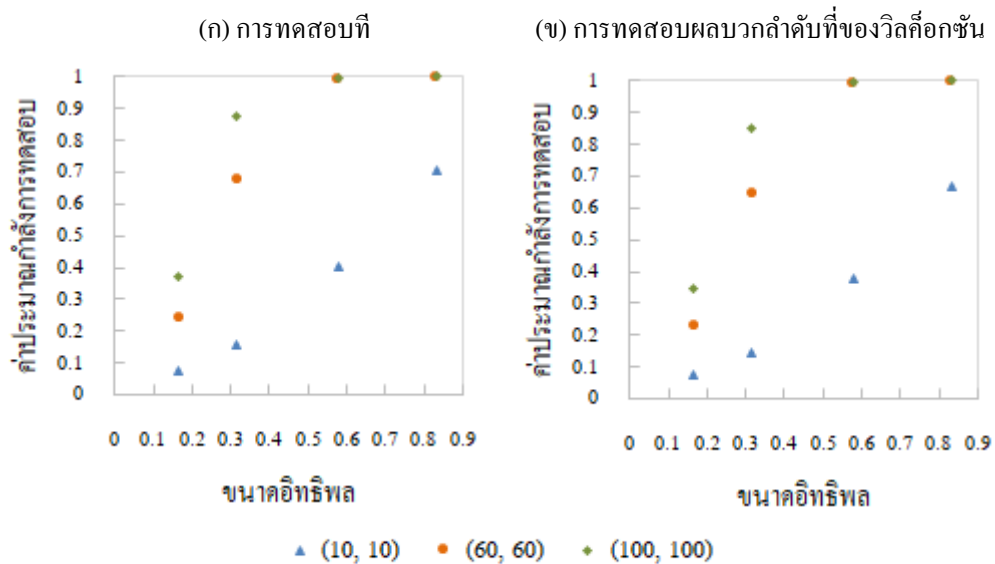
* ค่าประมาณกำลังการทดสอบของขนาดตัวอย่าง (n_1, n_2) เท่ากับ (100, 20) (90, 30) และ (80, 40)

^d ขนาดอิทธิพล (Effect size) คำนวณจาก $|(\lambda_1 - \lambda_2) / \sqrt{\lambda_1 + \lambda_2}|$

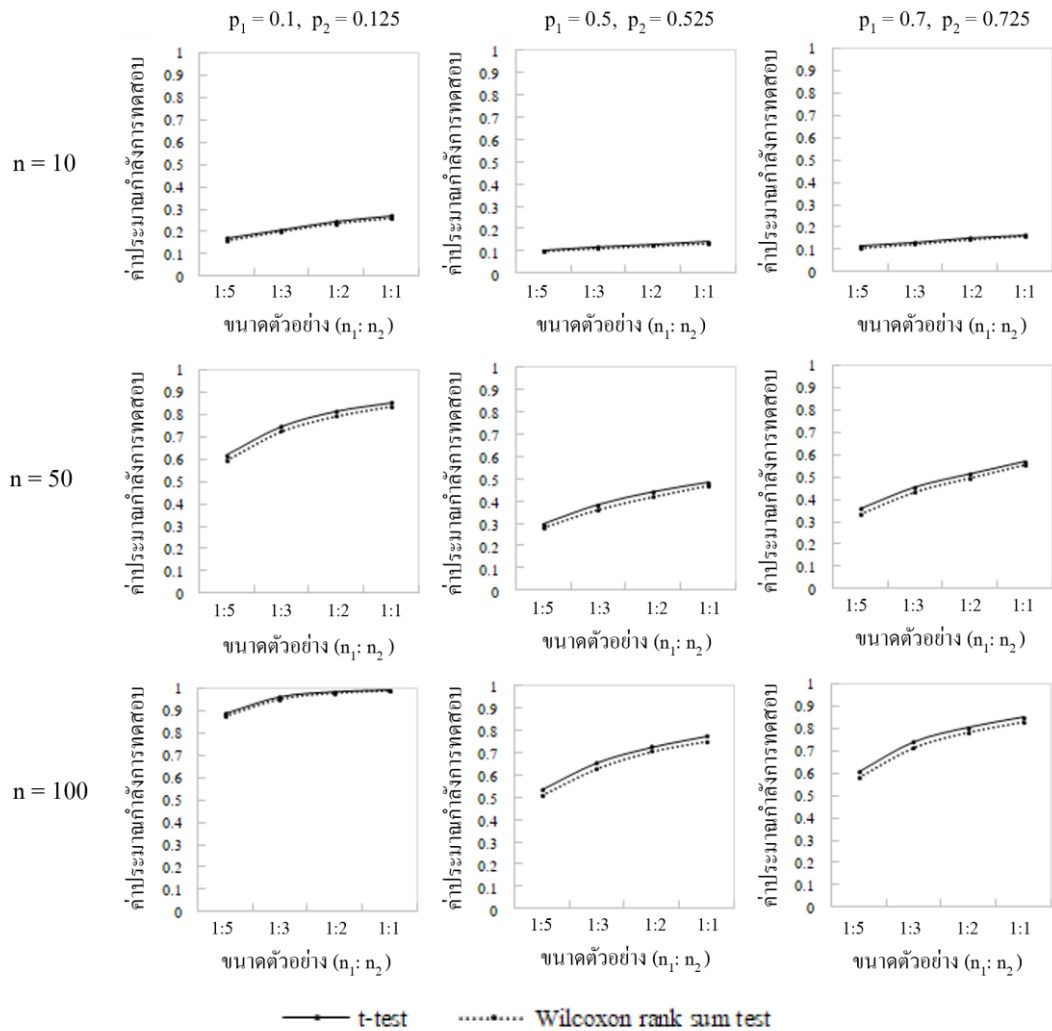
สำหรับข้อมูลที่มีการแจกแจงปัวซง เมื่อ $\lambda_1 < \lambda_2$ และ $n_1 < n_2$ (หรือ $n_1 > n_2$) พบว่า กำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันจะมีค่าสูงสุดเมื่ออัตราส่วนของ $n_1 : n_2$ เท่ากับ 1 : 1 เช่นเดียวกับข้อมูลที่มีการแจกแจงทวินาม นอกจากนี้ การทดสอบทีให้กำลังการทดสอบสูงกว่าการทดสอบผลบวกลำดับที่ของวิลค็อกซันทุกกรณี



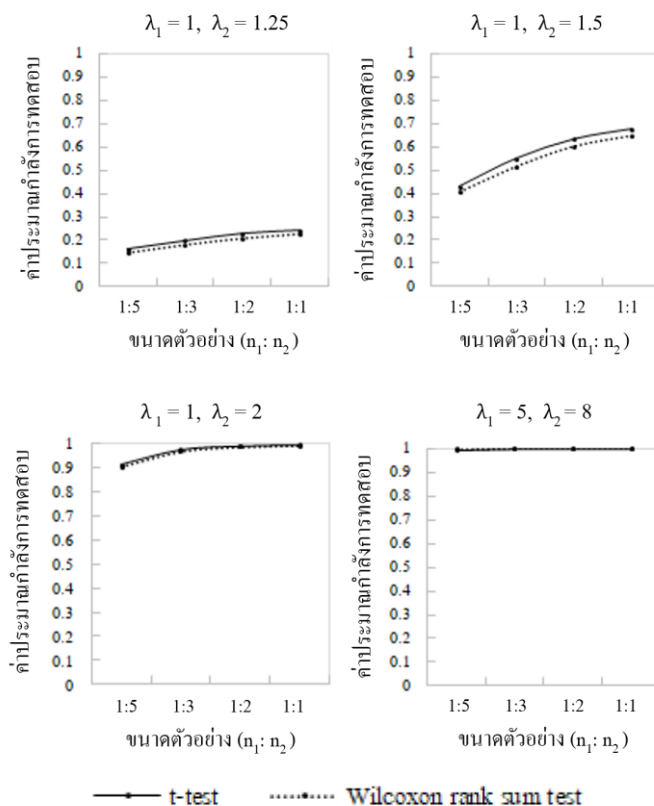
รูปที่ 3 ความสัมพันธ์ระหว่างขนาดอิทธิพลและค่าประมาณกำลังการทดสอบสำหรับข้อมูลที่มีการแจกแจงทวินาม



รูปที่ 4 ความสัมพันธ์ระหว่างขนาดอิทธิพลและค่าประมาณกำลังการทดสอบสำหรับข้อมูลที่มีการแจกแจงปัวซอง



รูปที่ 5 ค่าประมาณกำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซัน สำหรับทดสอบค่ากลางของข้อมูลที่มีการแจกแจงทวินาม กรณีขนาดตัวอย่าง 2 กลุ่มไม่เท่ากัน



รูปที่ 6 ค่าประมาณกำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซัน สำหรับทดสอบค่ากลางของข้อมูลที่มีการแจกแจงปัวซง กรณีขนาดตัวอย่าง 2 กลุ่มไม่เท่ากัน

สรุปและวิจารณ์ผลการวิจัย

การวิจัยในครั้งนี้มีวัตถุประสงค์เพื่อเปรียบเทียบประสิทธิภาพของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับทดสอบความแตกต่างระหว่างค่ากลางของประชากรสองกลุ่มที่เป็นอิสระกันเมื่อข้อมูลเป็นจำนวนนับ โดยการแจกแจงที่นำมาศึกษา คือ การแจกแจงทวินาม และการแจกแจงปัวซง การเปรียบเทียบประสิทธิภาพของการทดสอบพิจารณาจากความสามารถในการควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 และกำลังการทดสอบ โดยกำหนดระดับนัยสำคัญที่ 0.05

สำหรับความสามารถในการควบคุมความน่าจะเป็นของการเกิดความผิดพลาดแบบที่ 1 ของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงทวินามและการแจกแจงปัวซงพบว่า การทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสามารถควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 ได้ทุกกรณีที่ศึกษา แม้ในกรณีที่ตัวอย่างทั้ง 2 กลุ่มมีขนาดเล็ก

และในกรณีที่ขนาดตัวอย่างทั้งสองกลุ่มเท่ากันและไม่เท่ากัน ซึ่งผลการวิจัยที่ได้สอดคล้องกับงานวิจัยของ Araveeporn et al. (2017) ที่ได้ทำการศึกษาประสิทธิภาพของการทดสอบซี (z-test) การทดสอบที (t-test) การทดสอบแบบสุ่ม (randomization test) และการทดสอบแมนน์-วิทนียู (Mann-Whitney U test) สำหรับทดสอบค่าเฉลี่ยหรือค่ากลางของประชากร 2 กลุ่มที่เป็นอิสระกัน โดยผลการวิจัยพบว่า การทดสอบซี การทดสอบที และการทดสอบแมนน์-วิทนียูสามารถควบคุมความน่าจะเป็นของความผิดพลาดแบบที่ 1 ได้ทั้งในกรณีตัวอย่างขนาดเล็กและตัวอย่างขนาดใหญ่

เมื่อพิจารณากำลังการทดสอบพบว่า ในกรณีที่ขนาดตัวอย่างทั้งสองกลุ่มเท่ากัน กำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันสำหรับทดสอบความแตกต่างระหว่างค่ากลางของข้อมูลที่มีการแจกแจงทวินามและการแจกแจงปัวซองมีแนวโน้มเพิ่มขึ้นเมื่อขนาดตัวอย่างทั้งสองกลุ่ม และขนาดอิทธิพลเพิ่มขึ้น เมื่อพิจารณาอิทธิพลของอัตราส่วนของขนาดตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ที่มีต่อกำลังการทดสอบภายใต้ขนาดตัวอย่างรวมที่เท่ากันพบว่า กำลังการทดสอบของการทดสอบทีและการทดสอบผลบวกลำดับที่ของวิลค็อกซันจะมีค่าสูงสุดเมื่ออัตราส่วนของขนาดตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ($n_1 : n_2$) เท่ากับ 1 : 1 ยกเว้นกรณีที่ n_1 และ n_2 เท่ากับ 10 และจะมีค่าน้อยที่สุดเมื่ออัตราส่วนของขนาดตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ($n_1 : n_2$) เท่ากับ 1 : 5 (หรือ 5 : 1) ทั้งในกรณีที่ข้อมูลมีการแจกแจงทวินามและกรณีที่ข้อมูลมีการแจกแจงปัวซอง ซึ่งผลการวิจัยที่ได้สอดคล้องกับงานวิจัยของ Kim & Park (2019) ที่ได้ทำการศึกษาเกี่ยวกับข้อตกลงเบื้องต้นของการทดสอบทีสำหรับประชากร 2 กลุ่มที่เป็นอิสระกันและพบว่า ถ้าอัตราส่วนของขนาดตัวอย่างกลุ่มที่ 1 และกลุ่มที่ 2 ($n_1 : n_2$) เท่ากับ 1 : 1 จะให้กำลังการทดสอบที่สูงที่สุด คือ 0.88 รองลงมาคืออัตราส่วนของ $n_1 : n_2$ เท่ากับ 5 : 3, 3 : 1, และ 7 : 1 ซึ่งให้กำลังการทดสอบ 0.86, 0.78, และ 0.55 ตามลำดับ นอกจากนี้ในการวิจัยครั้งนี้ยังพบว่า การทดสอบทีซึ่งเป็นการทดสอบแบบอิงพารามิเตอร์ให้กำลังการทดสอบสูงกว่าการทดสอบผลบวกลำดับที่ของวิลค็อกซันทุกกรณี

จากผลการวิจัยในครั้งนี้จะเห็นได้ว่า สำหรับข้อมูลที่เป็นจำนวนนับที่มีการแจกแจงทวินามและการแจกแจงปัวซอง การทดสอบทีเป็นการทดสอบความแตกต่างระหว่างค่ากลางของประชากรสองกลุ่มที่เป็นอิสระกันที่มีประสิทธิภาพดีกว่าการทดสอบวิลค็อกซัน เนื่องจากสามารถควบคุมความน่าจะเป็นของการเกิดความผิดพลาดแบบที่ 1 ได้ และมีกำลังการทดสอบที่สูงกว่าถึงแม้ว่าข้อมูลจะไม่เป็นไปตามข้อตกลงเบื้องต้นของการทดสอบที กล่าวคือข้อมูลไม่ได้มาจากประชากรที่มีการแจกแจงปกติและไม่ได้มีมาตรการควบคุมตัวอยู่ในมาตราช่วงเป็นอย่างน้อย ทั้งนี้สาเหตุอาจเนื่องมาจากการทดสอบทีเป็นการทดสอบที่มีความแข็งแกร่งต่อการละเมิดข้อตกลงเบื้องต้น

เอกสารอ้างอิง

Araveeporn, A., Chuenarrom, C., Fuksab, Supawee Choedchitkusol, S., & Muangmeesup, A. (2017). An efficiency comparison of z-test, t-test, the randomization test and Mann-Whitney U test for testing two independent population means or medians. *Journal of Science Ladkrabang*, 26(2), 73-88.

- Blair, R. C., & Higgins, J. J. (1980). A comparison of the power of the Wilcoxon's rank-sum statistic to that of Student's t statistic under various non-normal distributions. *Journal of Educational Statistics*, 5(4), 309–335.
- Bridge, P. D., & Sawilowsky, S. S. (1999). Increasing physicians' awareness of the impact of statistics on research outcomes: comparative power of the t-test and Wilcoxon rank sum rank-sum test in small samples applied research. *Journal of clinical epidemiology*, 52(3), 229-235.
- Cochran, W. G. (1947). Some consequences when the assumptions for the analysis of variance are not satisfied. *Biometrics*, 3, 22-38.
- Kim, T. K., & Park, J. H. (2019). More about the basic assumptions of t-test: normality and sample size. *Korean journal of anesthesiology*, 72(4), 331-335.
- Moore, D. S., McCabe, G. P., & Craig, B. A. (2014). *Introduction to the Practice of Statistics* (8th ed.). New York, USA: W. H. Freeman and Company.
- Sinsomboonthong, S. (2020). *Nonparametric Statistics*. Bangkok, Thailand: Chamchuree Products Co., LTD.
- Sangthong, M. (2013). A study of efficiency of parametric and nonparametric statistics in testing of central difference between two populations. *KKU Science Journal*, 41(1), 226-238.
- Sangthong, M. (2014). Robustness and power of the test of parametric and nonparametric statistics in testing of central difference between two populations for Likert-type data 5 point. *Thai Science and Technology Journal*, 22(5), 605-619.

Research Article

การเปรียบเทียบตัวแบบพยากรณ์อนุกรมเวลาระหว่างตัวแบบผสมกับตัวแบบเดี่ยวเพื่อการพยากรณ์จำนวนสายโทรเข้ารายวัน

A Comparison of Time Series Forecasting Models between Hybrid Model and Individual Model for Forecasting Daily Incoming Call Volume

พรพิมล ชัยวุฒิศักดิ์^{1*}

Pornpimol Chaiwuttisak^{1*}

¹ภาควิชาสถิติ คณะวิทยาศาสตร์ สถาบันเทคโนโลยีพระจอมเกล้าเจ้าคุณทหารลาดกระบัง กรุงเทพฯ 10520 ประเทศไทย

¹Department of Statistics, Faculty of Science, King Mongkut's Institute of Technology Ladkrabang, Bangkok 10520, Thailand

*E-mail: pornpimol.ch@kmitl.ac.th

Received: 27/06/2021; Revised: 19/12/2021; Accepted: 28/03/2022

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อเปรียบเทียบความถูกต้องของวิธีการพยากรณ์จำนวนสายโทรเข้ารายวัน โดยวิธีการพยากรณ์ที่ทำการศึกษา 3 วิธี ได้แก่ 1) วิธีบ็อกซ์-เจนกินส์ประกอบด้วยตัวแบบ SARIMA และตัวแบบ SARIMAX 2) วิธีโครงข่ายประสาทเทียม และ 3) วิธีผสมระหว่างตัวแบบ SARIMAX และโครงข่ายประสาทเทียม ข้อมูลที่นำมาใช้ในการศึกษาเป็นข้อมูลอนุกรมเวลารายวันของจำนวนสายโทรเข้าไปที่ศูนย์บริการลูกค้าของบริษัทแห่งหนึ่ง ซึ่งแบ่งข้อมูลเป็น 2 ชุด คือ ข้อมูลในอดีตตั้งแต่เดือนมกราคม พ.ศ. 2559 ถึง เดือนธันวาคม พ.ศ. 2561 เพื่อคัดเลือกตัวแบบที่เหมาะสมที่สุด และ ข้อมูลในอดีตตั้งแต่เดือนมกราคม พ.ศ. 2562 ถึง เดือนธันวาคม พ.ศ. 2562 เพื่อเปรียบเทียบความแม่นยำของค่าพยากรณ์ที่ได้ โดยเกณฑ์ที่ใช้ในการเปรียบเทียบความคลาดเคลื่อนของการพยากรณ์คือ ค่าเฉลี่ยของค่าสัมบูรณ์เปอร์เซ็นต์ความคลาดเคลื่อน (MAPE) ผลการศึกษาพบว่า ตัวแบบที่มีค่า MAPE ต่ำสุดคือตัวแบบผสมระหว่าง SARIMAX และ โครงข่ายประสาทเทียม ซึ่งเท่ากับ 24.02% ซึ่งน้อยกว่าตัวแบบ SARIMAX ตัวแบบโครงข่ายประสาทเทียม และตัวแบบ SARIMA ซึ่งมีค่า MAPE เท่ากับ 24.06%, 43.70%,

44.97% ตามลำดับ จะเห็นได้ว่าตัวแบบผสมมีความแม่นยำในการพยากรณ์มากกว่าตัวแบบเดี่ยว โดยตัวแบบผสมที่สามารถนำไปพยากรณ์จำนวนสายโทรเข้ารายวันล่วงหน้าสำหรับใช้เป็นข้อมูลประกอบการวางแผนการกำลังคนทำงานของศูนย์บริการข้อมูลลูกค้าที่เหมาะสมในอนาคต

คำสำคัญ การพยากรณ์, จำนวนสายโทรเข้ารายวัน, วิธีบ็อกซ์-เจนกินส์, โครงข่ายประสาทเทียม, ตัวแบบผสม

Abstract

The research aimed to compare the accuracy of forecasting methods for daily incoming call volume. There were three forecasting methods that were used in the study: Box-Jenkins method with SRIMA model and SARIMAX model, Artificial Neural Network (ANN) model and the hybrid model combining SARIMAX and Artificial Neural Network (SARIMAX-ANN) model. The data used in this study is time series of daily incoming call volume to call center which can be divided into 2 data sets. The first data set which was the past data from January 2016 to December 2018 were used for selecting of the most suitable model and the second data set was the past data from January 2019 to December 2019 for the comparison of the accuracy of forecasting model by using Mean Absolute Percentage Error (MAPE). The results showed model with the lowest MAPE is hybrid model of SARIMAX-ANN (MAPE = 24.02%), while the MAPE values for SARIMAX, ANN and SARIMA were 24.06%, 43.70%, and 44.97% respectively. It indicates that the hybrid model is more accurate in forecasting than individual model. The hybrid model can be used to forecast daily incoming call volume which is supporting information for the suitable workforce planning of customer service center in the future.

Keywords: Forecasting, Daily Incoming Call Volume, Box-Jenkins, Artificial Neural Network, Hybrid model.

บทนำ

ปัจจุบันการดำเนินธุรกิจต่างๆ มักเน้นการสร้างภาพพจน์ให้แกลูกค้า โดยแนวทางหนึ่งในการสร้างความพึงพอใจคือการที่ลูกค้าสามารถติดต่อกับหน่วยงานหรือองค์กรได้ทุกที่ ทุกเวลา ดังนั้นการจัดตั้งศูนย์บริการข้อมูลลูกค้า หรือศูนย์ลูกค้าสัมพันธ์ ที่เรียกกันโดยทั่วไปว่า Call Center จะถูกใช้เป็นหนึ่งช่องทางในการตอบสนองความต้องการของลูกค้าในการติดต่อเข้ามาสอบถามข้อมูล ไม่ว่าจะเป็นเรื่องของผลิตภัณฑ์ หรือเมื่อต้องการความช่วยเหลือในเรื่องต่างๆ เพื่อให้ลูกค้าเกิดความพึงพอใจสูงสุดในบริการที่ได้รับ

บริษัทแห่งหนึ่งในธุรกิจประกันชีวิตมีศูนย์บริการข้อมูลลูกค้า (Call Center) ที่ถูกจัดตั้งขึ้นเพื่อเป็นการให้ข้อมูลในด้านต่างๆ แก่ลูกค้า เช่น ด้านผลิตภัณฑ์ หรือ ด้านการใช้งานผลิตภัณฑ์ แนะนำสินค้าและบริการ ประสานงานและช่วยเหลือในการแก้ไขปัญหาที่เกิดขึ้นของผู้ใช้บริการ เพื่อให้ตรงตามความต้องการและความพึงพอใจของลูกค้า โดยมีการให้บริการตลอด 24 ชั่วโมง ทั้งนี้ทางศูนย์บริการข้อมูลลูกค้ามีความประสงค์ที่จะพัฒนาและเพิ่มประสิทธิภาพในการให้บริการแก่ลูกค้า โดยอาศัยตัวแบบพยากรณ์จำนวนสายโทรเข้ารายวันสำหรับเป็นเครื่องมือหนึ่งที่จะช่วยในการวางแผนการปฏิบัติงานในด้านต่างๆ ในอนาคต รวมทั้งลดระยะเวลาการรอคอยในการรับบริการของลูกค้าได้ ซึ่งเป็นการส่งเสริมภาพลักษณ์และความเชื่อมั่นของลูกค้าที่มีต่อบริษัท (Ibrahim et al., 2016)

Taylor (2008) ศึกษาการเปรียบเทียบการพยากรณ์อนุกรมเวลาของปริมาณสายที่โทรเข้า ของศูนย์บริการลูกค้า โดยศึกษาจากกลุ่มการให้บริการ 5 กลุ่ม ของธนาคารพาณิชย์รายย่อยแห่งหนึ่งของสหราชอาณาจักร โดยทำการเปรียบเทียบ 2 วิธี คือ วิธีการของบ็อกซ์-เจนกินส์ และ วิธีปรับให้เรียบแบบเอ็กซ์โปเนนเชียลแบบไฮลท์-วินเทอร์ ผลการวิจัยพบว่าวิธีการของบ็อกซ์-เจนกินส์ ให้ความแม่นยำในการพยากรณ์ดีกว่าเนื่องจากให้ค่า MAPE ต่ำที่สุด ในขณะที่ Njeru (2018) ศึกษาการพยากรณ์ปริมาณการโทรเข้าของลูกค้าในอนาคต โดยทำการพยากรณ์ปริมาณสายโทรเข้ามาของลูกค้ารายวันของศูนย์บริการทั้งหมดในประเทศและศูนย์บริการอุตสาหกรรมโทรคมนาคมทั่วไป ประเทศเคนย่า 2 รูปแบบ คือ การโทรรายวันที่รับบนระบบเสียงแบบโต้ตอบ (IVR Hits) และการโทรที่ย้ายจากระบบ IVR ไปเพื่อรอการบริการ (Offered) ด้วยตัวแบบบ็อกซ์-เจนกินส์ โดยลักษณะข้อมูลมีความแปรผันของฤดูกาลเข้ามาเกี่ยวข้อง จึงเลือกใช้วิธี SARIMA โดยทำการเลือกตัวแบบ SARIMA ที่เหมาะสม และทำการเปรียบเทียบค่าพยากรณ์ใช้เกณฑ์ค่าเฉลี่ยความคลาดเคลื่อนกำลังสอง (MSE) ผลวิจัยพบว่าตัวแบบที่เหมาะสมที่สุดสำหรับ IVR คือ $ARIMA(0,0,0) \times SARIMA(5,1,3)$, ส่วนตัวแบบที่เหมาะสมที่สุดสำหรับ Offered คือ $ARIMA(0,0,0) \times SARIMA(6,3,1)$, จากนั้นนำตัวแบบที่ดีที่สุดไปพยากรณ์ 30 วันข้างหน้า เพื่อนำมาใช้ในการให้ข้อมูลเชิงลึกสำหรับการจัดการแรงงานที่เหมาะสม

การศึกษาวิจัยครั้งนี้ได้ทำการพยากรณ์จำนวนสายโทรเข้ารายวัน เนื่องจากการพยากรณ์รายวันจะเป็นประโยชน์ต่อการตัดสินใจ การวางแผน การเตรียมตัวจากการคาดการณ์ในอนาคต ผู้วิจัยได้สุ่มข้อมูลจำนวนสายโทรเข้ารายวันจากฐานข้อมูลของบริษัท มาทำการศึกษาเพื่อสร้างตัวแบบสำหรับการพยากรณ์จำนวนสายโทรเข้า โดยทำการวิเคราะห์ลักษณะของข้อมูล พร้อมทั้งจัดการกับข้อมูลเพื่อให้ข้อมูลมีความเหมาะสมที่จะนำไปใช้ทำตัวแบบพยากรณ์ทางสถิติเพื่อหาตัวแบบที่เหมาะสมที่สุดสำหรับการพยากรณ์ข้อมูลจำนวนสายโทรเข้าล่วงหน้าของบริษัทประกันชีวิตแห่งนี้

วิธีการดำเนินงานวิจัย

1. ลักษณะและแหล่งที่มาของข้อมูล

ข้อมูลที่ใช้ในงานวิจัยครั้งนี้ คือ ข้อมูลอนุกรมเวลาจำนวนปริมาณสายโทรเข้าแบบสุ่มของบริษัทประกันแห่งหนึ่ง ซึ่งข้อมูลนี้ถูกเก็บเป็นรายการ (Transaction) ของแต่ละสายที่โทรเข้ามาทาง Call Center โดยบริษัทมีการเก็บข้อมูลดังกล่าวตั้งแต่ พ.ศ.2559-2562 ในระบบฐานข้อมูลด้วยโปรแกรมจัดการฐานข้อมูล Oracle

2. ขั้นตอนการดำเนินงาน

การพิจารณาและคัดเลือกตัวแบบการพยากรณ์จำนวนสายเรียกเข้ารายวันที่เหมาะสมจำเป็นต้องจัดการกับข้อมูลอย่างเหมาะสมและวิเคราะห์ข้อมูลเพื่อทำให้เกิดความเข้าใจถึงลักษณะของข้อมูล และมีการจัดการกับข้อมูลอย่างเหมาะสม จากนั้นจึงจะสามารถเลือกวิธีการสร้างตัวแบบการพยากรณ์ที่เหมาะสมกับลักษณะข้อมูลของงานวิจัยนี้ โดยขั้นตอนการวิจัยสามารถแบ่งออกเป็น 3 ขั้นตอน

1) การจัดการข้อมูล

ผู้วิจัยต้องจัดการกับข้อมูลให้เหมาะสมก่อนที่จะนำข้อมูลไปวิเคราะห์ เพื่อใช้ในการสร้างตัวแบบการพยากรณ์ เช่น การสำรวจข้อมูล การหาความผิดปกติของข้อมูล (Outlier)

2) การวิเคราะห์ข้อมูล

การสร้างตัวแบบพยากรณ์รายวันนั้นจะต้องพิจารณาว่าข้อมูลที่จะพยากรณ์มีลักษณะแบบใด ดังนั้นควรตรวจสอบลักษณะข้อมูลก่อนเป็นอันดับแรกเพื่อที่จะได้เลือกวิธีการพยากรณ์ได้ถูกต้องเหมาะสม เช่น ทำการตรวจสอบข้อมูลนั้นว่ามีส่วนประกอบของแนวโน้มและความแปรผันของฤดูกาลหรือไม่ หรือข้อมูลมีลักษณะเป็นฟังก์ชันเชิงเส้นตรง (Linear Model) หรือไม่ ซึ่งสามารถตรวจสอบได้จากการวาดกราฟ หรือการทดสอบทางสถิติ และนำไปพยากรณ์ด้วยตัวแบบที่เหมาะสม

3) การสร้างตัวแบบการพยากรณ์

หลังจากข้อมูลได้ถูกเตรียมพร้อมสำหรับการสร้างตัวแบบการพยากรณ์จากข้อ 2. แล้ว ในงานวิจัยนี้จะทำการสร้างตัวแบบการพยากรณ์ด้วย 3 วิธีการ ตามข้อกำหนดขอบเขตงาน ดังนี้

1. การพยากรณ์ด้วยวิธีของบ็อกซ์-เจนกินส์

เป็นวิธีที่สร้างสมการตัวแบบจากความสัมพันธ์ของข้อมูลในเชิงเส้น ซึ่งข้อมูลที่น่ามาใช้ในการสร้างตัวแบบต้องมีลักษณะที่คงที่ซึ่งสามารถทดสอบด้วย Augmented Dickey-Fuller Test หากข้อมูลยังไม่คงที่ จะต้องทำการปรับข้อมูลให้ข้อมูลมีลักษณะคงที่ด้วยวิธีการหาผลต่าง ขั้นตอนต่อไปคือทำการหาพารามิเตอร์ในตัวแบบจากรูป

ฟังก์ชันสหสัมพันธ์ในตัวเอง (Autocorrelation Function: ACF) และรูปฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน (Partial Autocorrelation Function: PACF) และทำการทดสอบว่าตัวแบบที่ได้มีความเหมาะสมหรือไม่ด้วยการทดสอบ Ljung-Box (Ljung & Box, 1978) นอกจากนี้สามารถเพิ่มความแม่นยำของตัวแบบด้วยการเพิ่มปัจจัยภายนอกเข้าไปในตัวแบบ

2. การพยากรณ์ตัวแบบด้วยวิธีโครงข่ายประสาทเทียม

หากข้อมูลมีลักษณะไม่เป็นฟังก์ชันเชิงเส้นตรง วิธีที่เหมาะสมสำหรับข้อมูลประเภทนี้คือ โครงข่ายประสาทเทียมวิธีเพอร์เซ็ปตรอนหลายชั้น (Multi-Layer Perceptron Neural Networks) ซึ่งวิธีนี้สามารถแก้ปัญหาพยากรณ์ที่มีความซับซ้อนได้ดี มีความแม่นยำอยู่ในระดับสูง โดยปกติจะใช้แก้ปัญหาที่ไม่เป็นเชิงเส้นตรง วิธีการพยากรณ์คือ ใช้ข้อมูลเก่าในอดีตของข้อมูลชุดนั้นเพียงอย่างเดียว มาพยากรณ์ค่าในอนาคต โดยไม่ใช้ตัวแปรอื่นๆ มาวิเคราะห์

3. การพยากรณ์ตัวแบบด้วยวิธีตัวแบบผสม

เป็นการนำตัวแบบที่ใช้พยากรณ์ข้อมูลอนุกรมเวลาส่วนที่เป็นฟังก์ชันเชิงเส้นตรง คือ ตัวแบบบ็อกซ์-เจนกินส์ มารวมกับตัวแบบที่ใช้พยากรณ์ข้อมูลส่วนที่ไม่เป็นฟังก์ชันเชิงเส้นตรง คือ ตัวแบบโครงข่ายประสาทเทียม เพื่อให้ค่าพยากรณ์ที่ได้มีความถูกต้องและมีความแม่นยำมากยิ่งขึ้น

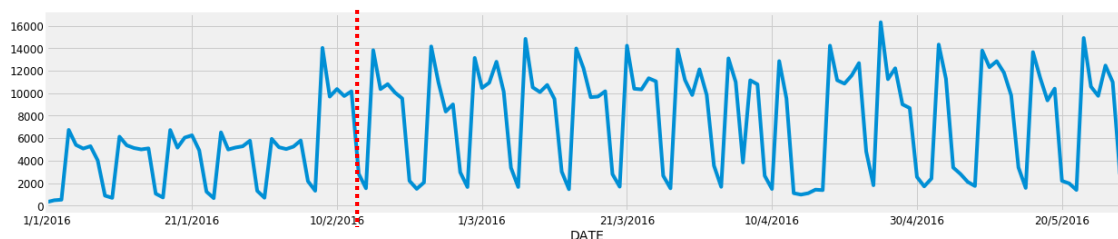
3. การวิเคราะห์ข้อมูล

เนื่องจากการสร้างตัวแบบพยากรณ์ จำเป็นต้องมีความเข้าใจในข้อมูลที่จะใช้ จึงต้องทำการวิเคราะห์ข้อมูลก่อนนำไปสร้างตัวแบบดังนี้

1) วิเคราะห์ข้อมูลอนุกรมเวลาว่ามีส่วนประกอบของแนวโน้มหรือความแปรผันของฤดูกาลเข้ามาเกี่ยวข้องหรือไม่

2) ถ้าข้อมูลอนุกรมเวลาเคลื่อนไหวรอบค่าเฉลี่ยหรือความแปรปรวนอยู่ในสถานะไม่คงที่จำเป็นต้องมีการแปลงข้อมูลนั้นให้มีสถานะคงที่ก่อนนำไปใช้งาน

ในการวิเคราะห์ข้อมูลได้อาศัยการเขียน โปรแกรมภาษา Python บน Jupyter Notebook

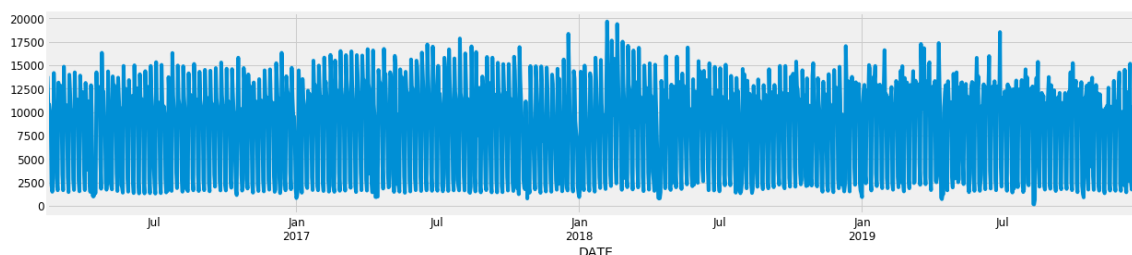


รูปที่ 1 ข้อมูลปริมาณสายโทรเข้ารายวันในช่วง 1/1/2016 - 29/5/2016

จากรูปที่ 1 จะสังเกตได้ว่าในช่วงต้นของกราฟตั้งแต่วันที่ 1/1/2016 ถึงวันที่ 14/2/2016 มีลักษณะของข้อมูลแตกต่างจากช่วงอื่น ดังนั้นจึงได้ทำการพล็อตกราฟในช่วงต้น เพื่อหาความแตกต่างระหว่างข้อมูลช่วงต้นกับข้อมูลที่เหลือ พบว่าข้อมูลมีความผิดปกติจริง เนื่องจากทางบริษัทประกันเปลี่ยนระบบในการจัดเก็บข้อมูล จึงทำให้ข้อมูลในช่วงต้นนี้ได้รับผลกระทบจากการเปลี่ยนระบบในการเก็บข้อมูล

ดังนั้น เพื่อลดความผิดพลาดและเพิ่มประสิทธิภาพของตัวแบบการพยากรณ์ ข้อมูลช่วงต้นนี้จะถูกตัดออก จากข้อมูลที่ใช้ในการวิเคราะห์ ทำให้ข้อมูลที่จะใช้ในการวิเคราะห์คือข้อมูลตั้งแต่วันที่ 15/2/2016 - 31/12/2019 (จำนวน 1,416 วัน หรือ 1,416 จุดเวลา) โดยนำมาทำการแบ่งข้อมูลออกเป็น 2 ชุด ได้แก่

- 1) ข้อมูลฝึกสอนตั้งแต่วันที่ 15/2/2016 - 31/12/2018 (จำนวน 1,051 จุดเวลา) สำหรับสร้างตัวแบบ
- 2) ชุดข้อมูลทดสอบ 1/1/2019 - 31/12/2019 (จำนวน 365 จุดเวลา) สำหรับประเมินความถูกต้องของตัวแบบ



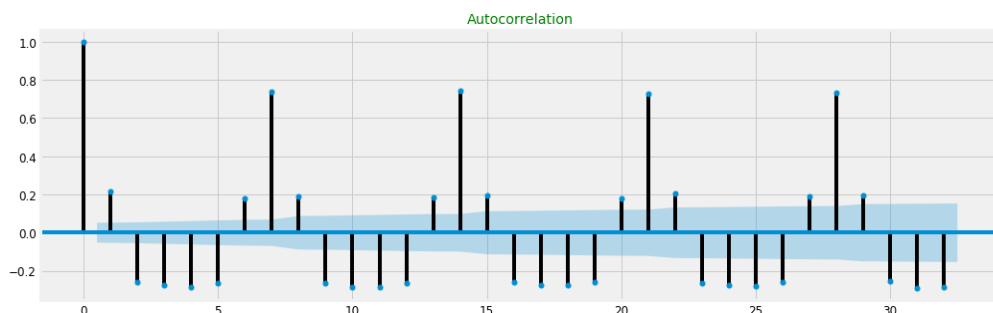
รูปที่ 2 ข้อมูลปริมาณสายโทรเข้ารายวันในช่วง 15/2/2016 - 31/12/2018

จากการพล็อตกราฟระหว่างวันที่และปริมาณสายโทรเข้าในระหว่างวันที่ 15/2/2016 ถึง วันที่ 31/12/2018 เพื่อวิเคราะห์ลักษณะข้อมูล ดังรูปที่ 2 จะพบว่าข้อมูลค่อนข้างมีลักษณะคงที่ กล่าวคือข้อมูลนี้มีค่าเฉลี่ยและความแปรปรวนคงที่ และเมื่อทำการพล็อตกราฟระหว่างวันที่และปริมาณสายโทรเข้าในระหว่างวันที่ 4/6/2016 ถึงวันที่ 25/5/2016 เพื่อวิเคราะห์ความแปรผันของฤดูกาล ซึ่งเป็นช่วงเวลาที่เกิดขึ้นการเปลี่ยนแปลงซ้ำ ๆ กันในช่วงเวลาหนึ่ง ดังรูปที่ 3 จะเห็นได้ว่าปริมาณสายโทรเข้ามีความแปรผันของฤดูกาลอย่างชัดเจน หากอนุกรมเวลาใดๆ ที่พิจารณา

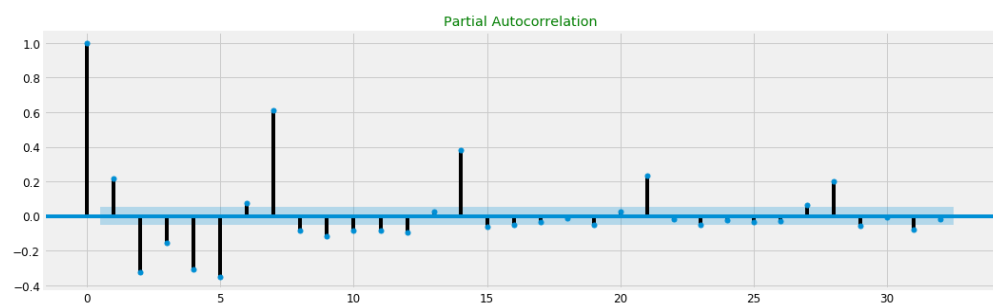
ความแปรผันทางฤดูกาลอยู่ด้วยจะต้องทำการกำจัดความแปรผันทางฤดูกาลออกไปก่อน ซึ่งวิธีกำจัดความผันแปรทางฤดูกาลนั้นมีหลายวิธี ในวิจัยนี้จะใช้วิธีหาผลต่างของฤดูกาล



รูปที่ 3 ข้อมูลปริมาณสายโทรเข้าในช่วง 4/6/2016-25/6/2016



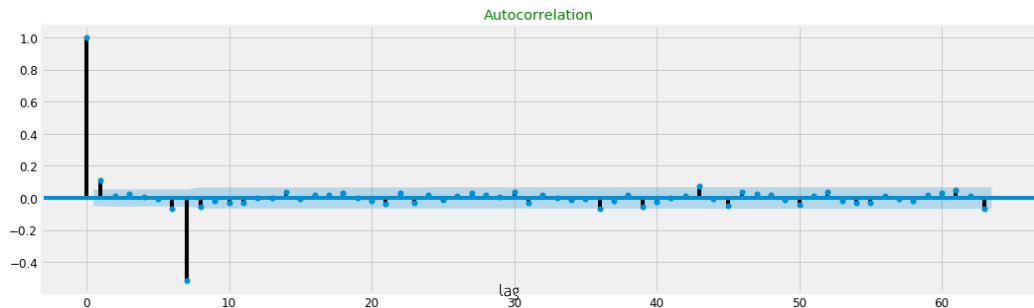
รูปที่ 4 ค่าฟังก์ชันสหสัมพันธ์ในตัวเอง (ACF) ของปริมาณสายโทรเข้ารายวัน



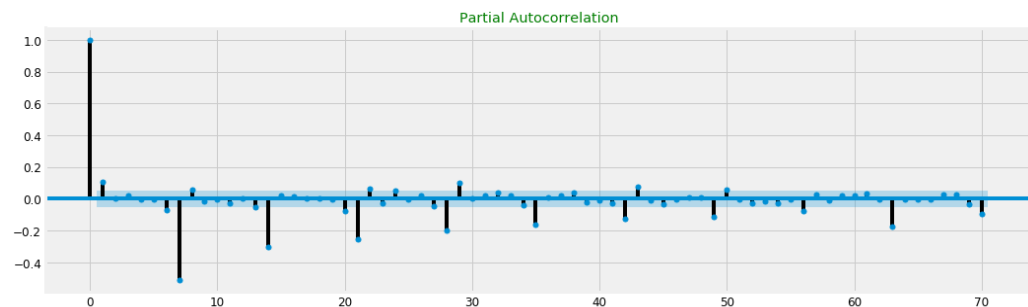
รูปที่ 5 ค่าฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน (PACF) ของปริมาณสายโทรเข้ารายวัน

เมื่อพิจารณารูปที่ 4 แสดงค่าฟังก์ชันสหสัมพันธ์ในตัวเอง (ACF) ของปริมาณสายโทรเข้ารายวัน พบว่าฟังก์ชันสหสัมพันธ์ในตัวเอง ลดลงเข้าสู่ศูนย์อย่างช้า ๆ และรูปที่ 5 พบว่ากราฟฟังก์ชันสหสัมพันธ์ในตัวเอง

บางส่วน (PACF) ลดลงเข้าสู่ศูนย์อย่างรวดเร็ว แสดงว่า อนุกรมเวลาของปริมาณสายโทรเข้ารายวัน ไม่คงที่ จึงทำการหาผลต่างครั้งที่ 1 ของความแปรผันเกี่ยวกับฤดูกาลของปริมาณสายโทรเข้ารายวัน



รูปที่ 6 ค่าฟังก์ชันสหสัมพันธ์ในตัวเอง (ACF) ของผลต่างฤดูกาลครั้งที่ 1



รูปที่ 7 ค่าฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน (PACF) ของผลต่างฤดูกาลครั้งที่ 1

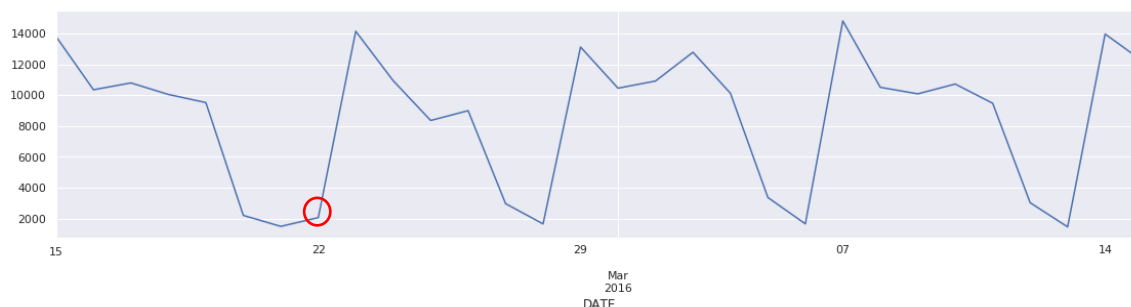
จากรูปที่ 6 และ 7 เมื่อทำการหาผลต่างครั้งที่ 1 ของความแปรผันเกี่ยวกับฤดูกาล พบว่าค่า ACF และ PACF ลดลงเข้าสู่ศูนย์อย่างรวดเร็ว แสดงว่าอนุกรมเวลาคงที่

เมื่อทำการพิจารณาวันหยุดต่างๆ ในปี 2016 – 2018 จากรูปที่ 8 พบว่าในวันที่ 5 ธันวาคม วันที่ 10 ธันวาคม 31 ธันวาคม และ 1 มกราคมของทุกปี มีปริมาณสายที่โทรเข้ามาน้อยกว่าปกติ เนื่องจากเป็นวันหยุดทางราชการ



รูปที่ 8 ปริมาณสายโทรเข้ามาในแต่ละวันตั้งแต่วันที่ 1/12/2016-10/1/2017

และจากรูปที่ 9 พบว่า ในวันที่ 22 กุมภาพันธ์ 2016 มีปริมาณสายที่โทรเข้ามาน้อยกว่าปกติ เนื่องจากเป็นวันชดเชยมาฆบูชา



รูปที่ 9 ปริมาณสายโทรเข้ามาในแต่ละวันตั้งแต่วันที่ 15/2/2016-15/3/2016

เนื่องจากสมบัติความคงที่ของข้อมูลเป็นสิ่งจำเป็นที่ข้อมูลพึงมีก่อนที่จะทำการพยากรณ์ด้วยแบบวิธีบ็อกซ์-เจนกินส์ ดังนั้นจึงทำการทดสอบด้วยตัวทดสอบทางสถิติ Augmented Dickey-Fuller Test (ADF) ที่ระดับนัยสำคัญ 0.05 พบว่าค่า $p\text{-value} = 0.0000002 < \alpha = 0.05$ สรุปได้ว่า ข้อมูลชุดนี้มีลักษณะคงที่ ที่ระดับนัยสำคัญ 0.05

4. ตัวแบบพยากรณ์ที่ใช้ในการวิจัย

1) การพยากรณ์วิธีการบ็อกซ์-เจนกินส์

วิธีการของบ็อกซ์-เจนกินส์ เหมาะกับข้อมูลที่มีลักษณะคงที่และไม่มีส่วนประกอบของแนวโน้มหรือความแปรผันเกี่ยวกับฤดูกาล หากข้อมูลยังมีลักษณะ ไม่คงที่หรือมีส่วนประกอบของแนวโน้มหรือความแปรผันเกี่ยวกับฤดูกาลอย่างชัดเจน จะต้องทำการหาผลต่างครั้งที่ 1 (1st differencing) ซึ่งแบบจำลองมีข้อดกลงเบื้องต้นคือค่าสังเกตปัจจุบันมีความสัมพันธ์เป็นกระบวนการเชิงเส้นระหว่างค่าสังเกตและค่าคลาดเคลื่อนในอดีต (Chatfield, 2016) เนื่องจากพบว่าข้อมูลมีความแปรผันทางฤดูกาลอย่างชัดเจน ทางผู้วิจัยจึงใช้การพยากรณ์วิธีบ็อกซ์-เจนกินส์ด้วยตัวแบบ SARIMA ซึ่งเป็นตัวแบบที่เหมาะสมกับข้อมูลที่มีลักษณะนี้

2) ตัวแบบโครงข่ายประสาทเทียมวิธีเพอร์เซ็ปตรอนหลายชั้น

ตัวแบบ ANN สามารถประยุกต์ใช้ในการพยากรณ์อนุกรมเวลาได้ ส่วนใหญ่นิยมใช้พยากรณ์ข้อมูลอนุกรมเวลาชุดเดียวซึ่งเรียกว่า “Univariate Time Series” โดยกำหนดชั้นอินพุตแสดงข้อมูลในอดีต จำนวน n ข้อมูล และชั้นเอาต์พุตแสดงการพยากรณ์วันข้างหน้า 1 วัน และมีชั้นซ่อนเร้น ที่สามารถกำหนดความซับซ้อนของตัวแบบเพื่อให้ได้ค่าพยากรณ์ที่แม่นยำขึ้น

การพยากรณ์ข้อมูลอนุกรมเวลามีเป้าหมาย คือการพยากรณ์ข้อมูลล่วงหน้า โดยใช้ข้อมูลเก่าในอดีตของข้อมูลชุดนั้นเพียงอย่างเดียวมาพยากรณ์ค่าในอนาคต โดยไม่ใช่ตัวแปรอื่นมาวิเคราะห์ โดยส่วนมากการพยากรณ์อนุกรมเวลาจะเกี่ยวข้องกับตัวเลข (Anansapsuk, 2017)

Zhang (2012) กล่าวว่าไว้ดังนี้ สมมติว่ามีข้อมูลอนุกรมเวลาอยู่ N ค่า $Y_1, Y_2, \dots, Y_N$ ในชุดข้อมูลฝึกสอน (Training set) และต้องการพยากรณ์ 1 ช่วงเวลาไปข้างหน้า (1-step-ahead) คือ Y_{t+1} เมื่อกำหนดจำนวน lag หรือจุดเวลาย้อนหลัง เช่น $\text{lag} = n = 1$ คือทำการ Y_{t+1} ด้วย Y_t , $\text{lag} = n = 2$ คือทำการพยากรณ์ Y_{t+1} ด้วย Y_t และ Y_{t-1} $\text{lag} = n = 3$ คือทำการพยากรณ์ Y_{t+1} ด้วย Y_t , Y_{t-1} และ Y_{t-2}

3) การพยากรณ์วิธีตัวแบบผสม

Zhang (2003) ได้สรุปหลักในการสร้างตัวแบบผสมไว้ดังนี้ กล่าวคือการสร้างตัวแบบผสมจะกำหนดว่าข้อมูลที่นำมาวิเคราะห์นั้นมี 2 องค์ประกอบ คือ องค์ประกอบที่เป็นเส้นตรง (Linear Component) และองค์ประกอบที่ไม่เป็นเส้นตรง (Nonlinear Component) เพื่อให้ค่าพยากรณ์ที่ได้มีความถูกต้องและแม่นยำมากขึ้น วิธีการสร้างตัวแบบผสมมีขั้นตอนดังนี้

1) พยากรณ์ข้อมูลในส่วนฟังก์ชันเชิงเส้นตรง โดยวิเคราะห์ข้อมูลอนุกรมเวลาในส่วนฟังก์ชันเชิงเส้นตรงจากตัวแบบ SARIMAX

2) คำนวณค่าส่วนเหลือ หรือ ความคลาดเคลื่อนจากการพยากรณ์ (Residuals) ที่ได้จากการพยากรณ์ด้วยตัวแบบ SARIMAX ไปพยากรณ์ในส่วนที่ไม่เป็นฟังก์ชันเชิงเส้นตรง

3) พยากรณ์ข้อมูลในส่วนที่ไม่เป็นฟังก์ชันเชิงเส้นตรง โดยนำส่วนที่ไม่เป็นฟังก์ชันเชิงเส้นตรงไปพยากรณ์ด้วยตัวแบบเครือข่ายประสาทเทียม ด้วยการกำหนดค่าไฮเปอร์พารามิเตอร์หรือกระบวนการในการสร้างตัวแบบ

4) คำนวณค่าพยากรณ์รวม (Total Forecasting) โดยนำตัวแบบมารวมกันที่ได้จากการพยากรณ์ข้อมูลในส่วนฟังก์ชันเชิงเส้นตรง และการพยากรณ์ข้อมูลในส่วนที่ไม่เป็นฟังก์ชันเชิงเส้นตรง

หลังจากนั้นทำการเปรียบเทียบประสิทธิภาพของตัวแบบพยากรณ์ โดยพิจารณาจากค่าเฉลี่ยเปอร์เซ็นต์ความคลาดเคลื่อนสมบูรณ์ (Mean Absolute Percentage Error : MAPE) เป็นเกณฑ์ในการเปรียบเทียบตัวแบบ ซึ่งตัวแบบที่ให้ค่า MAPE ที่ต่ำที่สุด จะเป็นตัวแบบที่ดีที่สุด และนำตัวแบบที่ดีที่สุดไปใช้ในการพยากรณ์ต่อไป

ผลการวิจัย

จากการตรวจสอบข้อมูล พบว่าข้อมูลมีความแปรผันทางฤดูกาลทุกๆ 7 วัน จึงกำหนดค่า $s=7$ ในการพิจารณาตัวแบบที่เหมาะสม จะทำการพิจารณาตัวแบบที่มีค่า AIC น้อยที่สุด

ตารางที่ 1 ค่า AIC ในแต่ละอันดับของตัวแบบ ARIMA(p,d,q)SARIMA(P,D,Q)_s

ARIMA(p,d,q)SARIMA(P,D,Q) _s	AIC
ARIMA(0, 0, 0)xSARIMA(0, 0, 0) ₇	30,037.485
ARIMA(0, 0, 0)xSARIMA(0, 0, 1) ₇	29,115.961
ARIMA(0, 0, 0)xSARIMA(0, 1, 0) ₇	26,953.149
ARIMA(0, 0, 0)xSARIMA(0, 1, 1) ₇	25,988.764
ARIMA(1, 1, 1)xSARIMA(0, 1, 0) ₇	26,935.634
ARIMA(1, 1, 1)xSARIMA(0, 1, 1)₇	25,900.738
...	...
ARIMA(1, 1, 2)xSARIMA(0, 0, 2) ₇	27,000.077
ARIMA(1, 1, 2)xSARIMA(0, 2, 0) ₇	28,211.030
ARIMA(1, 1, 2)xSARIMA(0, 2, 2) ₇	26,983.504
ARIMA(1, 1, 2)xSARIMA(0, 2, 2) ₇	26,983.504
ARIMA(1, 2, 0)xSARIMA(0, 0, 2) ₇	28,622.106
ARIMA(2, 0, 1)xSARIMA(3, 1, 1) ₇	27,134.785
ARIMA(2, 0, 2)xSARIMA(3, 1, 2) ₇	26,105.861
ARIMA(2, 1, 1)xSARIMA(3, 1, 1) ₇	26,953.248
ARIMA(2, 1, 2)xSARIMA(3, 1, 1) ₇	26,044.364
ARIMA(2, 1, 2)xSARIMA(3, 1, 2) ₇	25,935.634

เมื่อพิจารณาตารางที่ 1 ค่า AIC ที่น้อยที่สุด พบว่าอันดับที่ดีที่สุด ได้แก่ p=2 , d=0 , q=2 และ P=3 , D=1 ,Q=1 และตัวแบบที่เหมาะสมคือ ตัวแบบ ARIMA(2,0,2)xSARIMA(3,1,1)₇ ซึ่งมีรูปแบบทั่วไปดังสมการที่ 1

$$(1 - f_1 B)(1 - B)(1 - B^7) Y_t = (1 - q_1 B)(1 - Q_1 B^7) e_t \quad (1)$$

โดยกำหนด

Y_t แทน ค่าพยากรณ์ ณ เวลา t

e_t	แทน ค่าความคลาดเคลื่อน ณ เวลาที่ t
B	แทน ตัวดำเนินการย้อนหลัง
f_1	แทน การถดถอยในตัวเองแบบ ไม่มีฤดูกาลอันดับ 1
q_1	แทน ค่าเฉลี่ยเคลื่อนที่แบบ ไม่มีฤดูกาลอันดับ 1
Q_1	แทน ค่าเฉลี่ยเคลื่อนที่แบบ มีฤดูกาลอันดับ 1

หลังจากได้ตัวแบบที่เหมาะสมจากสมการที่ 4.1 แล้วจึงนำตัวแบบดังกล่าวมาประมาณค่าพารามิเตอร์ดังตารางที่ 2

ตารางที่ 2 การทดสอบค่าพารามิเตอร์ ARIMA(1,1,1)xSARIMA(0,1,1)₇

Statistics	Coef	SE Coef	t	p-value
f_1	0.1185	0.017	6.929	0.000
q_1	-0.9999	0.022	-46.418	0.000
Q_1	-1.0782	0.009	-115.309	0.000

จากตารางที่ 4.2 พบว่าค่า $p\text{-value} = 0.000 < \alpha = 0.05$ จึงปฏิเสธ H_0 ที่ระดับนัยสำคัญ 0.05 สรุปคือค่าพารามิเตอร์ f_1 , q_1 , Q_1 ในตัวแบบมีค่าไม่เท่ากับ 0 ดังนั้นตัวแบบ ARIMA(1,1,1) x SARIMA(0,1,1)₇ เป็นตัวแบบที่เหมาะสม

จากนั้นทำการตรวจสอบค่าความคลาดเคลื่อนว่าจะต้องไม่มีสหสัมพันธ์ในตัวเองด้วย Box-Pierce and Ljung-Box Tests พบว่า $p\text{-value} = 0.8 > \alpha = 0.05$ จึงยอมรับ H_0 ที่ระดับนัยสำคัญ 0.05 สรุปว่า ค่าความคลาดเคลื่อน ของตัวแบบ ที่พิจารณาไม่มีสหสัมพันธ์ในตัวเองทุกช่วงเวลา ดังนั้นแสดงว่าตัวแบบ ARIMA(1,1,1)xSARIMA(0,1,1)₇ เป็นตัวแบบที่เหมาะสม

ดังนั้นสมพยากรณ์ด้วยตัวแบบ SARIMA คือ

$$Y_t = (1 + f_1)Y_{t-1} - f_1Y_{t-2} + Y_{t-7} - (1 + f_1)Y_{t-8} + f_1Y_{t-9} + e_t - Q_1e_{t-7} - q_1e_{t-1} + q_1Q_1e_{t-8}$$

สำหรับการพยากรณ์วิธีการบ็อกซ์-เจนกินส์ด้วยตัวแบบ SARIMAX สำหรับปริมาณสายโทรเข้ามาในแต่ ละวันตั้งแต่วันที่ 1/12/2016-10/1/2017 พบว่าปริมาณสายที่โทรเข้ามาน้อยในวันหยุดทางราชการ และหาก

วันหยุดราชการตรงกับ เสาร์ - อาทิตย์ จะมีการชดเชยในวันจันทร์ถัดไป โดยประกอบไปด้วย วันสงกรานต์, วันปีใหม่, วันแม่, วันปิยมหาราช, วันพ่อ, วันจักรี, วันรัฐธรรมนูญ และวันแรงงาน รวมถึงวันหยุดทางศาสนา และ หากวันหยุดทางศาสนาตรงกับ เสาร์ - อาทิตย์ จะมีการชดเชยในวันจันทร์ถัดไป โดยประกอบไปด้วยวันมาฆบูชา, วันวิสาขบูชา และวันอาสาฬหบูชา ทั้งนี้ได้เพิ่มตัวแปรอิสระหรือตัวแปรภายนอก (Exogenous) ขึ้นมาโดยที่

X1 แทน วันหยุดทางราชการ

X2 แทน วันหยุดชดเชยทางราชการ

X3 แทน วันหยุดพิเศษ

X4 แทน วันหยุดชดเชยพิเศษ

X5 แทน วันหยุดสงกรานต์-ปีใหม่

X6 แทน วันหยุดสงกรานต์-ปีใหม่ตรงกับวันเสาร์-อาทิตย์

X7 แทน วันหยุดชดเชยสงกรานต์-ปีใหม่

โดยตัวแบบที่เหมาะสม คือ ARIMA (2,0,2)xSARIMA(3,1,1) with Exogenous ดังนั้นสมการพยากรณ์ด้วยตัวแบบ SARIMAX คือ

$$Y_t = (1 + f_1)Y_{t-1} - f_1Y_{t-2} + Y_{t-7} - (1 + f_1)Y_{t-8} + f_1Y_{t-9} + e_t - Q_1e_{t-7} - q_1e_{t-1} + q_1Q_1e_{t-8} \\ + 9,663.5124X_1 - 12,690X_2 - 8,982.4182X_3 - 12,420X_4 - 5,856.0354X_5 \\ + 1,480.2860X_6 - 12,410X_7$$

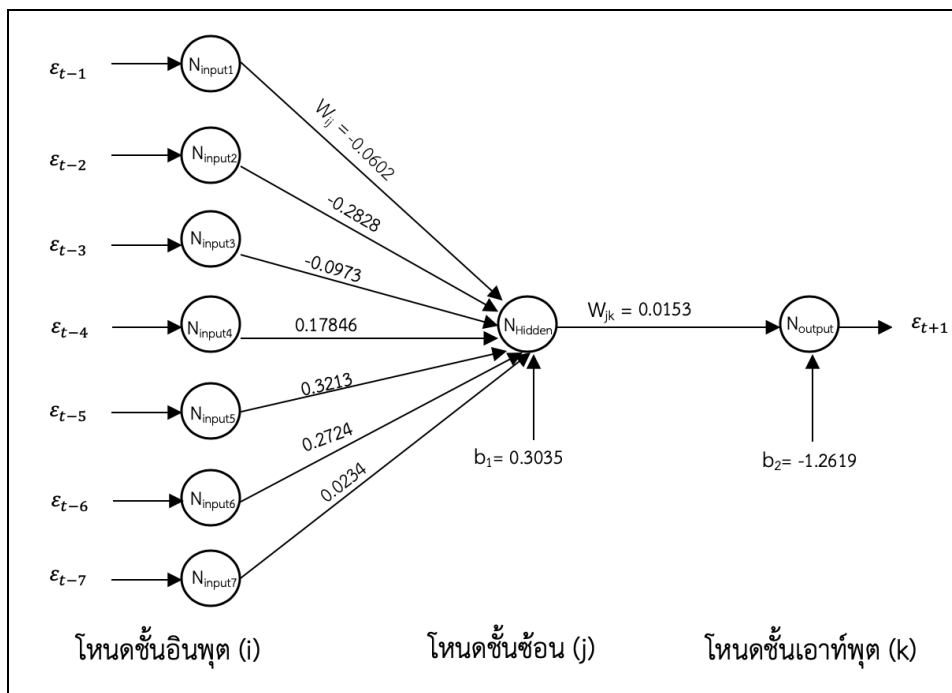
หลังจากนั้นคำนวณค่าส่วนเหลือ หรือ ความคลาดเคลื่อนจากการพยากรณ์ (Residuals) โดยส่วนเหลือนั้นจะเหลือเพียงฟังก์ชันส่วนไม่เป็นเส้นตรงด้วยตัวแบบ SARIMAX นำเข้าสู่ตัวแบบโครงข่ายประสาทเทียมเพอร์เซ็ปตรอนหลายชั้นด้วยการเรียนรู้แบบไปข้างหน้า (Feed forward) โดยเลือกกำหนดค่าไฮเปอร์พารามิเตอร์ดังตารางที่ 3

ตารางที่ 3 การกำหนดค่าไฮเปอร์พารามิเตอร์หรือกระบวนการในการสร้างตัวแบบ

ข้อกำหนดต่าง ๆ ในการสร้างตัวแบบ	ค่าไฮเปอร์พารามิเตอร์
อัลกอริทึม (Algorithm)	Neural Network
จำนวนโหนดในชั้นอินพุต (Input Node)	7
จำนวนชั้นซ่อน (Hidden Layer)	1

ข้อกำหนดต่าง ๆ ในการสร้างตัวแบบ	ค่าไฮเปอร์พารามิเตอร์
จำนวนโหนดในชั้นซ่อน (Hidden Node)	1
จำนวนโหนดในชั้นเอาต์พุต (Output Node)	1
ฟังก์ชันกระตุ้น (Activation Function)	Linear Function
Optimizer	Adam
โมเมนตัม (Momentum)	0.9
อัตราการเรียนรู้ (Learning Rate)	0.001

จากตารางที่ 3 แสดงการกำหนดค่าไฮเปอร์พารามิเตอร์สำหรับการสร้างตัวแบบโครงข่ายประสาทเทียม โดยสมการพยากรณ์วิธีผสมเป็นการรวมข้อมูลส่วนที่เป็นฟังก์ชันเชิงเส้นตรงที่ได้จากการพยากรณ์ด้วยตัวแบบ SARIMAX และข้อมูลส่วนที่ไม่เป็นฟังก์ชันเชิงเส้นตรงที่ได้จากการพยากรณ์ค่าส่วนเหลือด้วยตัวแบบโครงข่ายประสาทเทียมเพื่อใช้ป้อนหลายชั้นที่พยากรณ์ 1 วันข้างหน้าด้วยค่าส่วนเหลือ 7 วันก่อนหน้า หรืออยู่ในรูปแบบ 7-1-1 เข้าด้วยกัน มีรูปแบบโครงสร้างดังรูปที่ 10



รูปที่ 10 รูปแบบโครงสร้างในการพยากรณ์ส่วนไม่เป็นฟังก์ชันเชิงเส้นตรง

โดยกำหนดให้

e_{t-1} แทน ค่าส่วนเหลือก่อนหน้า 1 วัน

e_{t-2} แทน ค่าส่วนเหลือก่อนหน้า 2 วัน

e_{t-3} แทน ค่าส่วนเหลือก่อนหน้า 3 วัน

e_{t-4} แทน ค่าส่วนเหลือก่อนหน้า 4 วัน

e_{t-5} แทน ค่าส่วนเหลือก่อนหน้า 5 วัน

e_{t-6} แทน ค่าส่วนเหลือก่อนหน้า 6 วัน

e_{t-7} แทน ค่าส่วนเหลือก่อนหน้า 7 วัน

e_{t+1} แทน ค่าส่วนเหลือ 1 วันข้างหน้า

b_1 แทน ค่าไบเอส (Bias) ในชั้นซ่อน

b_2 แทน ค่าไบเอส (Bias) ในชั้นเอาต์พุต

w_{ij} แทน ค่าถ่วงน้ำหนักของเส้นเชื่อมโยงกับข้อมูลนำเข้าที่ i ไปยังโหนดที่ j หรือโหนดชั้นซ่อน

w_{jk} แทน ค่าถ่วงน้ำหนักของเส้นเชื่อมโยงกับโหนดชั้นซ่อนที่ j ไปยังโหนดที่ k หรือโหนดชั้น

เอาต์พุต

จากการทดลองได้ตัวแบบที่เหมาะสมในการพยากรณ์อนุกรมเวลาในส่วนไม่เป็นฟังก์ชันเชิงเส้นตรง ดังตารางที่ 4

ตารางที่ 4 ค่าพารามิเตอร์ของตัวแบบที่เหมาะสมที่สุดส่วนไม่เป็นฟังก์ชันเชิงเส้นตรง

ตัวแบบ	พารามิเตอร์		
	Lag	Hidden neurons	Output neurons
ANN (MLP)	7	1	1

การเปรียบเทียบประสิทธิภาพของตัวแบบพยากรณ์

หลังจากทดลองหาค่าพารามิเตอร์ที่ดีที่สุดของตัวแบบทั้ง 4 วิธี ได้แก่ ตัวแบบ SARIMA, ตัวแบบ SARIMAX, ตัวแบบ โครงข่ายประสาทเทียมเพอร์เซ็ปตรอนหลายชั้น และตัวแบบผสม Hybrid จากนั้นทำการเปรียบเทียบความแม่นยำของค่าพยากรณ์โดยใช้ค่าเฉลี่ยเปอร์เซ็นต์ความคลาดเคลื่อนสมบูรณ์ดังตารางที่ 5

ตารางที่ 5 ค่า MAPE ของแต่ละตัวแบบ

ตัวแบบ	MAPE
SARIMA	44.97%
SARIMAX	24.06%
ANN (MLP)	43.70%
Hybrid	24.02%

เมื่อพิจารณาค่า MAPE ที่ดีที่สุดของแต่ละตัวแบบ เพื่อเลือกแบบจำลองที่มีความแม่นยำที่สุดนั้น จะเห็นว่าแบบจำลองที่เหมาะสมที่สุดในการพยากรณ์ปริมาณสายโทรเข้า ได้แบบจำลองที่ดีที่สุดคือ ตัวแบบจำลองผสมระหว่างตัวแบบ SARIMA กับตัวแบบโครงข่ายประสาทเทียม ที่มีค่า MAPE เท่ากับ 24.02%

สรุปผลการวิจัย

ผลการเปรียบเทียบความแม่นยำของตัวแบบพยากรณ์ทั้ง 4 ตัวแบบ สามารถสรุปผลการดำเนินวิจัยได้ว่าตัวแบบผสม โดยนำค่าส่วนเหลือที่ได้จากการพยากรณ์ด้วยตัวแบบ ARIMA(1,1,1) x SARIMA(0,1,1) และ ตัวแปรปัจจัยภายนอกซึ่งประกอบไปด้วย X1 คือวันหยุดทางราชการ, X2 คือวันหยุดชดเชยทางราชการ, X3 คือวันหยุดพิเศษ, X4 คือวันหยุดชดเชยพิเศษ, X5 คือวันหยุดสงกรานต์-ปีใหม่, X6 คือวันหยุดสงกรานต์-ปีใหม่ตรงกับวันเสาร์อาทิตย์, X7 คือวันหยุดชดเชยสงกรานต์-ปีใหม่ มาสร้างตัวแบบโครงข่ายประสาทเทียมด้วยวิธีเพอร์เซ็ปตรอนหลายชั้น พบว่าตัวแบบโครงข่ายประสาทเทียมชั้นข้อมูลนำเข้ามีจำนวนโหนด 7 โหนด ชั้นซ่อนมีจำนวนโหนด 1 โหนด และชั้นผลลัพธ์มีจำนวนโหนด 1 โหนด หรือเรียกว่าตัวแบบโครงข่ายประสาทเทียม 7-1-1 เป็นตัวแบบที่เหมาะสมที่สุด และเมื่อนำไปทดสอบความแม่นยำของตัวแบบให้ค่า MAPE เท่ากับ 24.02% โดยตัวแบบดังกล่าวสามารถนำไปพยากรณ์ล่วงหน้า เพื่อเป็นแนวทางหรือข้อมูลเพิ่มเติมในการตัดสินใจในการวางแผนการทำงานของศูนย์บริการข้อมูลลูกค้า

เอกสารอ้างอิง

- Anansapsuk, A. (2017). *A Comparative Study of Hybrid Time Series Models for Forecasting Seasonal Time Series* [Unpublished master's thesis]. Chulalongkorn University. (in Thai)
- Ljung, G., & Box, G. (1978). On a Measure of Lack of Fit in Time Series Models. *Biometrika*, 65(2), 297-303.
- Chatfield, C. (2016). *The Analysis of Time Series: An Introduction* (6th ed.). New York, Chapman and Hall/CRC Press.
- Ibrahim, R., Ye, H., L'Ecuyer, P., & Shen, H. (2016). Modeling and forecasting call center arrivals: A literature survey and a case study. *International Journal of Forecasting*, 32(3), 865-874.
- Njeru, K., Na, K., & Ivivi, J. (2018). Forecasting future customer call volumes: A case study. *International Journal on Future Revolution in Computer Science & Communicaton Engineering*, 4(6), 12-16.
- Taylor, J. W. (2008). A comparison of univariate time series methods for forecasting intraday arrivals at a call center. *Management Science*, 54(2), 253–265.
- Zhang, G. P. (2003). Time series forecasting using a hybrid ARIMA and neural network model. *Neurocomputing*, 50, 159-175.
- Zhang, G. P. (2012). Neural Networks for Time-Series Forecasting. In: Rozenberg G., Bäck T., Kok J.N. (eds.) *Handbook of Natural Computing*. Berlin, Heidelberg, Springer. https://doi.org/10.1007/978-3-540-92910-9_14

Research Article

ตัวแบบพยากรณ์มูลค่าการส่งออกน้ำตาลของประเทศไทย

Forecasting Model for the Export Values for Sugar of Thailand

วารangkนา เรียนสุทธิ^{1*}

Warangkhan Riansut^{1*}

¹สาขาวิชาคณิตศาสตร์และสถิติ คณะวิทยาศาสตร์ มหาวิทยาลัยทักษิณ อำเภอป่าพะยอม จังหวัดพัทลุง 93210 ประเทศไทย

¹Department of Mathematics and Statistics, Faculty of Science, Thaksin University, Pa Phayom, Phatthalung 93210 Thailand

*E-mail: warang27@gmail.com

Received: 09/05/2021; Revised: 26/07/2021; Accepted: 12/05/2022

บทคัดย่อ

การศึกษานี้มีวัตถุประสงค์เพื่อสร้างตัวแบบพยากรณ์มูลค่าการส่งออกน้ำตาลของประเทศไทยด้วยวิธีการทางสถิติ โดยใช้ข้อมูลค่าเฉลี่ยรายเดือนจากเว็บไซต์ของสำนักงานเศรษฐกิจการเกษตร ตั้งแต่เดือนมกราคม 2554 ถึงเดือนพฤศจิกายน 2563 จำนวน 119 เดือน ผู้วิจัยได้แบ่งข้อมูลออกเป็น 2 ชุด ชุดที่ 1 ตั้งแต่เดือนมกราคม 2554 ถึงเดือนมิถุนายน 2563 จำนวน 114 เดือน สำหรับการสร้างตัวแบบพยากรณ์ด้วยวิธีการทางสถิติทั้งหมด 7 วิธี ได้แก่ วิธีบ็อกซ์-เจนกินส์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของโฮลต์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีแนวโน้มแบบแฉก วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ ข้อมูลชุดที่ 2 ตั้งแต่เดือนกรกฎาคมถึงเดือนพฤศจิกายน 2563 จำนวน 5 เดือน นำมาใช้สำหรับการเปรียบเทียบความแม่นยำของตัวแบบพยากรณ์ด้วยเกณฑ์รากของค่าคลาดเคลื่อนกำลังสองเฉลี่ยที่ต่ำที่สุด ผลการศึกษาพบว่า วิธีที่มีความแม่นยำมากที่สุด คือ วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ รองลงมา คือ วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์

คำสำคัญ: น้ำตาล, บ็อกซ์-เจนกินส์, การทำให้เรียบด้วยเลขชี้กำลัง, รากของค่าคลาดเคลื่อนกำลังสองเฉลี่ย

Abstract

The objective of this study is to construct the appropriate forecasting model for the export values for sugar of Thailand via the use of statistical methods. The monthly average data, which were gathered from the website of the Office of Agricultural Economics during January 2011 to November 2020 of 119 months were divided into 2 datasets. The first dataset, which consisted of 114 months from January 2011 to June 2020 was used for constructing the forecasting models via the use of 7 statistical methods, namely, Box-Jenkins method, Holt's exponential smoothing method, Brown's exponential smoothing method, damped trend exponential smoothing method, simple seasonal exponential smoothing method, Winters' additive exponential smoothing method, and Winters' multiplicative exponential smoothing method. The second dataset, which consisted of 5 months from July to November 2020 was used for comparing the accuracy of the forecasting model via the lowest root mean square error. The results indicated that the most accurate method was Winters' multiplicative exponential smoothing method, followed by Brown's exponential smoothing method.

Keywords: Sugar, Box-Jenkins, Exponential Smoothing, Root Mean Square Error

บทนำ

น้ำตาล คือ ผลึกซูโครสที่มีอยู่ในพืชต่าง ๆ เช่น อ้อย หัวบีท มะพร้าว และตาล แต่ที่นิยมนำมาทำเป็นอุตสาหกรรม คือ น้ำตาลทรายที่ทำมาจากอ้อย การผลิตน้ำตาลจากอ้อยจะใช้หีบเอาน้ำอ้อยออกจากอ้อย จากนั้นนำน้ำอ้อยไปทำให้สะอาด เกลวจนงวดเพื่อให้ตกผลึกซูโครส แล้วจึงแยกเอาออกจากส่วนที่เป็นของเหลวมาเป็นน้ำตาลที่ใช้บริโภค (Kavitanon, 2004) จากการสำรวจพื้นที่ปลูกอ้อย พบว่า มีปริมาณเพิ่มขึ้นทุกภาคของประเทศไทย (Foundation for Agricultural and Environmental Conservation in Thailand, 2021) เนื่องจากการสนับสนุนให้ปลูกอ้อยกันมาก เพราะสภาพภูมิอากาศในประเทศไทยเหมาะสมแก่การปลูกอ้อย จึงทำให้มีปริมาณอ้อยเข้าสู่โรงงานมากขึ้น ส่งผลให้ผลผลิตน้ำตาลทรายเพิ่มขึ้น ประเทศไทยจึงมีรายได้จากการส่งออกน้ำตาลทรายไปจำหน่ายยังต่างประเทศ ในขณะที่ตลาดโลกยังมีความต้องการน้ำตาลทรายในปริมาณที่เพิ่มขึ้น โดยตลอด (Teppiam & Kittichotipanit, 2015) การรักษาสภาพในการส่งออกน้ำตาลทรายจะเป็นการเพิ่มรายได้ให้กับประเทศได้เป็นอย่างดี ซึ่งการวางแผนการผลิตอย่างมีประสิทธิภาพโดยการพยากรณ์มูลค่าการส่งออกน้ำตาลในอนาคตเป็นประโยชน์อย่างยิ่งต่อหน่วยงานของภาครัฐและผู้ที่เกี่ยวข้องกับอุตสาหกรรมอ้อยและน้ำตาลทราย รวมทั้งกลุ่มบริษัทอ้อยและน้ำตาลทรายของไทย ดังนั้นการพยากรณ์จึงเป็นวิธีหนึ่งที่ช่วยให้ผู้ที่เกี่ยวข้องได้นำข้อมูลมาใช้ตัดสินใจสำหรับสถานการณ์ต่าง ๆ ที่อาจจะเกิดขึ้นในอนาคต โดยถ้าค่าพยากรณ์มีความคลาดเคลื่อนสูงจะส่งผลกระทบต่อหลาย ๆ ฝ่าย เช่น ส่งผลให้ภาครัฐวางแผนนโยบายการเพาะปลูกในปีนั้น ๆ ผิดพลาด อาจก่อให้เกิดภาวะผลผลิตล้นตลาดหรือขาด

ตลาดได้ (Luangtong & Kantanantha, 2016) ด้วยเหตุผลดังกล่าว จึงเป็นที่น่าสนใจที่ควรมีการศึกษาการพยากรณ์ข้อมูลในอนาคต โดยการศึกษาครั้งนี้ผู้วิจัยจะให้ความสนใจกับการพยากรณ์มูลค่าการส่งออกน้ำตาล จึงเริ่มต้นที่การทบทวนวรรณกรรม พบว่า Teppiam & Kittichotipanit (2015) ได้ศึกษาเปรียบเทียบตัวแบบพยากรณ์ปริมาณการส่งออกน้ำตาลทรายดิบของประเทศไทยโดยวิธีอนุกรมเวลา (วิธีบ็อกซ์-เจนกินส์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของโฮลต์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์ วิธีแยกส่วนประกอบอนุกรมเวลา) และวิธีการวิเคราะห์การถดถอยเชิงเส้นพหุ ผลการศึกษาพบว่า วิธีการวิเคราะห์การถดถอยเชิงเส้นพหุมีความเหมาะสมมากกว่า โดยตัวแปรที่มีผลต่อปริมาณการส่งออกน้ำตาลทรายดิบของประเทศไทย คือ ดัชนีการส่งออกน้ำตาลทรายดิบ ปริมาณการส่งออกกากน้ำตาลทราย ปริมาณน้ำฝนเฉลี่ยทั่วประเทศ มีค่าสัมประสิทธิ์การตัดสินใจเชิงพหุเท่ากับ 0.78 ค่าคลาดเคลื่อนกำลังสองเฉลี่ย (Mean Square Error: MSE) เท่ากับ 5.86×10^{15} และ Chantabutr (2004) ได้ศึกษาการพยากรณ์ราคาส่งออกน้ำตาล โดยวิธีบ็อกซ์-เจนกินส์ ผลการศึกษาพบว่า ตัวแบบ MA(1) MA(17) SMA(12) มีความเหมาะสมมากที่สุดที่จะเป็นตัวแทนของราคาส่งออกน้ำตาลดิบ และตัวแบบ AR(30) MA(30) มีความเหมาะสมที่จะเป็นตัวแทนของราคาน้ำตาลทรายขาว ตัวแบบทั้ง 2 ตัวแบบนี้แสดงให้เห็นว่าทิศทางของอนุกรมเวลาระหว่างข้อมูลที่แท้จริงและข้อมูลราคาที่ประมาณขึ้นมีทิศทางการขึ้นลงไปในทางเดียวกัน จึงทำให้ราคาที่พยากรณ์ได้สามารถช่วยตัดสินใจในการประกอบการของผู้ผลิตหรือผู้ที่เกี่ยวข้องในอุตสาหกรรมนี้ได้

จากการศึกษาที่ผ่านมา พบว่า มีการศึกษาพยากรณ์ปริมาณและราคาการส่งออกน้ำตาล แต่ยังไม่เคยมีการศึกษามูลค่าการส่งออกน้ำตาล ดังนั้นการศึกษานี้ ผู้วิจัยมีความสนใจที่จะศึกษาวิธีการสร้างตัวแบบพยากรณ์มูลค่าการส่งออกน้ำตาลด้วยวิธีการทางสถิติทั้งหมด 7 วิธี ได้แก่ วิธีบ็อกซ์-เจนกินส์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของโฮลต์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีแนวโน้มแบบแฉก วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ จากนั้นคัดเลือกตัวแบบพยากรณ์ที่มีความแม่นยำมากที่สุด 1 ตัวแบบ โดยใช้เกณฑ์รากของค่าคลาดเคลื่อนกำลังสองเฉลี่ย (Root Mean Square Error: RMSE) ที่ต่ำที่สุด เพื่อให้ได้ตัวแบบพยากรณ์สำหรับการพยากรณ์มูลค่าการส่งออกน้ำตาลในอนาคต ซึ่งจะส่งผลดีต่อการตัดสินใจ การบริหารจัดการด้านความเสี่ยงต่าง ๆ ช่วยในการประเมินการคาดการณ์มูลค่าการส่งออกน้ำตาลล่วงหน้า รวมถึงเป็นแนวทางให้รัฐบาลสามารถออกนโยบายในการสนับสนุนการส่งออกน้ำตาลได้อย่างเหมาะสม ซึ่งจะนำรายได้มาสู่ประเทศไทยมากยิ่งขึ้น และทำให้เกิดการเจริญเติบโตทางเศรษฐกิจของประเทศต่อไป

วิธีการทดลอง

การศึกษานี้ดำเนินการตามขั้นตอนดังนี้

1. เรียบเรียงข้อมูลมูลค่าการส่งออกน้ำตาลของประเทศไทย (บาท) โดยใช้ค่าเฉลี่ยต่อเดือนจากเว็บไซต์ของสำนักงานเศรษฐกิจการเกษตร (The Office of Agricultural Economics, 2021) ตั้งแต่เดือนมกราคม 2554 ถึง

เดือนพฤศจิกายน 2563 จำนวน 119 เดือน หลังจากนั้นแบ่งข้อมูลออกเป็น 2 ชุด ชุดที่ 1 คือ ข้อมูลตั้งแต่เดือนมกราคม 2554 ถึงเดือนมิถุนายน 2563 จำนวน 114 เดือน สำหรับการสร้างตัวแบบพยากรณ์ ชุดที่ 2 ตั้งแต่เดือนกรกฎาคมถึงเดือนพฤศจิกายน 2563 จำนวน 5 เดือน สำหรับการเปรียบเทียบความแม่นยำของตัวแบบพยากรณ์ด้วยเกณฑ์ RMSE ที่ต่ำที่สุด

2. ตรวจสอบการเคลื่อนไหวจากแนวโน้มและอิทธิพลของฤดูกาลของอนุกรมเวลามูลค่าการส่งออกน้ำตาลชุดที่ 1 ดังนี้ ถ้าอนุกรมเวลา มีการแจกแจงปกติ (ตรวจสอบการแจกแจงปกติด้วยการทดสอบชาฟิโร-วิลค์: Shapiro-Wilk Test เนื่องจากข้อมูลในแต่ละกลุ่มที่จะตรวจสอบมีจำนวนไม่เกิน 50 ค่า) และมีความแปรปรวนเท่ากัน (ตรวจสอบความแปรปรวนเท่ากันด้วยการทดสอบของเลวินภายใต้การใช้มัธยฐาน: Levene's Test Based on Median) จะใช้สถิติอิงพารามิเตอร์ (Parametric Statistics) คือ การวิเคราะห์ความแปรปรวนทางเดียว (One-Way Analysis of Variance: ANOVA) แต่ถ้าอนุกรมเวลาไม่มีการแจกแจงปกติหรือมีความแปรปรวนไม่เท่ากัน จะใช้สถิติไม่อิงพารามิเตอร์ (Nonparametric Statistics) คือ การวิเคราะห์ความแปรปรวนทางเดียวโดยลำดับที่ของครัสคาล-วอลลิส (Kruskal-Wallis's One-Way Analysis of Variance by Rank) ถ้าผลการตรวจสอบพบว่าอนุกรมเวลามีเฉพาะการเคลื่อนไหวจากแนวโน้ม วิธีการพยากรณ์ที่เหมาะสม ได้แก่ วิธีบ็อกซ์-เจนกินส์ที่มีตัวแบบ Autoregressive Integrated Moving Average: ARIMA(p, d, q) วิธีการทำให้เรียบด้วยเลขชี้กำลังของโฮลต์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ และวิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีแนวโน้มแบบแฉก อนุกรมเวลามีเฉพาะอิทธิพลของฤดูกาล วิธีการพยากรณ์ที่เหมาะสม ได้แก่ วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย และอนุกรมเวลามีทั้งการเคลื่อนไหวจากแนวโน้มและอิทธิพลของฤดูกาล วิธีการพยากรณ์ที่เหมาะสม ได้แก่ วิธีบ็อกซ์-เจนกินส์ที่มีตัวแบบ Seasonal Autoregressive Integrated Moving Average: SARIMA(p, d, q)(P, D, Q), วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ (Ket-iam, 2005; Manmin, 2006) จากผลการตรวจสอบการเคลื่อนไหวของแนวโน้มและอิทธิพลของฤดูกาลในผลการวิจัยแสดงว่า มูลค่าการส่งออกน้ำตาลไม่มีการเคลื่อนไหวจากแนวโน้ม แต่มีอิทธิพลของฤดูกาล อย่างไรก็ตาม การศึกษาครั้งนี้จะพิจารณาวิธีการสร้างตัวแบบพยากรณ์โดยวิธีการพยากรณ์ทางสถิติที่มีความเหมาะสมกับอนุกรมเวลาทุกรูปแบบ ได้แก่ วิธีบ็อกซ์-เจนกินส์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของโฮลต์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีแนวโน้มแบบแฉก วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ เพื่อให้ครอบคลุมตัวแบบพยากรณ์ที่ดีที่สุด (Riansut, 2018) ตัวแบบพยากรณ์ทั้ง 7 วิธีดังกล่าว แสดงรายละเอียดดังตารางที่ 1

ตารางที่ 1 ตัวแบบพยากรณ์

วิธี ที่	วิธี พยากรณ์	ตัวแบบพยากรณ์	ลักษณะ อนุกรมเวลา
1	บ็อกซ์- เจนกินส์	$SARIMA(p, d, q)(P, D, Q)_s :$ $\hat{\phi}_p(B)\hat{\phi}_p(B^s)(1-B)^d(1-B^s)^D \hat{Y}_t = \hat{\delta} + \hat{\theta}_q(B)\hat{\theta}_q(B^s)e_t$ (Box <i>et al.</i> , 1994)	มีแนวโน้ม และฤดูกาล
2	โฮลต์	$\hat{Y}_{t+m} = a_t + b_t(m)$ โดยที่ $a_t = \alpha Y_t + (1-\alpha)(a_{t-1} + b_{t-1})$, $b_t = \gamma(a_t - a_{t-1}) + (1-\gamma)b_{t-1}$ (Manmin, 2006)	มีเพียง แนวโน้ม
3	บรวาน	$\hat{Y}_{t+m} = a_t + b_t \left[(m-1) + \frac{1}{\alpha} \right]$ โดยที่ $a_t = \alpha Y_t + (1-\alpha)a_{t-1}$, $b_t = \alpha(a_t - a_{t-1}) + (1-\alpha)b_{t-1}$ (Ket-iam, 2005)	มีเพียง แนวโน้ม
4	แดม	$\hat{Y}_{t+m} = a_t + b_t \sum_{i=1}^m \phi^i$ โดยที่ $a_t = \alpha Y_t + (1-\alpha)(a_{t-1} + \phi b_{t-1})$, $b_t = \gamma(a_t - a_{t-1}) + (1-\gamma)\phi b_{t-1}$ (Manmin, 2006)	มีเพียง แนวโน้ม
5	ฤดูกาล อย่างง่าย	$\hat{Y}_t = a_t + \hat{S}_t$ โดยที่ $a_t = \alpha(Y_t - \hat{S}_{t-s}) + (1-\alpha)a_{t-1}$, $\hat{S}_t = \delta(Y_t - a_t) + (1-\delta)\hat{S}_{t-s}$ (Ket-iam, 2005)	มีเพียง ฤดูกาล
6	วินเทอร์ แบบบวก	$\hat{Y}_{t+m} = (a_t + b_t m) + \hat{S}_t$ โดยที่ $a_t = \alpha(Y_t - \hat{S}_{t-s}) + (1-\alpha)(a_{t-1} + b_{t-1})$, $b_t = \gamma(a_t - a_{t-1}) + (1-\gamma)b_{t-1}$, $\hat{S}_t = \delta(Y_t - a_t) + (1-\delta)\hat{S}_{t-s}$ (Ket-iam, 2005)	มีแนวโน้ม และฤดูกาล
7	วินเทอร์ แบบคูณ	$\hat{Y}_{t+m} = (a_t + b_t m)\hat{S}_t$ โดยที่ $a_t = \alpha \frac{Y_t}{\hat{S}_{t-s}} + (1-\alpha)(a_{t-1} + b_{t-1})$, $b_t = \gamma(a_t - a_{t-1}) + (1-\gamma)b_{t-1}$, $\hat{S}_t = \delta \frac{Y_t}{a_t} + (1-\delta)\hat{S}_{t-s}$ (Ket-iam, 2005)	มีแนวโน้ม และฤดูกาล

จากตารางที่ 1 มีความหมายของสัญลักษณ์ต่าง ๆ ดังนี้

$\hat{Y}_t$ และ $\hat{Y}_{t+m}$ แทนค่าพยากรณ์ ณ เวลา t และเวลา $t+m$ ตามลำดับ โดยที่ m แทนจำนวนช่วงเวลาที่ต้องการพยากรณ์ไปข้างหน้า

$\hat{\delta} = \hat{\mu}\hat{\phi}_p(B)\hat{\phi}_p(B^s)$ แทนค่าคงตัว (Constant) โดยที่ $\hat{\mu}$ แทนค่าเฉลี่ยของอนุกรมเวลาที่คงที่ (Stationary)

$\hat{\phi}_p(B) = 1 - \hat{\phi}_1 B - \hat{\phi}_2 B^2 - \dots - \hat{\phi}_p B^p$ แทนตัวดำเนินการสหสัมพันธ์ในลำดับที่ p กรณีไม่มีฤดูกาล (Non-Seasonal Autoregressive Operator of Order p : AR(p))

$\hat{\phi}_p(B^s) = 1 - \hat{\phi}_1 B^s - \hat{\phi}_2 B^{2s} - \dots - \hat{\phi}_p B^{ps}$ แทนตัวดำเนินการสหสัมพันธ์ในลำดับที่ P กรณีมีฤดูกาล (Seasonal Autoregressive Operator of Order P : SAR(P))

$\hat{\theta}_q(B) = 1 - \hat{\theta}_1 B - \hat{\theta}_2 B^2 - \dots - \hat{\theta}_q B^q$ แทนตัวดำเนินการเฉลี่ยเคลื่อนที่อันดับที่ q กรณีไม่มีฤดูกาล (Non-Seasonal Moving Average Operator of Order q : MA(q))

$\hat{\Theta}_Q(B^s) = 1 - \hat{\Theta}_1 B^s - \hat{\Theta}_2 B^{2s} - \dots - \hat{\Theta}_Q B^{Qs}$ แทนตัวดำเนินการเฉลี่ยเคลื่อนที่อันดับที่ Q กรณีมีฤดูกาล (Seasonal

Moving Average Operator of Order Q: SMA(Q))

t แทนช่วงเวลา ซึ่งมีค่าตั้งแต่ 1 ถึง n_1 โดยที่ n_1 แทนจำนวนข้อมูลในอนุกรมเวลาชุดที่ 1 ($n_1 = 114$)

s แทนจำนวนฤดูกาล ซึ่งอนุกรมเวลามูลค่าการส่งออกน้ำตาลเป็นข้อมูลรายเดือน ดังนั้น $s = 12$

d และ D แทนลำดับที่ของการหาผลต่างและผลต่างฤดูกาล ตามลำดับ

B แทนตัวดำเนินการถอยหลัง (Backward Operator) โดยที่ $B^s Y_t = Y_{t-s}$

a_t และ b_t แทนค่าประมาณระยะตัดแกน Y และความชันของแนวโน้ม ณ เวลา t ตามลำดับ

α , γ , ϕ และ δ แทนค่าคงตัวการทำให้เรียบ โดยที่ $0 < \alpha < 1$, $0 < \gamma < 1$, $0 < \phi < 1$ และ $0 < \delta < 1$

3. ตรวจสอบข้อสมมุติ (Assumption) ของตัวแบบพยากรณ์ คือ อนุกรมเวลาของค่าคลาดเคลื่อนจากการพยากรณ์มีการแจกแจงปกติ ตรวจสอบโดยใช้การทดสอบคอลโมโกรอฟ-สมิรโนฟ (Kolmogorov-Smirnov's Test: KS Test) มีการเคลื่อนไหวเป็นอิสระกัน ตรวจสอบโดยใช้การทดสอบรันส์ (Runs Test) มีค่าเฉลี่ยเท่ากับศูนย์ ตรวจสอบโดยใช้การทดสอบที (t-Test) และมีความแปรปรวนเท่ากันทุกช่วงเวลา ตรวจสอบโดยใช้การทดสอบของเลวินภายใต้การใช้มัธยฐาน หากพบว่าข้อสมมุติข้อใดข้อหนึ่งไม่เป็นจริงจะสรุปว่าตัวแบบพยากรณ์ไม่เหมาะสมและไม่สมควรนำไปใช้ในการพยากรณ์ต่อไป

4. เปรียบเทียบความแม่นยำของตัวแบบพยากรณ์ โดยการเปรียบเทียบมูลค่าการส่งออกน้ำตาลของข้อมูลชุดที่ 2 ตั้งแต่เดือนกรกฎาคมถึงเดือนพฤศจิกายน 2563 จำนวน 5 เดือน ($n = 5$) กับค่าพยากรณ์ เพื่อคำนวณค่า RMSE โดยตัวแบบพยากรณ์ที่ให้ค่า RMSE ต่ำที่สุด จัดเป็นตัวแบบที่มีความแม่นยำมากที่สุด เนื่องจากให้ค่าพยากรณ์ที่มีความแตกต่างกับข้อมูลจริงน้อยที่สุด สูตร RMSE แสดงดังนี้ (Ket-iam, 2005)

$$RMSE = \sqrt{\frac{1}{n} \sum_{t=1}^n e_t^2}$$

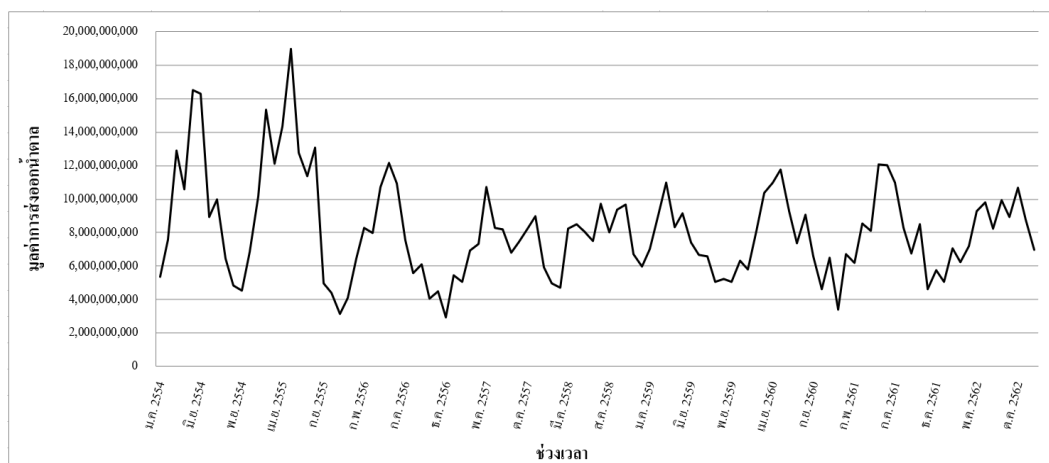
เมื่อ $e_t = Y_t - \hat{Y}_t$ แทนค่าคลาดเคลื่อนจากการพยากรณ์ ณ เวลา t

Y_t และ $\hat{Y}_t$ แทนอนุกรมเวลาและค่าพยากรณ์ ณ เวลา t ตามลำดับ

ผลและวิจารณ์ผลการทดลอง

จากการพิจารณาลักษณะการเคลื่อนไหวของอนุกรมเวลามูลค่าการส่งออกน้ำตาลชุดที่ 1 ตั้งแต่เดือนมกราคม 2554 ถึงเดือนมิถุนายน 2563 จำนวน 114 เดือน ดังรูปที่ 1 พบว่า มูลค่าการส่งออกน้ำตาลไม่มีการเคลื่อนไหวจากแนวโน้ม หรือการเคลื่อนไหวของมูลค่าการส่งออกน้ำตาลเป็นไปในลักษณะคงที่ ไม่มีแนวโน้มเพิ่มขึ้นและลดลง แต่มีความผันผวนของมูลค่าการส่งออกน้ำตาลในแต่ละเดือน และจากการทดสอบสมมุติฐานเพื่อตรวจสอบการเคลื่อนไหวจากแนวโน้มและอิทธิพลของฤดูกาล พบว่า อนุกรมเวลามูลค่าการส่งออกน้ำตาลในแต่ละปีมีการแจกแจงปกติ แต่มีความแปรปรวนไม่เท่ากันที่ระดับนัยสำคัญ 0.05 ดังนั้นจึงใช้การวิเคราะห์ความแปรปรวนทางเดียวโดยลำดับที่ของครัสคาล – วอลลิสในการตรวจสอบการเคลื่อนไหวจากแนวโน้ม พบว่า อนุกรม

เวลามีค่ามัธยฐานในแต่ละปีไม่แตกต่างกัน ($\chi^2 = 5.665$, p-value = 0.685) นั่นคือ อนุกรมเวลาไม่มีการเคลื่อนไหวจากแนวโน้ม และอนุกรมเวลามูลค่าการส่งออกน้ำตาลในแต่ละเดือนไม่มีการแจกแจงปกติ แต่มีความแปรปรวนเท่ากันที่ระดับนัยสำคัญ 0.05 ดังนั้นจึงใช้การวิเคราะห์ความแปรปรวนทางเดียวโดยลำดับที่ของครัสคอล-วอลลิส ในการตรวจสอบอิทธิพลของฤดูกาล พบว่า อนุกรมเวลามีค่ามัธยฐานในแต่ละเดือนแตกต่างกันที่ระดับนัยสำคัญ 0.05 ($\chi^2 = 46.532$, p-value < 0.0001) นั่นคือ อนุกรมเวลามีอิทธิพลของฤดูกาล เนื่องจากอนุกรมเวลามูลค่าการส่งออกน้ำตาลไม่มีการเคลื่อนไหวจากแนวโน้ม แต่มีอิทธิพลของฤดูกาล ดังนั้นวิธีการพยากรณ์ที่มีความเหมาะสมคือ วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย อย่างไรก็ตาม เมื่อพิจารณาค่า RMSE ของข้อมูลชุดที่ 1 ดังตารางที่ 2 พบว่า วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก ซึ่งเหมาะสมกับอนุกรมเวลาที่มีการเคลื่อนไหวจากแนวโน้มและอิทธิพลของฤดูกาล มีค่า RMSE ต่ำที่สุด (RMSE = 1,871,486,389) รองลงมาคือ วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย ซึ่งเหมาะสมกับอนุกรมเวลาที่มีเพียงอิทธิพลของฤดูกาลเท่านั้น (RMSE = 1,873,495,894) ดังนั้นการศึกษานี้จะพิจารณาวิธีการสร้างตัวแบบพยากรณ์ที่เหมาะสมกับอนุกรมเวลาทุกรูปแบบ เพื่อให้ครอบคลุมตัวแบบพยากรณ์ที่ดีที่สุด



รูปที่ 1 ลักษณะการเคลื่อนไหวของอนุกรมเวลามูลค่าการส่งออกน้ำตาล ตั้งแต่เดือนมกราคม 2554 ถึงเดือนมิถุนายน 2563

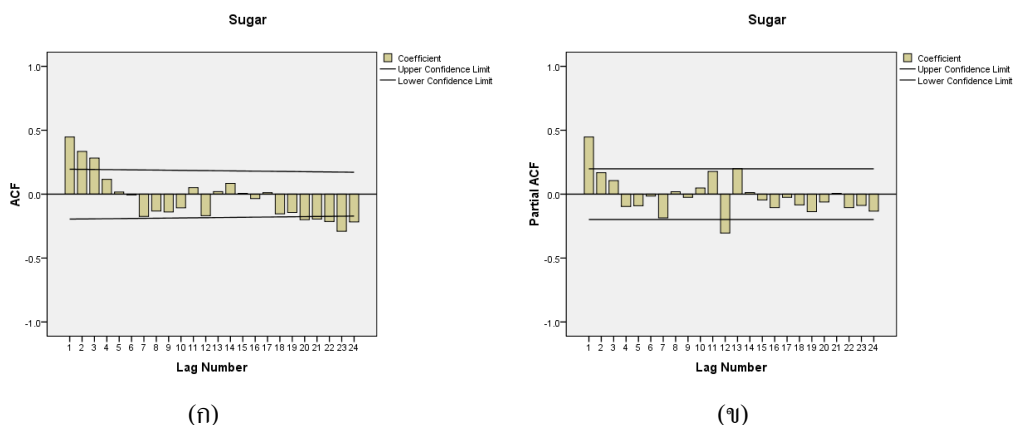
ตารางที่ 2 ค่า RMSE ของข้อมูลชุดที่ 1

วิธีพยากรณ์	บ็อกซ์-เจนกินส์	โซลต์	บราวน์	แดม
RMSE	2,099,373,319	2,323,687,226	2,590,797,343	2,314,183,280
วิธีพยากรณ์	ฤดูกาลอย่างง่าย	วินเทอร์แบบบวก	วินเทอร์แบบคูณ	
RMSE	<u>1,873,495,894</u>	<u>1,871,486,389</u>	2,070,839,781	

จากการตรวจสอบที่พบว่า อนุกรมเวลามูลค่าการส่งออกน้ำตาลไม่มีการเคลื่อนไหวจากแนวโน้ม แต่มีอิทธิพลของฤดูกาล ดังนั้นผู้วิจัยจึงแปลงข้อมูลด้วยการหาผลต่างฤดูกาลลำดับที่ 1 ($D = 1$) เพื่อสร้างตัวแบบพยากรณ์ โดยวิธีบ็อกซ์-เจนกินส์ ได้กราฟ Autocorrelation Function (ACF) และกราฟ Partial Autocorrelation Function (PACF) ของอนุกรมเวลาที่แปลงข้อมูลแล้ว แสดงดังรูปที่ 2 (ก) และ (ข) ตามลำดับ ซึ่งพบว่า อนุกรมเวลามีลักษณะคงที่ เพราะค่าสัมประสิทธิ์สหสัมพันธ์ในตัวจากกราฟ ACF รูปที่ 2 (ก) และค่าสัมประสิทธิ์สหสัมพันธ์ในตัวบางส่วนจากกราฟ PACF รูปที่ 2 (ข) ตกอยู่ภายในขอบเขตความเชื่อมั่นที่กำหนด ($\pm 2/\sqrt{n_1}$) ยกเว้นเฉพาะ Lag ที่ 1, 2, 3 และกลุ่มของ Lag ฤดูกาล Lag ที่ 20 – 24 ของกราฟ ACF รูปที่ 2 (ก) มีค่าสัมประสิทธิ์สหสัมพันธ์ในตัวเกินขอบเขต จึงกำหนดตัวแบบพยากรณ์ที่เป็นไปได้เริ่มต้น มีค่า $q = 3$, $Q = 1$ และ Lag ที่ 1, 12 ของกราฟ PACF รูปที่ 2 (ข) มีค่าสัมประสิทธิ์สหสัมพันธ์ในตัวบางส่วนเกินขอบเขต จึงกำหนดตัวแบบพยากรณ์ที่เป็นไปได้เริ่มต้น มีค่า $p = 1$, $P = 1$ ดังนั้นตัวแบบพยากรณ์เริ่มต้น คือ ตัวแบบ SARIMA(1, 0, 3)(1, 1, 1)₁₂ จากการคัดเลือกตัวแบบให้เหลือเฉพาะพารามิเตอร์ที่มีนัยสำคัญที่ระดับ 0.05 พบว่า ตัวแบบพยากรณ์ที่เหมาะสม คือ ตัวแบบ SARIMA(1, 0, 0)(0, 1, 1)₁₂ ไม่มีพจน์ค่าคงตัว ซึ่งสามารถเขียนเป็นตัวแบบของวิธีบ็อกซ์-เจนกินส์ได้ดังนี้

$$\begin{aligned}(1 - \phi_1 B)(1 - B^{12})Y_t &= (1 - \Theta_1 B^{12})\varepsilon_t \\ (1 - B^{12} - \phi_1 B + \phi_1 B^{13})Y_t &= \varepsilon_t - \Theta_1 \varepsilon_{t-12} \\ Y_t &= \phi_1 (Y_{t-1} - Y_{t-13}) + Y_{t-12} + \varepsilon_t - \Theta_1 \varepsilon_{t-12}\end{aligned}$$

จากการแทนค่าประมาณพารามิเตอร์ จะได้ตัวแบบพยากรณ์ของวิธีบ็อกซ์-เจนกินส์และวิธีการพยากรณ์อื่น ๆ แสดงดังตารางที่ 3 และค่าดัชนีฤดูกาลจากวิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ แสดงดังตารางที่ 4 ซึ่งพบว่า วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่ายและวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก แสดงว่า มูลค่าการส่งออกน้ำตาลมีค่าสูงในช่วงเดือนมีนาคมถึงสิงหาคม เนื่องจากดัชนีฤดูกาลมีค่ามากกว่า 0 และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ แสดงว่า มูลค่าการส่งออกน้ำตาลมีค่าสูงในช่วงเดือนพฤษภาคมถึงมิถุนายน เนื่องจากดัชนีฤดูกาลมีค่ามากกว่า 1 ผลการตรวจสอบข้อสมมุติของตัวแบบพยากรณ์ แสดงดังตารางที่ 5 ซึ่งพบว่า ตัวแบบพยากรณ์ที่สร้างขึ้นทั้ง 7 วิธี มีข้อสมมุติเป็นจริงทุกข้อที่ระดับนัยสำคัญ 0.05 กล่าวคือ อนุกรมเวลาของค่าตลาดเคลื่อนไหวจากการพยากรณ์มีการแจกแจงปกติ เคลื่อนไหวเป็นอิสระกัน มีค่าเฉลี่ยเท่ากับศูนย์ และความแปรปรวนเท่ากันทุกช่วงเวลา



รูปที่ 2 กราฟ ACF และ PACF ของอนุกรมเวลามูลค่าการส่งออกน้ำตาล เมื่อแปลงข้อมูลด้วยผลต่างฤดูกาลลำดับที่ 1

ตารางที่ 3 ผลการสร้างตัวแบบพยากรณ์

วิธีที่	วิธีพยากรณ์	ตัวแบบพยากรณ์
1	บ็อกซ์-เจนกินส์	$\hat{Y}_t = 0.61148(Y_{t-1} - Y_{t-13}) + Y_{t-12} - 0.73324e_{t-12}$ โดยที่ Y_{t-j} แทนอนุกรมเวลา ณ เวลา $t-j$ และ e_{t-j} แทนค่าคลาดเคลื่อนจากการพยากรณ์ ณ เวลา $t-j$
2	โสมล์	$\hat{Y}_{t+m} = 4,833,966,655.71732 - 18,599,127.77542(m)$ โดยที่ $m = 1$ แทนเดือนกรกฎาคม 2563
3	บรานน์	$\hat{Y}_{t+m} = 5,719,091,961.84232 - 756,737,309.23653 \left[(m-1) + \frac{1}{0.54881} \right]$ โดยที่ $m = 1$ แทนเดือนกรกฎาคม 2563
4	แดม	$\hat{Y}_{t+m} = 5,456,281,444.30739 - 1,665,208,404.40762 \sum_{i=1}^m (0.34861)^i$ โดยที่ $m = 1$ แทนเดือนกรกฎาคม 2563
5	ฤดูกาลอย่างง่าย	$\hat{Y}_t = 3,234,742,234.44463 + \hat{S}_t$ โดยที่ $\hat{S}_t$ แสดงดังตารางที่ 4
6	วินเทอร์แบบบวก	$\hat{Y}_{t+m} = (3,134,853,742.62817 - 19,005,787.21141m) + \hat{S}_t$ โดยที่ $m = 1$ แทนเดือนกรกฎาคม 2563 และ $\hat{S}_t$ แสดงดังตารางที่ 4
7	วินเทอร์แบบคูณ	$\hat{Y}_{t+m} = (5,116,323,417.74919 - 17,133,371.75470m)\hat{S}_t$ โดยที่ $m = 1$ แทนเดือนกรกฎาคม 2563 และ $\hat{S}_t$ แสดงดังตารางที่ 4

ตารางที่ 4 คำนวณฤดูกาลของอนุกรมเวลามูลค่าการส่งออกน้ำตาล จากวิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ

เดือน	$\hat{S}_t$ ของวิธีฤดูกาลอย่างง่าย	$\hat{S}_t$ ของวิธีวินเทอร์แบบบวก	$\hat{S}_t$ ของวิธีวินเทอร์แบบคูณ
มกราคม	-1,794,874,056	-1,790,141,417	0.57646
กุมภาพันธ์	-270,386,075	-246,648,750	0.74566
มีนาคม	<u>1,147,583,411</u>	<u>1,190,322,819</u>	0.94533
เมษายน	<u>1,216,465,122</u>	<u>1,278,203,923</u>	0.98521
พฤษภาคม	<u>3,521,095,907</u>	<u>3,601,826,879</u>	<u>1.23931</u>
มิถุนายน	<u>1,822,346,519</u>	<u>1,922,046,167</u>	<u>1.06668</u>
กรกฎาคม	<u>708,664,280</u>	<u>713,466,580</u>	0.99980
สิงหาคม	<u>527,817,327</u>	<u>551,575,933</u>	0.95057
กันยายน	-1,209,177,856	-1,166,430,799	0.75988
ตุลาคม	-1,376,351,320	-1,314,605,988	0.71835
พฤศจิกายน	-2,209,741,586	-2,128,994,494	0.57202
ธันวาคม	-2,710,375,073	-2,610,623,378	0.50837

ตารางที่ 5 ผลการตรวจสอบข้อสมมุติของตัวแบบพยากรณ์

วิธีที่	วิธีพยากรณ์	KS test	p-value	Runs test	p-value	t-test	p-value	Levene statistic	p-value
1	บ็อกซ์-เจนกินส์	0.578	0.892	0.603	0.546	-1.114	0.268	0.451	0.928
2	โซลต์	0.572	0.899	1.407	0.159	0.058	0.954	0.265	0.991
3	บราวน์	0.603	0.861	1.407	0.159	-0.338	0.736	0.520	0.885
4	แดม	0.566	0.906	1.407	0.159	-0.049	0.961	0.286	0.987
5	ฤดูกาลอย่างง่าย	0.567	0.905	0.201	0.841	-0.315	0.753	0.287	0.987
6	วินเทอร์แบบบวก	0.591	0.876	0.603	0.546	-0.169	0.866	0.287	0.987
7	วินเทอร์แบบคูณ	0.383	0.999	-1.809	0.070	-0.306	0.760	0.290	0.986

เมื่อใช้ตัวแบบพยากรณ์ที่สร้างขึ้นในตารางที่ 3 สำหรับการพยากรณ์มูลค่าการส่งออกน้ำตาลชุดที่ 2 ตั้งแต่เดือนกรกฎาคมถึงเดือนพฤศจิกายน 2563 จากนั้นเปรียบเทียบค่าพยากรณ์กับค่าจริงโดยการคำนวณค่า RMSE ได้ผล

แสดงดังตารางที่ 6 ผลการเปรียบเทียบพบว่า ตัวแบบพยากรณ์ของวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณมีความแม่นยำมากที่สุดในการพยากรณ์มูลค่าการส่งออกน้ำตาล เนื่องจากมีค่า RMSE ต่ำที่สุด (RMSE = 680,806,354) หรือมีความผิดพลาดในการพยากรณ์มูลค่าการส่งออกน้ำตาล 680,806,354 บาท และตัวแบบพยากรณ์ที่มีความแม่นยำรองลงมา คือ ตัวแบบพยากรณ์ของวิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ (RMSE = 695,330,560) หรือมีความผิดพลาดในการพยากรณ์มูลค่าการส่งออกน้ำตาล 695,330,560 บาท

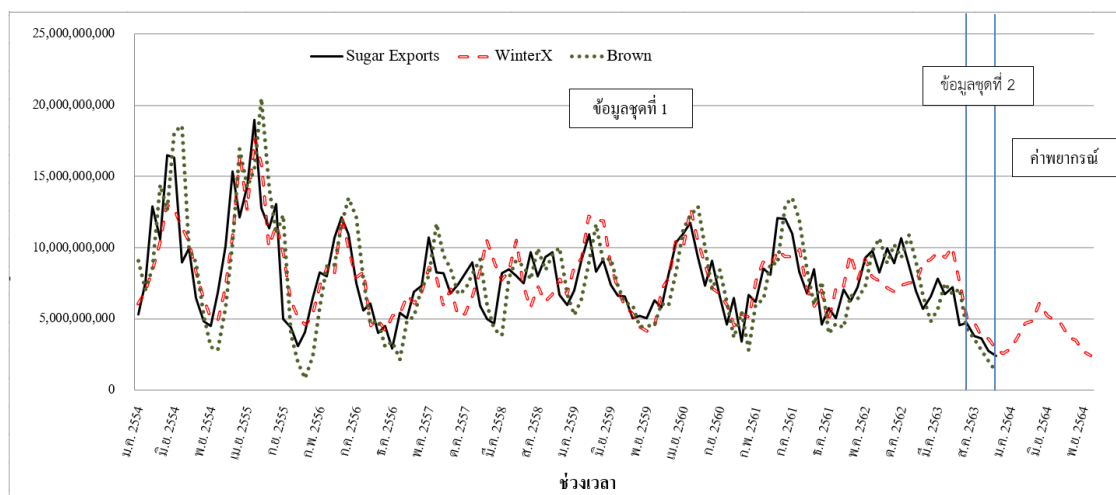
ตารางที่ 6 ค่า RMSE ของข้อมูลชุดที่ 2

วิธีพยากรณ์	บ็อกซ์-เจนกินส์	โซลต์	บราวน์	แถม
RMSE	2,974,057,968	1,547,488,130	695,330,560	1,407,647,097
วิธีพยากรณ์	ฤดูกาลอย่างง่าย	วินเทอร์แบบบวก	วินเทอร์แบบคูณ	
RMSE	1,075,529,292	1,175,280,036	680,806,354	

จากผลการตรวจสอบแนวโน้มและอิทธิพลของฤดูกาลของอนุกรมเวลามูลค่าการส่งออกน้ำตาลชุดที่ 1 ที่พบว่า อนุกรมเวลาชุดนี้ไม่มีการเคลื่อนไหวจากแนวโน้ม แต่มีอิทธิพลของฤดูกาล ดังนั้นวิธีการพยากรณ์ที่เหมาะสมควรจะเป็นวิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย (Ket-iam, 2005; Manmin, 2006) ขัดแย้งกับผลการศึกษารุ่นนี้ที่พบว่า วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณให้ตัวแบบพยากรณ์ที่มีความแม่นยำมากที่สุด แต่การศึกษารุ่นนี้ได้ผลสอดคล้องกับการศึกษาของ Riansut (2018) ที่สรุปว่าควรศึกษาวิธีการพยากรณ์ที่หลากหลายเพื่อให้ครอบคลุมตัวแบบพยากรณ์ที่ดีที่สุด ดังนั้นการศึกษาเกี่ยวกับการพยากรณ์ทุกครั้งควรพิจารณาวิธีการพยากรณ์หลาย ๆ วิธี และใช้เกณฑ์การเปรียบเทียบความแม่นยำของตัวแบบพยากรณ์ในการคัดเลือกวิธีการพยากรณ์ที่เหมาะสม เช่น เกณฑ์ RMSE ซึ่งเกณฑ์นี้จะพิจารณาความแตกต่างระหว่างข้อมูลจริงและค่าพยากรณ์ ผู้วิจัยเชื่อว่ามีค่าน่าเชื่อถือมากกว่าการตรวจสอบแนวโน้มและฤดูกาล เนื่องจากการตรวจสอบแนวโน้มและฤดูกาลจะพิจารณาข้อมูลทั้งหมด โดยบางช่วงเวลาข้อมูลอาจมีแนวโน้มหรือมีฤดูกาล แต่เมื่อตรวจสอบข้อมูลทั้งหมดอาจตรวจไม่พบแนวโน้ม/ฤดูกาลได้

การศึกษารุ่นนี้ได้ใช้เกณฑ์ RMSE ในการเปรียบเทียบความแม่นยำของตัวแบบพยากรณ์ที่สร้างขึ้นทั้ง 7 วิธี สำหรับเกณฑ์การเปรียบเทียบความแม่นยำนั้นยังคงมีเกณฑ์อื่น ๆ อีก เช่น เกณฑ์ร้อยละค่าคลาดเคลื่อนสัมบูรณ์เฉลี่ย (Mean Absolute Percentage Error: MAPE) ซึ่งการใช้เกณฑ์การเปรียบเทียบความแม่นยำแตกต่างกันอาจส่งผลให้การคัดเลือกตัวแบบพยากรณ์ที่ดีที่สุดเหมือนหรือแตกต่างกันได้ อย่างไรก็ตาม การศึกษารุ่นนี้ได้ทดลองใช้เกณฑ์ MPAE แล้วพบว่า วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณยังคงมีความแม่นยำมากที่สุด (MAPE = 18.9745)

ผลการเปรียบเทียบอนุกรมเวลามูลค่าการส่งออกน้ำตาลกับค่าพยากรณ์จากวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณและวิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ ดังรูปที่ 3 พบว่า วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณมีความแม่นยำมากกว่าวิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ สอดคล้องกับการพิจารณาจากเกณฑ์ RMSE เนื่องจากวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณให้ค่าพยากรณ์ที่มีความแม่นยำสูงในช่วงปี 2554 – 2556 และ 2559 – 2561 แต่ช่วงปี 2557 – 2558 และ 2562 – 2563 จะมีความแม่นยำค่อนข้างต่ำ ขณะที่วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ให้ค่าพยากรณ์ที่ค่อนข้างแตกต่างจากค่าจริงตลอดทุกช่วงเวลา และเมื่อพิจารณาค่าพยากรณ์ในข้อมูลชุดที่ 2 พบว่า วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณให้ค่าพยากรณ์สูงกว่าข้อมูลจริงทุกค่า ขณะที่วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ให้ค่าพยากรณ์ต่ำกว่าข้อมูลจริงทุกค่า



รูปที่ 3 การเปรียบเทียบอนุกรมเวลามูลค่าการส่งออกน้ำตาลกับค่าพยากรณ์จากวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณและวิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์

เมื่อใช้วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณในการพยากรณ์มูลค่าการส่งออกน้ำตาลตั้งแต่เดือนธันวาคม 2563 ถึงเดือนธันวาคม 2564 ได้ผลแสดงดังตารางที่ 7 และรูปที่ 3 ซึ่งพบว่า มูลค่าการส่งออกน้ำตาลมีการเคลื่อนไหวจากแนวโน้มในทิศทางลดลง อย่างไรก็ตาม ค่าพยากรณ์ของวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณมีความแตกต่างจากข้อมูลจริงอยู่บ้าง อาจเนื่องมาจากการศึกษาครั้งนี้ได้พิจารณาเพียงปัจจัยเวลาเท่านั้นในการสร้างตัวแบบพยากรณ์ ซึ่งมูลค่าการส่งออกน้ำตาลมีการเปลี่ยนแปลงอยู่เสมอ และการเปลี่ยนแปลงอาจเกิดจากปัจจัยอื่น ๆ นอกเหนือจากปัจจัยเวลา ดังนั้น เมื่อมีปัจจัยที่มีผลกระทบต่อเปลี่ยนแปลงของมูลค่า

การส่งออกน้ำตาลหรือมีข้อมูลที่เป็นปัจจุบันมากขึ้น ผู้วิจัยควรนำมาปรับปรุงตัวแบบเพื่อให้ได้ตัวแบบพยากรณ์ที่มีความแม่นยำมากยิ่งขึ้น สำหรับใช้ในการพยากรณ์มูลค่าการส่งออกน้ำตาลในอนาคตต่อไป

ตารางที่ 7 ค่าพยากรณ์ของมูลค่าการส่งออกน้ำตาล (บาท) ตั้งแต่เดือนธันวาคม 2563 ถึงเดือนธันวาคม 2564

ช่วงเวลา	ค่าพยากรณ์	ช่วงเวลา	ค่าพยากรณ์	ช่วงเวลา	ค่าพยากรณ์
ธ.ค. 2563	2,548,702,331	พ.ค. 2564	6,107,162,463	ต.ค. 2564	3,478,378,483
ม.ค. 2564	2,880,239,687	มิ.ย. 2564	5,238,148,620	พ.ย. 2564	2,760,048,755
ก.พ. 2564	3,712,843,600	ก.ค. 2564	4,892,604,813	ธ.ค. 2564	2,444,182,145
มี.ค. 2564	4,690,835,016	ส.ค. 2564	4,635,436,824		
เม.ย. 2564	4,871,862,106	ก.ย. 2564	3,692,485,532		

สรุปผลการทดลอง

การศึกษานี้ได้นำเสนอวิธีการสร้างและคัดเลือกตัวแบบพยากรณ์ที่เหมาะสมกับอนุกรมเวลามูลค่าการส่งออกน้ำตาลของประเทศไทยด้วยวิธีการทางสถิติ 7 วิธี ได้แก่ วิธีบ็อกซ์-เจนกินส์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของโฮลด์ วิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์ วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีแนวโน้มแบบแฉก วิธีการทำให้เรียบด้วยเลขชี้กำลังที่มีฤดูกาลอย่างง่าย วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบบวก และวิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณ ผลการศึกษาพบว่า วิธีการทำให้เรียบด้วยเลขชี้กำลังของวินเทอร์แบบคูณมีความแม่นยำมากที่สุด (RMSE = 680,806,354) โดยมีตัวแบบพยากรณ์ ดังนี้

$$\hat{Y}_{t+m} = (5,116,323,417.74919 - 17,133,371.75470m)\hat{S}_t$$

และวิธีการทำให้เรียบด้วยเลขชี้กำลังของบราวน์มีความแม่นยำรองลงมา (RMSE = 695,330,560) โดยมีตัวแบบพยากรณ์ ดังนี้

$$\hat{Y}_{t+m} = 5,719,091,961.84232 - 756,737,309.23653 \left[(m-1) + \frac{1}{0.54881} \right]$$

โดยที่ $m = 1$ แทนเดือนกรกฎาคม 2563 และ $\hat{S}_t$ แทนดัชนีฤดูกาล

เอกสารอ้างอิง

- Box, G. E. P., Jenkins, G. M., & Reinsel, G. C. (1994). *Time Series Analysis: Forecasting and Control* (3rd ed.). Prentice Hall.
- Chantabutr, P. (2004). *Sugar export-price forecasting by ARIMA method* [Unpublished Master's thesis]. Chiangmai University. (in Thai)
- Foundation for Agricultural and Environmental Conservation in Thailand. (2021). *Sugarcane* <https://www.aecth.org/upload/13823/kVLyMFsUd8.pdf> (in Thai)
- Kavitanon, S. (2004). *Sugar industry in Thailand*. <https://www.lib.ku.ac.th/KUCONF/KC1507004.pdf> (in Thai)
- Ket-iam, S. (2005). *Forecasting Technique* (2nd ed.). Thaksin University. (in Thai)
- Luangtong, N., & Kantanantha, N. (2016). Selection of the appropriate agricultural yield forecasting models. *Thai Science and Technology Journal*, 24(3), 370-381. (in Thai)
- Manmin, M. (2006). *Time Series and Forecasting*. Foreprinting. (in Thai)
- Riansut, W. (2018). Comparison of tangerine prices forecast model by exponential smoothing methods. *Thai Journal of Science and Technology*, 7(5), 460-470. (in Thai)
- Teppiam, T., & Kittichotipanit, N. (2015). Comparison of forecasting models for quantity of exporting Thailand's raw sugar using time series analysis and multiple linear regression analysis. *Journal of Science Ladkrabang*, 24(2), 77-92. (in Thai)
- The Office of Agricultural Economics. (2021). *Export Statistics of Sugar from 2011 to 2020*. https://impexp.oae.go.th/service/export.php?S_YEAR=2554&E_YEAR=2564&PRODUCT_GROUP=5254&PRODUCT_ID=&wf_search=&WF_SEARCH=Y#export (in Thai)

Research Article

ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุมผลรวมสะสมสำหรับตัวแบบ SARX(P, 1)_L โดยวิธีสมการปริพันธ์เชิงตัวเลข

Average Run Length of CUSUM Control Chart for SARX(P,1)_L Model using Numerical Integral Equation Method

พรพิมล บุรณพล¹ และ สุวิมล พันธุ์แย้ม^{1*}

Phornpimon Buranapol¹ and Suvimol Phanyaem^{1*}

¹ภาควิชาสถิติประยุกต์ คณะวิทยาศาสตร์ประยุกต์ มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ บางซื่อ กรุงเทพฯ 10800 ประเทศไทย

¹Department of Applied Statistics, Faculty of Applied Science, King Mongkut's University of Technology North Bangkok, Bangsue, Bangkok 10800, Thailand

*E-mail: suvimol.p@sci.kmutnb.ac.th

Received: 05/11/2021; Revised: 07/04/2022; Accepted: 26/04/2022

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อประมาณค่าความยาวรันเฉลี่ยของแผนภูมิควบคุมผลรวมสะสม (Cumulative Sum: CUSUM) สำหรับตัวแบบอัตโนมัติแบบมีฤดูกาลที่มีตัวแปรภายนอก (SARX(P,1)_L) เมื่อค่าความคลาดเคลื่อนมีการแจกแจงแบบเลขชี้กำลัง โดยวิธีสมการปริพันธ์เชิงตัวเลข (Numerical Integration Equation: NIE) 4 วิธี ได้แก่ วิธีกฎซิมป์สัน (Simpson Rule) วิธีกฎค่ากลาง (Midpoint Rule) วิธีกฎสี่เหลี่ยมคางหมู (Trapezoidal Rule) และวิธีกฎเกาส์ (Gaussian Rule) เปรียบเทียบประสิทธิภาพของแผนภูมิควบคุมวัดจากค่าความยาวรันเฉลี่ย (Average Run Length : ARL) เมื่อกระบวนการไม่อยู่ภายใต้การควบคุม โดยทำการเปรียบเทียบค่า ARL ที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลข ทั้ง 4 วิธี และเปรียบเทียบเวลาที่ใช้ในการประมวลผล (CPU Times) ผลการวิจัยพบว่า ค่า ARL ที่ได้จากวิธีกฎซิมป์สัน วิธีกฎค่ากลาง วิธีกฎสี่เหลี่ยมคางหมู และวิธีกฎเกาส์ ให้ค่าความยาวรันเฉลี่ยไม่แตกต่างกัน ในขณะที่เวลาที่ใช้ในการประมวลผลโดยวิธีกฎค่ากลางมีค่าน้อยที่สุดในบรรดาทั้ง 4 วิธี

คำสำคัญ: ค่าความยาวรันเฉลี่ย, แผนภูมิควบคุมผลรวมสะสม, วิธีสมการปริพันธ์เชิงตัวเลข, ตัวแบบอัตโนมัติลดทอนแบบมีฤดูกาลที่มีตัวแปรภายนอก

Abstract

The main objective of this paper is to approximate the Average Run Length (ARL) of Cumulative Sum (CUSUM) control chart for a Seasonal Autoregressive with exogenous variable model; SARX(P,1)₄ with Exponential white noise using the Numerical Integral Equation (NIE) method including Simpson rule, Midpoint rule, Trapezoidal rule and Gauss-Legendre Quadrature rule. The control chart's performance are measured by Average Run Length (ARL) when the process is out of control by comparing the ARL values obtained from the four NIE methods and comparing the central processing unit processing time (CPU Times). The results show that the ARL value obtained from the NIE methods by Simpson rule, Midpoint rule, Trapezoidal rule and Gaussian rule methods are not difference. While, the processing time by the Midpoint rule method was the lowest among the four methods.

Keywords : Average Run Length, Cumulative Sum Control Chart, Numerical Integral Equations, Seasonal Autoregressive with Exogenous Variable

บทนำ

โดยปัจจุบัน โรงงานอุตสาหกรรมได้มีการนำเครื่องจักรมาใช้ในกระบวนการผลิต เพื่อผลิตสินค้าให้มีคุณภาพดีตรงตามมาตรฐานและตรงตามความต้องการของผู้บริโภค ซึ่งในกระบวนการผลิตมักจะเกิดความผันแปรเกิดขึ้นอยู่เสมอ การควบคุมคุณภาพกระบวนการผลิตทางสถิติ (Statistical Process Control: SPC) ถือเป็นสิ่งสำคัญสำหรับกระบวนการผลิต โดยเครื่องมือพื้นฐานของการควบคุมคุณภาพ มี 7 ชนิด ได้แก่ ใบตรวจสอบ (Check Sheet) แผนภาพพาเรโต (Pareto Diagram) แผนภาพก้างปลา (Fishbone Diagram) ฮิสโตแกรม (Histogram) แผนภาพการกระจาย (Scatter Diagram) กราฟชนิดต่างๆ (Graph) และแผนภูมิควบคุม (Control Chart) ซึ่งใช้เป็นเครื่องมือตรวจจับการเปลี่ยนแปลงของกระบวนการผลิตเพื่อแก้ไขปัญหาด้านคุณภาพได้อย่างรวดเร็ว จึงทำให้มีส่วนช่วยลดความผันแปรของกระบวนการผลิตทำให้ผลิตภัณฑ์มีคุณภาพสม่ำเสมอ ซึ่งแผนภูมิควบคุมเป็นหนึ่งในเครื่องมือที่สามารถนำไปประยุกต์ใช้กับงานในศาสตร์ต่างๆ มากมาย เช่น ด้านอุตสาหกรรม ด้านการแพทย์ ด้านสาธารณสุข ด้านการเงิน และด้านเศรษฐศาสตร์ เป็นต้น แผนภูมิควบคุมทำหน้าที่หลัก 3 ประการ คือ ประการแรกเพื่อกำหนดมาตรฐานในการผลิต ประการที่สอง เพื่อช่วยให้การผลิตบรรลุเป้าหมาย และประการสุดท้าย เพื่อปรับปรุงคุณภาพการผลิต โดยแผนภูมิควบคุมแบ่งออกเป็น 2 ประเภท คือ แผนภูมิควบคุมเชิงผันแปร (Control Chart for Variables) และแผนภูมิควบคุมเชิงลักษณะ (Control Chart for Attributes) โดยหลักการของแผนภูมิควบคุมจะจำกัดคุณภาพที่เหมาะสมของกระบวนการด้วยขีดจำกัดควบคุมบน (Upper Control Limit: UCL) ขีดจำกัดควบคุมล่าง (Lower Control Limit: LCL) และเส้นกลาง (Center Line: CL) ระยะห่างจากเส้นกลางถึงขีดจำกัดควบคุมบน และขีดจำกัดควบคุมล่างมีค่าเท่ากับ 3σ จึงเรียกแผนภูมิควบคุมนี้ว่า แผนภูมิควบคุม 3σ

ในปี ค.ศ. 1931 Walter A. Shewhart ได้เสนอแผนภูมิควบคุมชีวฮาร์ต (Shewhart Control Chart) ซึ่งเป็นแผนภูมิควบคุมที่มีประสิทธิภาพในการตรวจจับการเปลี่ยนแปลงค่าเฉลี่ยของกระบวนการผลิต กรณีที่ค่าเฉลี่ยมี

การเปลี่ยนแปลงขนาดใหญ่ (Large Shift) จากแนวคิดดังกล่าวนำไปสู่การพัฒนาแผนภูมิควบคุมที่มีประสิทธิภาพในการตรวจจับการเปลี่ยนแปลงของกระบวนการผลิตที่มีขนาดเล็ก (Small Shift) อาทิเช่น แผนภูมิควบคุมผลรวมสะสม (Cumulative Sum: CUSUM) และแผนภูมิควบคุมค่าเฉลี่ยเคลื่อนที่ถ่วงน้ำหนักเอ็กซ์โปเนนเชียล (Exponentially Weighted Moving Average: EWMA) เป็นต้น

ในปี ค.ศ. 1954 Page (1954) ได้เสนอแผนภูมิควบคุมผลรวมสะสม (Cumulative Sum: CUSUM) เพื่อใช้ในการตรวจจับการเปลี่ยนแปลงค่าเฉลี่ยของกระบวนการผลิต และเปรียบเทียบประสิทธิภาพแผนภูมิควบคุมชีวฮาร์ต โดยพิจารณาความยาวรันเฉลี่ย (Average Run Length: ARL) เป็นเกณฑ์ประเมินประสิทธิภาพของแผนภูมิควบคุม พบว่าแผนภูมิควบคุม CUSUM มีประสิทธิภาพตรวจจับการเปลี่ยนแปลงขนาดเล็กที่เกิดขึ้นในกระบวนการผลิตได้ดีกว่าแผนภูมิควบคุมชีวฮาร์ต เนื่องจากแผนภูมิควบคุม CUSUM ใช้ข้อมูลตลอดช่วงเวลาของการเก็บข้อมูลในกระบวนการมาวิเคราะห์แทนการใช้ข้อมูลในแต่ละคาบเวลา จึงส่งผลให้แผนภูมิควบคุม CUSUM สามารถส่งสัญญาณเตือนถึงการเปลี่ยนแปลงของค่าเฉลี่ยของกระบวนการได้อย่างทันทั่วถึง

โดยส่วนใหญ่ข้อมูลมักจะมีลักษณะการแจกแจงปกติ แต่ในการใช้งานบางสถานการณ์ที่ข้อมูลมีลักษณะที่เกิดอัตสหสัมพันธ์ (Autocorrelated) ที่มีรูปแบบอนุกรมเวลาที่แตกต่างกัน เช่น ข้อมูลมีลักษณะเกิดอัตสหสัมพันธ์ในรูปแบบของตัวแบบอัตถดถอย (Autoregressive: AR) ตัวแบบค่าเฉลี่ยเคลื่อนที่ (Moving Average: MA) ตัวแบบอัตถดถอยเคลื่อนที่ (Autoregressive and Moving Average: ARMA) และตัวแบบอัตถดถอยเคลื่อนที่แบบมีฤดูกาล (Seasonal Autoregressive Moving Average: SARMA) ดังนั้นการมีแผนภูมิควบคุมที่เหมาะสมกับข้อมูลที่เกิดอัตสหสัมพันธ์จะทำให้การใช้ประโยชน์ของแผนภูมิควบคุมมีประสิทธิภาพมากยิ่งขึ้น

ในปี ค.ศ. 2014 Phanyaem S. และคณะ ได้เสนอสูตรสำเร็จค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM กรณีที่ข้อมูลเกิดอัตถดถอยเคลื่อนที่ (Autoregressive and Moving Average: ARMA (1,1)) ผลการวิจัยพบว่า ค่าความยาวรันเฉลี่ยที่คำนวณได้จากสูตรสำเร็จมีค่าใกล้เคียงกับผลลัพธ์ที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลข ต่อมาปี ค.ศ. 2017 Phanyaem S. ได้ทำการเปรียบเทียบผลลัพธ์ที่ได้จากสูตรสำเร็จกับผลลัพธ์ที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลข (Integral Equation Approach) สำหรับแผนภูมิควบคุม CUSUM ในกรณีที่ข้อมูลเกิดอัตถดถอยเคลื่อนที่แบบมีฤดูกาล (Seasonal Autoregressive Moving Average: SARMA(1,1)_L) เมื่อค่าสังเกตมีการแจกแจงแบบเลขชี้กำลัง ผลการวิจัยพบว่า ค่าความยาวรันเฉลี่ยที่คำนวณจากสูตรสำเร็จให้ผลลัพธ์ที่ใกล้เคียงกับวิธีสมการปริพันธ์ โดยพิจารณาจากค่าสัมบูรณ์ของเปอร์เซ็นต์ความแตกต่าง (Absolute Percentage Difference) มีค่าต่ำกว่า 0.1% ต่อมาในปี ค.ศ. 2020 เพชรรัตน์ และคณะ (2563) ได้ประมาณค่าความยาวรันเฉลี่ยด้วยวิธีสมการปริพันธ์เชิงตัวเลข (Numerical Integral Equations) 4 วิธี ได้แก่ วิธีกฎซิมป์สัน (Simpson Rule) วิธีกฎค่ากลาง (Midpoint Rule) วิธีกฎสี่เหลี่ยมคางหมู (Trapezoidal Rule) และวิธีกฎเกาส์ (Gaussian Rule) ของแผนภูมิควบคุม EWMA เมื่อข้อมูลมีการแจกแจงเลขชี้กำลัง

การแจกแจงลาปลัส การแจกแจงโลจิสติกส์ และการแจกแจงพาร์โต โดยเปรียบเทียบผลลัพธ์กับวิธีการจำลองผลการวิจัยพบว่า เมื่อข้อมูลมีการแจกแจงเลขชี้กำลัง การแจกแจงลาปลัส และการแจกแจงโลจิสติกส์ วิธีกึ่งกลาง

มีค่าสัมบูรณ์ของเปอร์เซ็นต์ความแตกต่างน้อยที่สุด และเมื่อข้อมูลมีการแจกแจงพาร์โต วิธีกึ่งกลางมีค่าสัมบูรณ์ของเปอร์เซ็นต์ความแตกต่างน้อยกว่าวิธีอื่น โดยทั้งสองวิธีให้ค่าการประมาณที่ถูกต้องเทียบเท่ากับผลลัพธ์ที่ได้จากวิธีจำลองมอนติคาร์โลทุกกรณีศึกษา

ในการศึกษาครั้งนี้ผู้วิจัยสนใจศึกษากรณีที่ข้อมูลเป็นอนุกรมเวลาที่มีตัวแบบ SARX(P,1)_L เมื่อค่าความคลาดเคลื่อนมีการแจกแจงแบบเลขชี้กำลัง โดยใช้แผนภูมิควบคุม CUSUM ในการตรวจจับการเปลี่ยนแปลงของค่าเฉลี่ย เกณฑ์ที่ใช้ในการประเมินประสิทธิภาพของแผนภูมิควบคุมจะพิจารณาจากค่า ARL ซึ่งแบ่งออกเป็น 2 สถานะ คือ ค่าความยาวรันเฉลี่ยเมื่อกระบวนการอยู่ภายใต้การควบคุมใช้สัญลักษณ์ ARL_0 และค่าความยาวรันเฉลี่ยเมื่อกระบวนการออกนอกการควบคุมใช้สัญลักษณ์ ARL_1 ซึ่งวิธีการหาค่าความยาวรันเฉลี่ยมีหลายวิธี และวิธีที่นิยมใช้ทั่วไป คือ วิธีจำลองมอนติคาร์โล (Monte Carlo Simulations: MC) วิธีนี้มีข้อดี คือ ผลลัพธ์ของวิธีนี้สามารถนำมาตรวจสอบความถูกต้องและสามารถใช้เปรียบเทียบกับผลลัพธ์ของวิธีที่ต้องการศึกษาได้ แต่ข้อเสียของวิธีนี้ คือ ค่าตอบที่ได้นั้นต้องใช้เวลาเป็นอย่างมากในการประมวลผล ดังนั้นในการศึกษาครั้งนี้ผู้วิจัยจึงนำเสนอการประมาณค่าความยาวรันเฉลี่ยโดยวิธีสมการปริพันธ์เชิงตัวเลข (Numerical Integral Equations: NIE) 4 วิธี ได้แก่ วิธีกึ่งกลาง (Simpson Rule) วิธีกึ่งกลาง (Midpoint Rule) วิธีกึ่งกลาง (Trapezoidal Rule) และวิธีกึ่งกลาง (Gaussian Rule) โดยเปรียบเทียบค่าความยาวรันเฉลี่ยระหว่างวิธีสมการปริพันธ์เชิงตัวเลข และเวลาที่ใช้ในการประมวลผล (CPU Times) จากทั้ง 4 วิธี ในการศึกษาครั้งนี้ใช้โปรแกรม Mathematica ในการประมวลผล

แผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(P,1)_L

ส่วนนี้จะนำเสนอแผนภูมิควบคุม CUSUM ของตัวแบบ SARX(P,1)_L เมื่อค่าความคลาดเคลื่อนมีการแจกแจงแบบเลขชี้กำลัง และสมการปริพันธ์ของแผนภูมิควบคุม CUSUM

กำหนดให้ตัวแบบ SARX(P,1)_L มีรูปสมการ ดังนี้

$$Y_t = \beta X_t + \mu + \phi_1 Y_{t-1} + \dots + \phi_P Y_{t-P} + \varepsilon_t \quad (1)$$

โดยที่ X_t แทน ตัวแปรภายนอก ณ คาบเวลา t

Y_t แทน ค่าสังเกต ณ คาบเวลา t

β แทน ค่าพารามิเตอร์ของตัวแปรภายนอก

μ แทน ค่าพารามิเตอร์ของกระบวนการผลิต

ϕ_i แทน ค่าสัมประสิทธิ์ของการถดถอยในตัวแบบมีฤดูกาล เมื่อ $|\phi_i| < 1$; $i = 1, 2, \dots, P$

ε_t แทน ค่าความคลาดเคลื่อนที่มีการแจกแจงแบบเลขชี้กำลัง

กำหนดให้ค่าสถิติของแผนภูมิควบคุม CUSUM แสดงได้ดังนี้

$$C_t = \max(C_{t-1} + \xi_t - a, 0), \quad t = 1, 2, \dots \quad (2)$$

โดยที่ ξ_t แทน ลำดับของตัวแปรสุ่มที่เป็นอิสระต่อกันและมีการแจกแจงเหมือนกัน

C_0 แทน ค่าเริ่มต้นของค่าสถิติ CUSUM มีค่าเท่ากับ u

a แทน ค่าอ้างอิง (ค่าคงที่)

กำหนดให้ τ_b แทน คาบเวลาที่แผนภูมิควบคุมส่งสัญญาณเตือนว่ากระบวนการออกนอกการควบคุม

$$\tau_b = \inf \{ t > 0; C_t > b \}, \quad b > u, \quad (3)$$

โดยที่ b แทน ขีดจำกัดควบคุมบนของแผนภูมิควบคุม CUSUM

กำหนดให้สมการปริพันธ์ $H(u) = ARL = E_\infty(\tau_b) < \infty$ แทนค่า ARL ของแผนภูมิควบคุม CUSUM เมื่อ
กำหนดให้ $u \in [0, b]$ สมการปริพันธ์ของค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM ได้ดังนี้

$$H(u) = 1 + E_c \left[I \{ 0 < C_1 < b \} H(C_1) \right] + P_c \{ C_1 = 0 \} H(0) \quad (4)$$

สมการปริพันธ์ $H(u)$ ของแผนภูมิควบคุม CUSUM แสดงได้ดังนี้

$$H(u) = 1 + \alpha e^{\alpha(u-a+\mu+\beta X_t+\phi_1 y_{t-L}+\dots+\phi_P y_{t-PL})} \int_0^b H(y) e^{-\alpha y} dy + \left(1 - e^{-\alpha(a-u-\mu-\beta X_t-\phi_1 y_{t-L}-\dots-\phi_P y_{t-PL})} \right) H(0) \quad (5)$$

ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM โดยวิธีสมการปริพันธ์เชิงตัวเลข

ในงานวิจัยนี้ ผู้วิจัยสนใจที่จะทำการเปรียบเทียบค่าประมาณโดยใช้วิธีการประมาณ 4 วิธี ได้แก่ วิธีกฎซิมป์สัน (Simpson Rule) วิธีกฎค่ากลาง (Midpoint Rule) วิธีกฎสี่เหลี่ยมคางหมู (Trapezoidal Rule) และวิธีกฎเกาส์ (Gaussian Rule) วิธีสมการปริพันธ์เชิงตัวเลข (Numerical Integral Equation: NIE) ทั้ง 4 วิธี มีรายละเอียดดังนี้

1. วิธีกฎซิมป์สัน (Simpson Rule)

สำหรับกฎซิมป์สัน กำหนดให้ฟังก์ชัน $w(y) = 1$ และช่วงของการอินทิกรัล $[0, b]$ แบ่งเป็น $2m$ ช่วงย่อย $\{[y_{k-1}, y_k], k = 1, 2, \dots, 2m\}$ โดยมีความกว้างเท่ากับ $h = \frac{b-0}{2m}$ มีระยะห่างเท่ากัน $\{y_k = y_0 + kh, k = 1, 2, \dots, 2m\}$

ในงานวิจัยนี้กำหนดให้ $y_0 = 0$ และ $y_{2m} = b$ จากกฎพื้นฐานของซิมป์สันที่ใช้กับช่วงย่อย 2 ช่วงย่อย และกำหนดตัวถ่วงน้ำหนัก w_k คือ $y_{2k}, y_{2k+1}, y_{2k+2}$ ดังนั้นพหุนามองศาที่ $\{1, y, y^2\}$ ได้ถูกรวมโดยกฎของซิมป์สัน มีค่าน้ำหนักของจุดย่อยทั้ง 3 คือ $\{\frac{1}{3}h, \frac{4}{3}h, \frac{1}{3}h\}$ สำหรับกฎของคอมโพสิตซิมป์สัน เป็นการแบ่งช่วง $[0, b]$ ออกเป็นช่วงย่อยด้วยความกว้างเท่ากับ $2h$ ดังนั้นจากกฎของคอมโพสิตซิมป์สันสามารถเขียนได้ในรูป

$$S(f, h) = \frac{h}{3} (f(0) - f(b)) + \frac{2h}{3} \sum_{k=1}^m f(y_{2k}) + \frac{4h}{3} \sum_{k=1}^m f(y_{2k-1}); k = 0, 1, 2, \dots, m$$

วิธีการประมาณค่าปริพันธ์เชิงตัวเลขโดยวิธีกฎซิมป์สัน (Simpson Rule) อยู่ในรูป $\int_0^b f(y)dy \approx S(f, h)$

2. วิธีกฎค่ากลาง (Midpoint Rule)

กฎค่ากลางกำหนดให้ฟังก์ชัน $w(y)=1$ และ $f(y)$ บนช่วง $[0, b]$ แบ่งเป็นช่วงย่อย m ช่วง $\{[y_{k-1}, y_k], k=1, 2, \dots, m\}$ มีระยะห่างที่เท่ากัน $\{y_k = y_0 + kh, k=0, 1, 2, \dots, m\}$ ดังนั้น $y_0 = 0$ และ $y_m = b$ โดยช่วงมีความกว้าง $h = \frac{b-0}{m}$ จากกฎค่ากลาง (Midpoint Rule) สามารถเขียนได้ในรูป

$$M(f, h) = h \sum_{k=1}^m f\left(a + \left(k - \frac{1}{2}\right)h\right); k=0, 1, 2, \dots, m$$

วิธีการประมาณค่าปริพันธ์เชิงตัวเลขโดยวิธีกฎค่ากลาง (Midpoint Rule) อยู่ในรูป $\int_0^b f(y)dy \approx M(f, h)$

3. วิธีกฎสี่เหลี่ยมคางหมู (Trapezoidal Rule)

กฎสี่เหลี่ยมคางหมู กำหนดให้ฟังก์ชัน $w(y)=1$ และ $f(y)$ บนช่วง $[0, b]$ แบ่งเป็นช่วงย่อย m ช่วง $\{[y_{k-1}, y_k], k=1, 2, \dots, m\}$ โดยช่วงมีความกว้าง $h = \frac{b-0}{m}$ ระยะห่างที่เท่ากัน $\{y_k = y_0 + kh, k=0, 1, 2, \dots, m\}$ ดังนั้น $y_0 = 0$ และ $y_m = b$ ด้วยวิธีกฎสี่เหลี่ยมคางหมู (Trapezoidal Rule) สามารถเขียนได้ในรูป

$$T(f, h) = \frac{h}{2}(f(0) + f(b)) + h \sum_{k=1}^m f(y_k); k=0, 1, 2, \dots, m$$

วิธีการประมาณค่าปริพันธ์เชิงตัวเลขโดยวิธีกฎสี่เหลี่ยมคางหมู (Trapezoidal Rule) อยู่ในรูป

$$\int_0^b f(y)dy \approx T(f, h)$$

4. วิธีกฎเกาส์ (Gaussian Rule)

สำหรับกฎเกาส์ กำหนดให้ฟังก์ชัน $w(y)=1$ และ $f(y)$ บนช่วง $[0, b]$ แบ่งออกเป็นช่วงย่อย m ช่วง $\{y_k, k=1, 2, \dots, m\}$ กฎของเกาส์ควอดราเจอร์จะนิยามเซตของ $a_k = \frac{h}{m}\left(k - \frac{1}{2}\right); k=1, 2, \dots, m$ โดยที่ $0 \leq a_1 \leq \dots \leq a_m \leq b$ และเซตของค่าน้ำหนัก $w_k = b/m \geq 0; k=0, 1, 2, \dots, m$

วิธีการประมาณค่าปริพันธ์เชิงตัวเลขโดยวิธีกฎเกาส์ (Gaussian Rule) อยู่ในรูป $\int_0^b w(y)f(y)dy \approx \sum_{k=1}^m w_k f(a_k)$

กำหนดให้ $\tilde{H}(u)$ แทนสมการปริพันธ์เชิงตัวเลขของฟังก์ชัน $H(u)$ จากสมการปริพันธ์ของแผนภูมิควบคุม CUSUM ในสมการที่ (5) สามารถเขียนได้อีกรูปแบบหนึ่ง ดังนี้

$$\begin{aligned} \tilde{H}(a_i) &= 1 + H(0)F(a - a_i - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ &\quad + \int_0^b H(y)f(y + a - a_i - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL})dy \end{aligned} \quad (6)$$

เมื่อ $F(u) = 1 - e^{-au}$ และ $f(u) = \frac{dF(u)}{du} = \lambda e^{-\lambda u}$

จากสมการที่ (6) สามารถหาคำตอบได้โดยใช้วิธีการแก้ระบบสมการพีชคณิตเชิงเส้น ดังนี้

$$\begin{aligned}\tilde{H}(a_i) &= 1 + H(0)F(a - a_i - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ &\quad + \sum_{j=1}^m w_j \tilde{H}(a_j) f(a_j + a - a_1 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL})\end{aligned}\quad (7)$$

นั่นคือ

$$\begin{aligned}\tilde{H}(a_1) &= 1 + \tilde{H}(a_1)[F(a - a_1 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ &\quad + w_1 f(a - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL})] \\ &\quad + \sum_{j=2}^m w_j \tilde{H}(a_j) f(a_j + a - a_1 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ \tilde{H}(a_2) &= 1 + \tilde{H}(a_1)[F(a - a_2 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ &\quad + w_1 f(a_1 + a - a_2 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL})] \\ &\quad + \sum_{j=2}^m w_j \tilde{H}(a_j) f(a_j + a - a_2 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ &\quad \vdots \\ \tilde{H}(a_m) &= 1 + \tilde{H}(a_1)[F(a - a_m - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ &\quad + w_1 f(a_1 + a - a_m - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL})] \\ &\quad + \sum_{j=2}^m w_j \tilde{H}(a_j) f(a_j + a - a_m - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL})\end{aligned}$$

สามารถเขียนในรูปเมตริกซ์ R มิติ $m \times m$ ดังนี้

$$H_{m \times 1} = 1_{m \times 1} + R_{m \times m} H_{m \times 1} \quad (8)$$

$$\text{โดยที่ } H_{m \times 1} = \begin{pmatrix} \tilde{H}(a_1) \\ \tilde{H}(a_2) \\ \vdots \\ \tilde{H}(a_m) \end{pmatrix}, 1_{m \times 1} = \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix}$$

$$R = \begin{pmatrix} F(a - a_1 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) + w_1 f(a - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) & \dots & w_m f(a_m + a - a_1 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ F(a - a_2 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) + w_1 f(a_1 + a - a_2 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) & \dots & w_m f(a_m + a - a_2 - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \\ \vdots & & \vdots \\ F(a - a_m - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) + w_1 f(a_1 + a - a_m - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) & \dots & w_m f(a_m + a - a_m - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL}) \end{pmatrix}$$

โดยที่ $I_m = \text{diag}(1, 1, \dots, 1)$ ถ้า $(I_m - R_{m \times m})^{-1}$ มีอยู่จริง จะได้ว่า

$$H_{m \times 1} = (I_m - R_m)^{-1} 1_{m \times 1} \quad (9)$$

จากสมการที่ (8) สามารถแก้ปัญหาค่าเชิงตัวเลขเพื่อให้ได้คำตอบของเวกเตอร์ $\tilde{H}(a_i) = H_j$ จากนั้นแทนค่า $\tilde{H}(a_i)$ ลงในสมการที่ (7) และแทนค่า a_i ด้วย u จะได้สมการประมาณค่าสำหรับสมการปริพันธ์ $H(u)$ ดังนี้

$$\tilde{H}(u) = 1 + \tilde{H}(a_i)F(a - u - \mu - \beta X_t - \phi_1 y_{t-L} - \dots - \phi_P y_{t-PL})$$

$$+ \sum_{j=1}^m w_j \tilde{H}(a_j) f(a_j + a - u - \mu - \beta X_t - \phi_1 Y_{t-L} - \dots - \phi_P Y_{t-PL}) \quad (10)$$

ผลการวิจัย

งานวิจัยนี้มีวัตถุประสงค์เพื่อเปรียบเทียบค่าความยาวรันเฉลี่ย (Average Run Length: ARL) และความเร็วในการประมวลผล (CPU Times) ที่ได้จากการวิธีสมการปริพันธ์เชิงตัวเลข (Numerical Integral Equations: NIE) 4 วิธี ได้แก่ วิธีกฎซิมป์สัน (Simpson Rule) วิธีกฎค่ากลาง (Midpoint Rule) วิธีกฎสี่เหลี่ยมคางหมู (Trapezoidal Rule) และวิธีกฎเกาส์ (Gaussian Rule) สำหรับตัวแบบฮัตตอดอยแบบมีฤดูกาลที่มีตัวแปรภายนอก (SARX(P,1)_L) เมื่อค่าความคลาดเคลื่อนมีการแจกแจงแบบเลขชี้กำลัง ในกรณีที่กระบวนการอยู่ภายใต้การควบคุม (in-control state) กำหนดค่าพารามิเตอร์ของการแจกแจงแบบเลขชี้กำลัง (α) มีค่าเท่ากับ 1 และกรณีที่กระบวนการไม่อยู่ภายใต้การควบคุม (out-of-control state) กำหนดค่าระดับการเปลี่ยนแปลงค่าเฉลี่ย (δ) มีค่าเท่ากับ 1.50, 1.60, 1.70, 1.80, 1.90, 2.00, 2.10, 2.20, 2.30, 2.40, 2.50 และ 3.00 ตามลำดับ และกำหนดให้การแบ่งช่วงของวิธีสมการปริพันธ์เชิงตัวเลข (m) มีค่าเท่ากับ 500

ตารางที่ 1 และตารางที่ 2 นำเสนอค่า ARL และเวลาในการประมวลผลที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลขของตัวแบบ SARX(1,1)₄ โดยกำหนดค่าสัมประสิทธิ์ของการถดถอยในตัวเองแบบมีฤดูกาล ϕ_1 มีค่าเท่ากับ 0.1 ค่าสัมประสิทธิ์ของตัวแปรภายนอก β มีค่าเท่ากับ 1.0 ค่าอ้างอิง a มีค่าเท่ากับ 2.5 และขีดจำกัดควบคุมบนของแผนภูมิควบคุม b มีค่าเท่ากับ 3.976 เมื่อกำหนด ค่า $ARL_0 = 370$ แสดงผลลัพธ์ค่า ARL_1 ในแต่ละระดับการเปลี่ยนแปลงค่าเฉลี่ยได้ดังตารางที่ 1 และเมื่อกำหนดค่า $ARL_0 = 500$ ค่าขีดจำกัดควบคุมบนของแผนภูมิควบคุม b มีค่าเท่ากับ 4.326 แสดงผลลัพธ์ดังตารางที่ 2 ตามลำดับ

ตารางที่ 3 และตารางที่ 4 นำเสนอค่า ARL และเวลาในการประมวลผลที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลขของตัวแบบ SARX(2,1)₄ โดยกำหนดค่าสัมประสิทธิ์ของการถดถอยในตัวเองแบบมีฤดูกาล ϕ_1 และ ϕ_2 มีค่าเท่ากับ 0.1 ตามลำดับ ค่าสัมประสิทธิ์ของตัวแปรภายนอก β มีค่าเท่ากับ 1.0 ค่าอ้างอิง a มีค่าเท่ากับ 2.5 และขีดจำกัดควบคุมบนของแผนภูมิควบคุม b มีค่าเท่ากับ 4.151 เมื่อกำหนด ค่า $ARL_0 = 370$ แสดงผลลัพธ์ค่า ARL_1 ในแต่ละระดับการเปลี่ยนแปลงค่าเฉลี่ยดังตารางที่ 3 และเมื่อกำหนดค่า $ARL_0 = 500$ ค่าขีดจำกัดควบคุมบนของแผนภูมิควบคุม b มีค่าเท่ากับ 4.515 แสดงผลลัพธ์ดังตารางที่ 4 ตามลำดับ

ตารางที่ 5 และตารางที่ 6 นำเสนอค่า ARL และเวลาในการประมวลผลที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลขของตัวแบบ SARX(3,1)₄ โดยกำหนดค่าสัมประสิทธิ์ของการถดถอยในตัวเองแบบมีฤดูกาล ϕ_1, ϕ_2 และ ϕ_3 มีค่าเท่ากับ 0.1 ตามลำดับ ค่าสัมประสิทธิ์ของตัวแปรภายนอก β มีค่าเท่ากับ 1.0 ค่าอ้างอิง a มีค่าเท่ากับ 2.5 และขีดจำกัดควบคุมบนของแผนภูมิควบคุม b มีค่าเท่ากับ 4.349 เมื่อกำหนด ค่า $ARL_0 = 370$ แสดงผลลัพธ์ค่า ARL_1 ในแต่ละระดับการเปลี่ยนแปลงค่าเฉลี่ยได้ดังตารางที่ 5 และเมื่อกำหนดค่า $ARL_0 = 500$ ค่าขีดจำกัดควบคุมบนของแผนภูมิควบคุม b มีค่าเท่ากับ 4.731 แสดงผลลัพธ์ดังตารางที่ 6 ตามลำดับ

นอกจากนี้ได้แสดงการเปรียบเทียบค่า ARL_1 ของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(1,1)₄ ตัวแบบ SARX(2,1)₄ และตัวแบบ SARX(3,1)₄ ตามลำดับ เมื่อกำหนดให้ค่า $ARL_0 = 370$ ดังภาพที่ 1 และนำเสนอผลการ

เปรียบเทียบเวลาที่ใช้ในการประมวลผล (CPU Time) ของแผนภูมิควบคุม CUSUM ของตัวแบบ SARX(1,1)₄ ตัวแบบ SARX(2,1)₄ และตัวแบบ SARX(3,1)₄ ตามลำดับ โดยวิธีสมการปริพันธ์เชิงตัวเลขทั้ง 4 วิธี เมื่อกำหนดค่า $ARL_0 = 370$ แสดงดังภาพที่ 2

ตารางที่ 1 ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(1,1)₄ เมื่อกำหนดให้ $ARL_0 = 370$, $a = 2.5$ และ $b = 3.976$

δ	Numerical Integration Equation			
	Midpoint Rule	Simpson Rule	Trapezoidal Rule	Gaussian Rule
0.00	370.000 (11.68)*	370.000 (106.03)	370.000 (13.11)	370.000 (13.46)
1.50	7.917 (12.66)	7.891 (89.56)	7.963 (12.49)	8.025 (12.95)
1.60	7.262 (12.05)	7.240 (106.48)	7.300 (12.55)	7.351 (13.35)
1.70	6.709 (11.85)	6.691 (101.86)	6.741 (12.97)	6.784 (13.16)
1.80	6.237 (11.87)	6.222 (102.87)	6.264 (13.05)	6.300 (13.72)
1.90	5.831 (11.96)	5.818 (96.43)	5.855 (14.45)	5.885 (13.72)
2.00	5.479 (11.66)	5.468 (89.07)	5.499 (13.19)	5.526 (13.24)
2.10	5.171 (11.65)	5.161 (86.16)	5.189 (12.86)	5.211 (13.62)
2.20	4.899 (12.30)	4.891 (85.19)	4.915 (13.39)	4.935 (13.47)
2.30	4.659 (11.91)	4.652 (93.52)	4.673 (13.37)	4.691 (13.83)
2.40	4.448 (12.89)	4.438 (107.16)	4.457 (13.20)	4.473 (13.09)
2.50	4.253	4.247	4.264	4.278

	(12.57)	(106.71)	(13.44)	(13.91)
3.00	3.533	3.530	3.540	3.548
	(12.38)	(102.98)	(12.96)	(13.07)

* ตัวเลขที่อยู่ในเครื่องหมายวงเล็บ หมายถึง CPU Times (หน่วย: นาที)

ตารางที่ 2 ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(1,1)₄ เมื่อกำหนดให้ $ARL_0 = 500$,

$a = 2.5$ และ $b = 4.326$

δ	Numerical Integration Equation			
	Midpoint Rule	Simpson Rule	Trapezoidal Rule	Gaussian Rule
0.00	500.000 (11.77)*	500.000 (102.50)	500.000 (12.79)	500.000 (13.01)
1.50	8.541 (10.34)	8.508 (86.93)	8.607 (13.09)	8.698 (13.19)
1.60	7.809 (12.73)	7.782 (91.72)	7.864 (13.05)	7.938 (13.49)
1.70	7.194 (11.85)	7.172 (112.65)	7.239 (12.63)	7.301 (13.55)
1.80	6.672 (11.78)	6.653 (92.56)	6.710 (12.48)	6.762 (13.63)
1.90	6.224 (11.66)	6.208 (96.31)	6.256 (12.72)	6.300 (13.44)
2.00	5.836 (9.89)	5.823 (110.71)	5.864 (12.74)	5.902 (13.46)
2.10	5.498 (9.97)	5.486 (90.44)	5.522 (10.68)	5.555 (13.05)
2.20	5.201 (10.58)	5.191 (106.60)	5.222 (11.42)	5.250 (13.67)
2.30	4.939 (10.06)	4.929 (88.81)	4.957 (13.39)	4.981 (13.14)
2.40	4.705 (10.44)	4.698 (92.63)	4.721 (13.05)	4.728 (13.08)
2.50	4.496 (10.26)	4.486 (93.37)	4.510 (13.11)	4.529 (13.98)

3.00	3.716	3.711	3.724	3.736
	(10.33)	(84.84)	(11.49)	(12.54)

* ตัวเลขที่อยู่ในเครื่องหมายวงเล็บ หมายถึง CPU Times (หน่วย: นาที)

ตารางที่ 3 ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(2,1)₄ เมื่อกำหนดให้ $ARL_0 = 370$, $a = 2.5$ และ $b = 4.151$

δ	Numerical Integration Equation			
	Midpoint Rule	Simpson Rule	Trapezoidal Rule	Gaussian Rule
0.00	370.000 (12.51)*	370.000 (91.36)	370.000 (13.69)	370.000 (13.09)
1.50	7.714 (12.72)	7.688 (225.78)	7.768 (13.87)	7.841 (13.52)
1.60	7.087 (13.07)	7.065 (222.78)	7.134 (13.64)	7.192 (14.36)
1.70	6.557 (13.02)	6.539 (223.05)	6.595 (13.93)	6.645 (14.04)
1.80	6.105 (13.11)	6.089 (221.70)	6.137 (13.82)	6.179 (14.31)
1.90	5.716 (13.02)	5.703 (186.91)	5.743 (13.56)	5.779 (14.45)
2.00	5.378 (13.17)	5.366 (187.22)	5.401 (14.11)	5.432 (13.59)
2.10	5.082 (13.37)	5.072 (187.87)	5.102 (13.69)	5.129 (14.58)
2.20	4.821 (13.14)	4.812 (201.85)	4.839 (15.30)	4.862 (14.98)
2.30	4.590 (13.30)	4.582 (224.12)	4.605 (14.30)	4.626 (14.66)
2.40	4.383 (12.97)	4.376 (194.83)	4.397 (14.20)	4.415 (14.56)
2.50	4.198 (12.55)	4.191 (193.65)	4.209 (13.70)	4.226 (14.21)
3.00	3.501	3.498	3.508	3.518

(12.30)	(192.95)	(13.30)	(14.28)
---------	----------	---------	---------

* ตัวเลขที่อยู่ในเครื่องหมายวงเล็บ หมายถึง CPU Times (หน่วย: นาที)

ตารางที่ 4 ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(2,1)₄ เมื่อกำหนดให้ $ARL_0 = 500$, $a = 2.5$ และ $b = 4.515$

δ	Numerical Integration Equation			
	Midpoint Rule	Simpson Rule	Trapezoidal Rule	Gaussian Rule
0.00	500.000 (12.82)*	500.000 (116.93)	500.000 (14.67)	500.000 (14.37)
1.50	8.271 (10.90)	8.238 (101.13)	8.348 (14.01)	8.456 (14.47)
1.60	7.579 (12.33)	7.551 (98.48)	7.641 (14.13)	7.729 (14.57)
1.70	6.996 (11.93)	6.973 (115.16)	7.047 (14.10)	7.121 (15.49)
1.80	6.500 (12.23)	6.481 (116.28)	6.543 (13.84)	6.605 (14.91)
1.90	6.074 (11.36)	6.058 (95.36)	6.112 (14.46)	6.163 (14.22)
2.00	5.706 (10.97)	5.691 (109.17)	5.737 (14.10)	5.781 (13.66)
2.10	5.386 (12.32)	5.371 (101.07)	5.411 (14.22)	5.449 (13.67)
2.20	5.100 (13.18)	5.089 (96.67)	5.124 (13.94)	5.157 (14.24)
2.30	4.850 (12.84)	4.840 (92.15)	4.871 (13.72)	4.899 (14.97)
2.40	4.626 (13.29)	4.618 (95.58)	4.644 (13.81)	4.641 (14.41)
2.50	4.426 (13.03)	4.418 (97.10)	4.442 (13.68)	4.464 (14.27)
3.00	3.676	3.671	3.686	3.698

(12.75)	(116.36)	(13.58)	(13.70)
---------	----------	---------	---------

* ตัวเลขที่อยู่ในเครื่องหมายวงเล็บ หมายถึง CPU Times (หน่วย: นาที)

ตารางที่ 5 ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(3,1)₄ เมื่อกำหนดให้ $ARL_0 = 370$, $a = 2.5$ และ $b = 4.349$

δ	Numerical Integration Equation			
	Midpoint Rule	Simpson Rule	Trapezoidal Rule	Gaussian Rule
0.00	370.000 (12.41)*	370.000 (127.54)	370.000 (14.00)	370.000 (13.36)
1.50	7.485 (13.61)	7.458 (111.12)	7.548 (14.87)	7.636 (15.21)
1.60	6.891 (12.84)	6.869 (108.04)	6.943 (14.94)	7.016 (15.44)
1.70	6.389 (13.90)	6.371 (119.34)	6.432 (15.11)	6.493 (15.54)
1.80	5.960 (12.78)	5.944 (114.71)	5.997 (14.78)	6.048 (15.04)
1.90	5.590 (13.14)	5.577 (119.63)	5.621 (15.02)	5.665 (15.34)
2.00	5.269 (13.14)	5.257 (117.55)	5.295 (14.99)	5.332 (14.73)
2.10	4.986 (12.83)	4.976 (110.66)	5.010 (15.60)	5.042 (14.57)
2.20	4.737 (12.21)	4.728 (100.17)	4.758 (14.75)	4.786 (14.48)
2.30	4.516 (12.39)	4.508 (118.67)	4.534 (14.91)	4.559 (14.91)
2.40	4.318 (12.79)	4.311 (124.63)	4.334 (15.21)	4.355 (14.95)
2.50	4.141 (12.77)	4.134 (120.35)	4.155 (14.55)	4.174 (14.84)
3.00	3.471	3.467	3.459	3.491

(12.43)	(122.52)	(14.95)	(14.29)
---------	----------	---------	---------

* ตัวเลขที่อยู่ในเครื่องหมายวงเล็บ หมายถึง CPU Times (หน่วย: นาที)

ตารางที่ 6 ค่าความยาวรันเฉลี่ยของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(3,1)₄ เมื่อกำหนดให้ $ARL_0 = 500$, $a = 2.5$ และ $b = 4.731$

δ	Numerical Integration Equation			
	Midpoint Rule	Simpson Rule	Trapezoidal Rule	Gaussian Rule
0.00	500.000 (13.09)	500.000 (121.94)	500.000 (15.04)	500.000 (15.94)*
1.50	7.951 (12.09)	7.921 (116.53)	8.044 (13.55)	8.175 (14.94)
1.60	7.309 (13.08)	7.281 (119.09)	7.382 (13.02)	7.488 (14.30)
1.70	6.765 (13.33)	6.742 (100.20)	6.826 (13.26)	6.914 (14.55)
1.80	6.302 (12.91)	6.283 (111.15)	6.353 (13.87)	6.426 (15.11)
1.90	5.904 (12.23)	5.887 (121.85)	5.947 (13.75)	6.009 (15.12)
2.00	5.558 (12.99)	5.543 (124.02)	5.595 (15.32)	5.647 (15.06)
2.10	5.255 (12.09)	5.242 (122.54)	5.287 (14.44)	5.332 (15.07)
2.20	4.988 (12.46)	4.977 (103.13)	5.016 (14.71)	5.055 (14.60)
2.30	4.751 (12.56)	4.741 (118.11)	4.776 (14.03)	4.810 (15.36)
2.40	4.540 (12.22)	4.531 (114.60)	4.561 (14.08)	4.591 (15.38)
2.50	4.350 (12.18)	4.342 (116.21)	4.369 (14.12)	4.396 (15.64)
3.00	3.636	3.631	3.647	3.662

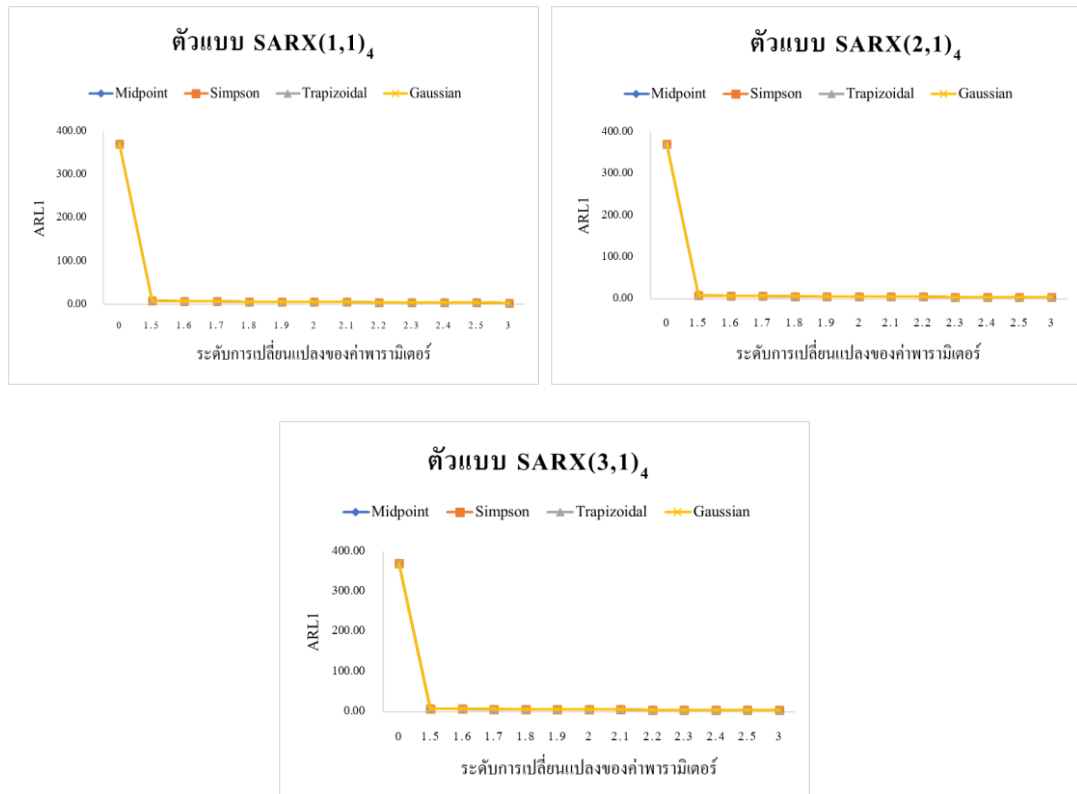
(12.84)

(108.27)

(13.84)

(15.88)

* ตัวเลขที่อยู่ในเครื่องหมายวงเล็บ หมายถึง CPU Times (หน่วย: นาที)



ภาพที่ 1 การเปรียบเทียบค่า ARL_1 ของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ $SARX(P,1)_4$ โดย
วิธีสมการปริพันธ์เชิงตัวเลข เมื่อกำหนดให้ $ARL_0=370$

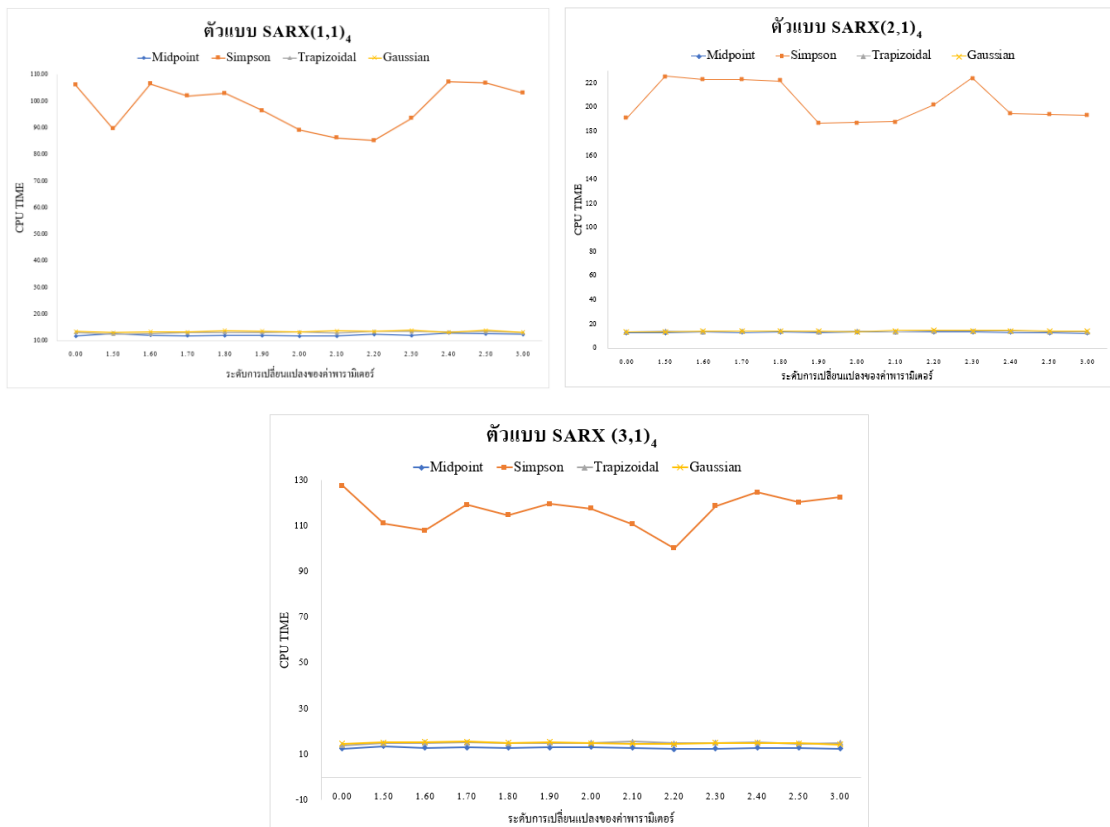
อภิปรายผลการวิจัย

จากตารางที่ 1 และตารางที่ 2 แสดงผลลัพธ์ค่า $ARL_0 = 370$ และ 500 ที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลขทั้ง 4 วิธีของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ $SARX(1,1)_4$ ผลการวิจัยพบว่า ผลลัพธ์ที่ได้จากวิธีปริพันธ์เชิงตัวเลขทั้ง 4 วิธี ให้ค่า ARL ไม่แตกต่างกันและเมื่อระดับการเปลี่ยนแปลงของค่าเฉลี่ย (δ) เพิ่มมากขึ้น ค่า ARL จะลดลงเรื่อยๆ และเมื่อพิจารณาจากเวลาที่ใช้ในการประมวลผล (CPU Times) พบว่า โดยส่วนใหญ่วิธีกฎค่ากลาง (Midpoint Rule) ใช้เวลาในการประมวลผลน้อยที่สุด

ตารางที่ 3 และตารางที่ 4 แสดงผลลัพธ์ค่า $ARL_0 = 370$ และ 500 ที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลขทั้ง 4 วิธีของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ $SARX(2,1)_4$ ผลการวิจัยพบว่า ผลลัพธ์ที่ได้จากวิธีปริพันธ์เชิง

ตัวเลขทั้ง 4 วิธี ให้ค่า ARL ไม่แตกต่างกัน และเมื่อระดับการเปลี่ยนแปลงของค่าเฉลี่ย (δ) เพิ่มขึ้น ค่า ARL จะลดลงเรื่อยๆ สำหรับการเปรียบเทียบเวลาที่ใช้ในการประมวลผลพบว่า โดยส่วนใหญ่วิธีกฎค่ากลาง (Midpoint Rule) ใช้เวลาในการประมวลผลน้อยที่สุด

ตารางที่ 5 และตารางที่ 6 แสดงผลลัพธ์ค่า $ARL_0 = 370$ และ 500 ที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลขทั้ง 4 วิธีของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(3,1)₄ ผลการวิจัยพบว่า ผลลัพธ์ที่ได้จากวิธีปริพันธ์เชิงตัวเลขทั้ง 4 วิธี ให้ค่า ARL ไม่แตกต่างกัน และเมื่อระดับการเปลี่ยนแปลงของค่าเฉลี่ย (δ) เพิ่มขึ้น ค่า ARL จะลดลงเรื่อยๆ สำหรับการเปรียบเทียบเวลาที่ใช้ในการประมวลผลพบว่า วิธีกฎค่ากลาง (Midpoint Rule) ใช้เวลาในการประมวลผลน้อยที่สุด



ภาพที่ 2 การเปรียบเทียบเวลาที่ใช้ในการประมวลผล (CPU Time) ของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ SARX(P,1)₄ โดยวิธีสมการปริพันธ์เชิงตัวเลข เมื่อกำหนดให้ $ARL_0=370$

จากภาพที่ 1 พบว่า ค่า ARL_1 ของแผนภูมิควบคุม CUSUM สำหรับตัวแบบ $SARX(1,1)_4$, $SARX(2,1)_4$ และ $SARX(3,1)_4$ โดยวิธีสมการปริพันธ์เชิงตัวเลขทั้ง 4 วิธี มีค่าไม่แตกต่างกัน และผลลัพธ์จากภาพที่ 2 แสดงให้เห็นว่าการประมาณค่าความยาวรันเฉลี่ยโดยวิธีกึ่งกลาง (Midpoint Rule) ใช้เวลาในการประมวลผล (CPU Times) น้อยที่สุดในบรรดาวิธีการประมาณค่าทั้ง 4 วิธี

สรุปผลการวิจัย

งานวิจัยนี้มีวัตถุประสงค์เพื่อประมาณค่าความยาวรันเฉลี่ยของแผนภูมิควบคุมผลรวมสะสม (Cumulative Sum: CUSUM) สำหรับตัวแบบอัตโนมัติแบบมีฤดูกาลที่มีตัวแปรภายนอก ($SARX(P,1)_L$) เมื่อค่าความคลาดเคลื่อนมีการแจกแจงแบบเลขชี้กำลัง โดยวิธีสมการปริพันธ์เชิงตัวเลข (Numerical Integration Equation: NIE) 4 วิธี ได้แก่ วิธีกฎซิมป์สัน (Simpson Rule) วิธีกึ่งกลาง (Midpoint Rule) วิธีกึ่งสี่เหลี่ยมคางหมู (Trapezoidal Rule) และวิธีกฎเกาส์ (Gaussian Rule) เนื่องจากเป็นวิธีที่ใช้คำนวณหาความยาวรันเฉลี่ย ที่ใช้เวลาในการประมวลผลน้อยเมื่อเทียบกับวิธีการหาค่าตอบโดยวิธีจำลองมอนติคาร์โล (Monte Carlo Simulations: MC) โดยผลการวิจัยพบว่า ค่าความยาวรันเฉลี่ย (ARL) ที่ได้จากวิธีสมการปริพันธ์เชิงตัวเลขทั้ง 4 วิธี และในแต่ละระดับการเปลี่ยนแปลงพารามิเตอร์ ผลลัพธ์ที่ได้จากวิธีกฎซิมป์สัน (Simpson Rule), วิธีกึ่งกลาง (Midpoint Rule), วิธีกึ่งสี่เหลี่ยมคางหมู (Trapezoidal Rule) และวิธีกฎเกาส์ (Gaussian Rule) ให้ค่า ARL ไม่แตกต่างกัน ในขณะที่เวลาที่ใช้ในการประมวลผล (CPU Times) ของวิธีกึ่งกลาง (Midpoint Rule) มีค่าน้อยที่สุด

เอกสารอ้างอิง

Page, E. S. (1954). Continuous Inspection Schemes. *Biometrika*. 41, 100-114.

Phanyaem, S. (2012). Analytical Expression and Numerical Solutions of Average Run Length for $SARMA(P,Q)_L$ Process on CUSUM procedure. *International Journal of Scientific & Technology research*, 1.

Phanyaem, S., Areepong, Y. & Sukparungsee, S. (2014). Explicit Formulas of Average Run Length for $ARMA(1,1)$ Process of CUSUM Control Chart. *Far East Journal of Mathematical Sciences*. 90(2), 211-224.

Phanyaem, S. (2017). Average Run Length of Cumulative Sum Control Charts for $SARMA(1,1)_L$ Models. *Thailand Statistician*, 15(2), 184-195.

Shewhart, W. A. (1931). *Economic control of quality of manufactured product* (1st ed.). Macmillan and Co. Ltd., London.

Poonpipat, P., Phanyaem, S., & Areepong, Y. (2020). Average run length values using the numerical integral equation method for exponential moving average control charts when the data is symmetrically distributed and is not symmetrical. *13th National Conference*, Srinakharinwirot University.

Research Article

ตัวแบบพยากรณ์ปริมาณการส่งออกถุงยางอนามัยของประเทศไทย ช่วงสถานการณ์ โควิด-19

Forecasting model for Export Condom Quantity of Thailand on the COVID-19 situation

ปรารถนา มินแสน^{1*}

Pradthana Minsan^{1*}

¹ภาควิชาคณิตศาสตร์และสถิติ คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยราชภัฏเชียงใหม่ เชียงใหม่ ประเทศไทย

¹Department of Mathematics and Statistics Faculty of Science and Technology, Chiang Mai Rajabhat University, Chiang Mai
Thailand

*E-mail: pradthana_min@g.cmru.ac.th

Received: 27/05/2021; Revised: 23/09/2021; Accepted: 23/09/2021

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อศึกษาการเคลื่อนไหว และสร้างตัวแบบพยากรณ์ที่เหมาะสมกับอนุกรมเวลา ปริมาณการส่งออกถุงยางอนามัยรายเดือนของประเทศไทยช่วงสถานการณ์โควิด-19 ตั้งแต่เดือนมกราคม พ.ศ. 2558 ถึงเดือนมีนาคม พ.ศ. 2564 จำนวน 75 ค่า โดยได้แบ่งข้อมูลออกเป็น 2 ชุด คือ ชุดข้อมูลฝึกฝน ข้อมูลตั้งแต่เดือน มกราคม พ.ศ. 2558 ถึงเดือนธันวาคม พ.ศ. 2562 จำนวน 60 ค่า สำหรับการสร้างตัวแบบพยากรณ์ด้วยวิธีการทาง สถิติ 4 วิธี ได้แก่วิธีการแยกส่วนประกอบ วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์ วิธีบอซซ์-เจนกินส์ และวิธีการพยากรณ์รวม และชุดข้อมูลทดสอบได้แก่ข้อมูลตั้งแต่เดือนมกราคม พ.ศ. 2563 ถึงเดือนมีนาคม พ.ศ. 2564 จำนวน 15 ค่า สำหรับเปรียบเทียบความคลาดเคลื่อนของวิธีการพยากรณ์ทั้ง 4 วิธี โดยใช้เกณฑ์รากของค่า คลาดเคลื่อนกำลังสองเฉลี่ย (Root Mean Square Error: RMSE) และ ร้อยละความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (Mean Absolute Percentage Error: MAPE) ผลการวิจัยพบว่า การพยากรณ์ด้วยวิธีการพยากรณ์รวมถ่วงน้ำหนักด้วย สัมประสิทธิ์การถดถอย เป็นตัวแบบที่เหมาะสมกับอนุกรมเวลาชุดนี้มากที่สุด เนื่องจากให้ค่า RMSE และ MAPE ต่ำ ที่สุด ซึ่งมีตัวแบบสมการพยากรณ์ $\hat{Y}_t = -10,200 + 0.584\hat{Y}_{1t} + 0.597\hat{Y}_{2t} - 0.198\hat{Y}_{3t}$ เมื่อ $\hat{Y}_{1t}, \hat{Y}_{2t}$, และ $\hat{Y}_{3t}$ แทนค่าพยากรณ์เดี่ยวในเวลา t จากวิธีการวิธีการแยกส่วนประกอบ วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของ โฮลต์และวินเตอร์ และวิธีบอซซ์-เจนกินส์ ตามลำดับ

คำสำคัญ: การแยกส่วนประกอบ, การปรับเรียบด้วยเส้นโค้งเลขชี้กำลัง, บอซ-เจนกินส์, การพยากรณ์รวม, รากของค่าคลาดเคลื่อนกำลังสองเฉลี่ย, ร้อยละความคลาดเคลื่อนสัมบูรณ์เฉลี่ย

Abstract

The objectives of this research are to study the movement and construct the most suitable forecasting model for the monthly condoms export volumes of Thailand on the COVID-19 situation. The data in this study gathered from the Thai Customs that recorded in monthly during January 2015 to March 2021 a total of 75 months. Then the data were classified into two sets. The training data set, January 2015 to December 2019, for 60 months were used to construct model by Decomposition, Holt-Winters Exponential Smoothing, Box-Jenkins and Combined method. The test data set, January 2020 to March 2021, a total of 15 months were used to compare forecasting accuracy model the earlier three models via the criteria of Root Mean Square Error (RMSE) and Mean Absolute Percentage Error (MAPE). Research results indicated that the combined forecasting method where the weighted are using a regression analysis method has given the lowest RMSE and MAPE, therefore this method is the most suitable method of monthly export condoms quantity. The forecasting model is:

$\hat{Y}C_t = -10,200 + 0.584\hat{Y}_{1t} + 0.597\hat{Y}_{2t} - 0.198\hat{Y}_{3t}$, where $\hat{Y}_{1t}$, $\hat{Y}_{2t}$, and $\hat{Y}_{3t}$ represent the single forecast at time t from Decomposition, Holt-Winters Exponential Smoothing and Box-Jenkins method, respectively.

Keywords: Decomposition Holt-Winters, Exponential Smoothing, Box-Jenkins, Root Mean Square, Error Mean, Absolute Percentage Error

บทนำ

ประเทศไทย มาเลเซีย และอินโดนีเซีย เป็นประเทศผู้ผลิตยางพาราธรรมชาติมากกว่าร้อยละ 70 ของโลก แนวโน้มความต้องการยางอนามัยขยายตัวเพิ่มขึ้นในช่วงระบาดของไวรัสโคโรนาสายพันธุ์ใหม่ (โควิด-19) โดยมีการคาดการณ์ว่าในปี พ.ศ. 2564 มูลค่าตลาดยางอนามัยจะมีมูลค่า 9,410 ล้านดอลลาร์สหรัฐ ขยายตัวเพิ่มขึ้นจากปี พ.ศ. 2562 ถึงร้อยละ 18.07 (Department of International Trade Promotion, 2020) อย่างไรก็ตามจากปัญหาการแพร่ระบาดของโควิด-19 ทั่วโลก ส่งผลให้บางประเทศเลือกใช้มาตรการล็อกดาวน์ (ปิดเมือง) เพื่อควบคุมสถานการณ์ และมีผลกระทบต่อกำลังการผลิตสินค้าในอุตสาหกรรมบางประเภทลดลง โดยในส่วนของอุตสาหกรรมการผลิตยางอนามัยนั้น มีโรงงานของผู้ผลิตรายใหญ่ในต่างประเทศบางแห่งที่ไม่สามารถเดินเครื่องจักรได้อย่างเต็มที่ ทำให้มีความเสี่ยงเกิดภาวะขาดแคลนสินค้า สำหรับประเทศไทย ยางอนามัยถูกจัดเป็นกลุ่มของเครื่องมือทางการแพทย์ ทำให้ถ้ามีการล็อกดาวน์ จะจัดเป็นกลุ่มสุดท้ายที่ต้องปิดโรงงาน จึงทำให้ประเทศไทยมีความเสี่ยงน้อย

ที่สุดในเรื่องการขาดแคลนถุงยางอนามัย และยังสามารถส่งไปตลาดโลกได้ (Phadungkarn, 2020) ซึ่งสาเหตุที่คาดว่าความต้องการถุงยางอนามัยเพิ่มขึ้นนั้น อาจสืบเนื่องมาจากภาวะเศรษฐกิจที่มีแนวโน้มหดตัวจากการว่างงานในช่วงที่เกิดวิกฤตการระบาดต่อเนื่องเป็นเวลานาน ทำให้ประชาชนส่วนใหญ่ต้องการมีบุตรน้อยลง ครอบครัวจึงเริ่มวางแผนคุมกำเนิดมากขึ้น (Department of International Trade Promotion, 2020)

จากความไม่แน่นอนในช่วงสถานการณ์ โควิด-19 ที่อาจจะส่งผลกระทบต่อความต้องการถุงยางอนามัยทำให้การพยากรณ์มีบทบาทสำคัญที่จะช่วยในการคาดการณ์เหตุการณ์ที่อาจจะเกิดขึ้น การพยากรณ์เชิงปริมาณที่ใช้ข้อมูลในอดีตที่จัดเก็บอย่างต่อเนื่องจึงเป็นเทคนิคการพยากรณ์ทางสถิติที่นิยมถูกนำมาใช้ วิธีการแยกส่วนประกอบ วิธีการปรับเรียบ วิธีบอซซ์-เจนกินส์ และวิธีการพยากรณ์รวม เป็นวิธีการพยากรณ์ที่สามารถใช้พยากรณ์ได้ครอบคลุมข้อมูลอนุกรมเวลาหลากหลายลักษณะและหลากหลายส่วนประกอบของอนุกรมเวลา จึงมีนักวิจัยมากมายนำเทคนิคเหล่านี้มาใช้ในการวิจัยเพื่อการพยากรณ์ เช่น Yonar et al. (2020) ได้มีการวิจัยโดยประยุกต์ใช้เทคนิคการแยกส่วนประกอบเมื่อข้อมูลมีส่วนประกอบของแนวโน้มโดยตัวแบบการประมาณค่าเส้นโค้ง เทคนิคการปรับเรียบด้วยโดยใช้การปรับเรียบเส้นโค้งเลขชี้กำลังของบราวน์ และการปรับเรียบเส้นโค้งเลขชี้กำลังของโฮลต์ รวมทั้งวิธีบอซซ์-เจนกินส์ ในการสร้างและเปรียบเทียบตัวแบบพยากรณ์จำนวนของผู้ติดเชื้อไวรัสโควิด -19 ในประเทศตุรกีและบางประเทศในกลุ่ม G8 Riansut (2020) ได้สร้างตัวแบบพยากรณ์เทคนิควิธีการปรับเรียบ 3 วิธี และวิธีบอซซ์-เจนกินส์ ในการพยากรณ์อัตราเร็วลมรายวัน และ Keerativibool (2013) ได้สร้างตัวแบบพยากรณ์เทคนิคการปรับเรียบ วิธีบอซซ์-เจนกินส์ และการวิธีการพยากรณ์รวม ในการพยากรณ์อุณหภูมิเฉลี่ยต่อเดือนในเขตกรุงเทพมหานคร เป็นต้น ดังนั้นผู้วิจัยจึงมีความสนใจที่จะนำเทคนิคการพยากรณ์ทางสถิติเป็นเครื่องมือช่วยการพยากรณ์การส่งออกของปริมาณถุงยางอนามัยของประเทศไทยในอนาคต เพื่อเป็นประโยชน์ต่อการการบริหารจัดการทั้งภาครัฐและเอกชนรวมทั้งผู้ที่เกี่ยวข้อง โดยในการวิจัยครั้งนี้ ผู้วิจัยได้สนใจศึกษาวิธีการสร้างตัวแบบพยากรณ์ทั้งหมด 4 วิธี ได้แก่ วิธีการแยกส่วนประกอบ วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์ วิธีบอซซ์-เจนกินส์ และวิธีการพยากรณ์รวม เพื่อให้ครอบคลุมตัวแบบพยากรณ์ที่เหมาะสมกับลักษณะข้อมูลปริมาณการส่งออกถุงยางอนามัยของประเทศไทย

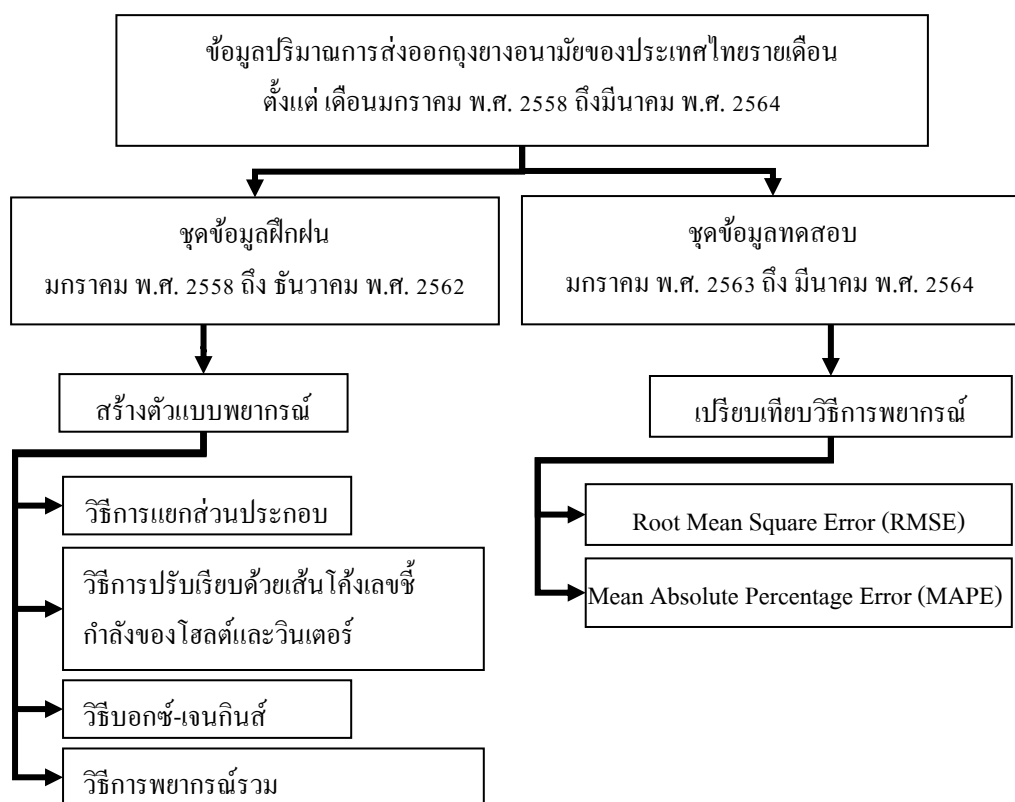
วิธีดำเนินการวิจัย

การศึกษาวิจัยครั้งนี้มีขั้นตอนวิธีการดำเนินการวิจัยดังนี้

1. การจัดเตรียมข้อมูล

การวิจัยครั้งนี้ได้ใช้ข้อมูลปริมาณการส่งออกถุงยางอนามัยรายเดือนของประเทศไทย ช่วงสถานการณ์โควิด-19 จากกรมศุลกากร ตั้งแต่เดือนมกราคม พ.ศ. 2558 ถึงเดือนมีนาคม พ.ศ. 2564 จำนวน 75 ค่า (Thai Customs, 2020) โดยจัดแบ่งข้อมูลดังกล่าวออกเป็น 2 ชุด ชุดที่ 1 เป็นชุดข้อมูลฝึกฝน (Training Data Set) ได้แก่ปริมาณการส่งออกถุงยางอนามัย ตั้งแต่เดือน มกราคม พ.ศ. 2558 ถึง เดือนธันวาคม พ.ศ. 2562 จำนวน 60 ค่า คิดเป็นร้อยละ 80

ของจำนวนข้อมูลทั้งหมด เพื่อใช้สำหรับหาตัวแบบที่เหมาะสมของแต่ละเทคนิคการพยากรณ์ ด้วยโปรแกรม Microsoft Excel (2020) และ โปรแกรม Minitab v.18 (2020) ส่วนข้อมูลชุดที่ 2 เป็นชุดข้อมูลทดสอบ (Test Data Set) ได้แก่ ปริมาณการส่งออกกุ้งยางอนามัยตั้งแต่เดือนมกราคม พ.ศ. 2563 ถึงเดือนมีนาคม พ.ศ. 2564 จำนวน 15 ค่า คิดเป็นร้อยละ 20 ของจำนวนข้อมูลทั้งหมด เพื่อใช้ตรวจสอบความแม่นยำ (Accuracy) ของตัวแบบพยากรณ์ที่สร้างขึ้น โดยใช้เกณฑ์รากของค่าคลาดเคลื่อนกำลังสองเฉลี่ย (Root Mean Square Error: RMSE) และ ร้อยละความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (Mean Absolute Percentage Error: MAPE) ที่ต่ำที่สุด ซึ่งได้แสดงกรอบแนวคิดในการพยากรณ์ปริมาณการส่งออกกุ้งยางอนามัยรายเดือนของประเทศไทย ดังรูปที่ 1



รูปที่ 1 กรอบแนวคิดการวิจัยตัวแบบพยากรณ์ปริมาณการส่งออกกุ้งยางอนามัยของประเทศไทย

2. การศึกษาความเคลื่อนไหวของอนุกรมเวลา

การศึกษาความเคลื่อนไหวของอนุกรมเวลาเป็นการพิจารณาเบื้องต้นว่าอนุกรมเวลานั้นๆ มีลักษณะเป็นแบบใด โดยพิจารณาจากกราฟ (t, Y_t) (Panichkitkosolkul, 2009)

3. การสร้างตัวแบบพยากรณ์

ในการวิจัยครั้งนี้ได้กำหนดสัญลักษณ์ไว้ดังนี้

Y_t และ $\hat{Y}_t$ แทนอนุกรมเวลา และค่าพยากรณ์ ณ เวลา t ตามลำดับ

ε_t แทนความคลาดเคลื่อนที่มีการแจกแจงปกติและเป็นอิสระกัน ที่มีเฉลี่ยเท่ากับศูนย์ และความแปรปรวนคงที่ทุกช่วงเวลา

t แทนช่วงเวลามีค่าตั้งแต่ 1 ถึง n เมื่อ n แทนจำนวนข้อมูลในอนุกรมเวลาชุดข้อมูลฝึกฝน

s แทนจำนวนคาบของฤดูกาล ($s = 12$)

3.1 วิธีการแยกส่วนประกอบ (Abraham & Ledolter, 1983) เป็นวิธีการที่แยกอนุกรมเวลาออกเป็น ส่วนประกอบต่างๆ คือ แนวโน้ม (Trend: T_t) ฤดูกาล (Seasonal: S_t) วงจร (Cycle: C_t) และส่วนประกอบไม่ปกติ (Irregular: I_t) โดยทั่วไปการพยากรณ์ในระยะสั้น ส่วนประกอบที่มีผลต่อการพยากรณ์จะมีเพียงแนวโน้ม และ ฤดูกาลเท่านั้น เนื่องจากอิทธิพลของวัฏจักรมีผลต่อการพยากรณ์ระยะสั้นน้อย และส่วนประกอบไม่ปกติก็ไม่สามารถที่จะคาดการณ์ได้ว่าเกิดขึ้นในช่วงเวลาใด (Minsan, 2014) วิธีแยกส่วนประกอบแบ่งเป็น 2 รูปแบบ ได้แก่ รูปแบบบวก (Additive Decomposition) เหมาะกับอนุกรมเวลาที่มีความผันแปรของฤดูกาลคงที่ และ รูปแบบคูณ (Multiplicative Decomposition) เหมาะกับอนุกรมเวลาที่มีความผันแปรของฤดูกาลไม่คงที่ สำหรับในการวิจัยครั้งนี้ ใช้วิธีแยกส่วนประกอบรูปแบบบวก เนื่องจากอนุกรมเวลาปริมาณการส่งออกยางอนามัยของชุดข้อมูลฝึกฝน ดัง ภาพ 2 พบว่ามีทั้งส่วนประกอบของแนวโน้มและความผันแปรของฤดูกาล โดยที่ความผันแปรของฤดูกาลค่อนข้าง คงที่ นั่นคือความผันแปรตามฤดูกาล มีได้เพียงพึ่งต่อแนวโน้มของข้อมูล (Lorchirachoonkul & Jitthavech, 2005) นอกจากนั้นในงานวิจัยนี้ได้พิจารณาจากการทดสอบความผันแปรของฤดูกาล โดยสถิติทดสอบ Levene's Test (Levene, 1960) (Brown & Forsythe, 1974) ควบคู่ไปด้วย ซึ่งจะได้อธิบายรายละเอียดในหัวข้อผลการศึกษาดังนั้นตัวแบบและตัวแบบพยากรณ์แสดงดังสมการ (1) และ (2) ตามลำดับ (Ket-iam, 2005)

$$Y_t = \beta_0 + \beta_1 t + S_t + \varepsilon_t \quad (1)$$

$$\hat{Y}_t = b_0 + b_1 t + \hat{S}_t \quad (2)$$

เมื่อพารามิเตอร์ β_0, β_1 แทนระดับของข้อมูล และความชัน ตามลำดับ S_t แทนความผันแปรของฤดูกาล ณ เวลา t ส่วน b_0, b_1 และ $\hat{S}_t$ เป็นตัวประมาณของ β_0, β_1 และ S_t ตามลำดับ

3.2 วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์ เทคนิคการสร้างสมการพยากรณ์ด้วยการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์เหมาะสำหรับอนุกรมเวลาที่มีการเคลื่อนไหวเนื่องจาก แนวโน้มและฤดูกาล ทั้งรูปแบบแนวโน้มฤดูกาลแบบบวกและแบบคูณ วิธีของโฮลต์และวินเตอร์ ใช้ค่าทำให้เรียบ 3 ค่าได้แก่ α, γ และ δ ซึ่งต่างมีค่าอยู่ระหว่าง 0 ถึง 1 สำหรับค่าแนวโน้ม ค่าความลาดชัน และค่าวัดอิทธิพลของ ฤดูกาลหรือดัชนีฤดูกาลตามลำดับ (Bowman & O'Connell, 1993) ซึ่งในงานวิจัยครั้งนี้ได้เลือกใช้รูปแบบแนวโน้ม

ฤดูกาลแบบบวก เนื่องจากเหตุผลเดียวกันที่ได้กล่าวไปแล้วในหัวข้อการพยากรณ์โดยวิธีแยกส่วนประกอบ ซึ่งมีรูปแบบดังสมการที่ (1) และตัวแบบพยากรณ์ ดังสมการที่ (3) (Taesombat, 2006)

$$\hat{Y}_{t+k} = \hat{T}_{t+k} + \hat{S}_{t+k-s} \quad \text{สำหรับ } k = 1, 2, \dots \quad (3)$$

เมื่อ $\hat{T}_{t+k} = \hat{T}_t + k\hat{\beta}_t$
 $\hat{Y}_{t+k}$ แทนค่าพยากรณ์ ณ เวลา $t+k$ โดยที่ k แทนจำนวนช่วงเวลาที่ต้องการพยากรณ์ไปล่วงหน้า
 $\hat{T}_t, \hat{\beta}_t$ และ $\hat{S}_t$ แทนค่าประมาณพารามิเตอร์ β_0, β_1 และ S_t ณ เวลา t ตามลำดับ โดยสามารถคำนวณได้ดังสมการที่ (4) (5) และ (6) (Tirkeş et al., 2017),

$$\hat{T}_t = \alpha(Y_t - \hat{S}_{t-s}) + (1-\alpha)(\hat{T}_{t-1} + \hat{\beta}_{t-1}), \quad (4)$$

$$\hat{\beta}_t = \gamma(\hat{T}_t - \hat{T}_{t-1}) + (1-\gamma)\hat{\beta}_{t-1}, \quad (5)$$

$$\hat{S}_t = \delta(Y_t - \hat{T}_t) + (1-\delta)\hat{S}_{t-s}. \quad (6)$$

3.3 วิธีบอกซ์-เจนกินส์ (Box et al., 1994) เป็นวิธีที่ให้ค่าพยากรณ์ที่มีความถูกต้องสูงกว่าค่าพยากรณ์จากวิธีการอื่นในการพยากรณ์ระยะสั้น การกำหนดตัวแบบพยากรณ์ทำได้โดยการตรวจสอบฟังก์ชันสหสัมพันธ์ในตัวเอง (Autocorrelation Function: ACF) และฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน (Partial Autocorrelation Function: PACF) ของอนุกรมเวลาแบบคงที่ (Stationary Time Series) หรืออนุกรมเวลาที่มีค่าเฉลี่ยของความแปรปรวนคงที่ (Taesombat, 2006) ตัวแบบทั่วไปวิธีบอกซ์-เจนกินส์ คือ Seasonal Autoregressive Integrated Moving Average หรือ SARIMA(p,d,q)(P,D,Q)_s ดังสมการที่ (7)

$$\phi_p(B)\Phi_p(B)(1-B^d(1-B)^D)Y_t = \delta + \theta_q(B)\Theta_q(B)\varepsilon_t \quad (7)$$

เมื่อ $\phi_p(B) = 1 - \phi_1 B - \phi_2 B^2 - \dots - \phi_p B^p$ แทนการถดถอยในตัวเองไม่มีฤดูกาลอันดับ p (Non-seasonal Autoregressive Operator of Order p : AR(p))

$\Phi_p(B) = 1 - \Phi_1 B - \Phi_2 B^s - \dots - \Phi_p B^{ps}$ แทนการถดถอยในตัวเองมีฤดูกาลอันดับ p (Seasonal Autoregressive: SAR(P))

$(1-B^d)$ แทนอันดับของผลต่างแบบไม่มีฤดูกาล (Non-seasonal Difference)

$(1-B)^D$ แทนอันดับของผลต่างแบบมีฤดูกาล (Seasonal Difference)

$\delta = \phi_p(B)\Phi_p(B)\mu$ แทนค่าคงที่ โดยที่ μ แทนค่าเฉลี่ยของอนุกรมเวลาที่คงที่

$\theta_q(B) = 1 - \theta_1 B - \theta_2 B^2 - \dots - \theta_q B^q$ แทนค่าเฉลี่ยเคลื่อนที่แบบไม่มีฤดูกาลอันดับ q (Non-seasonal Moving Average: MA(q))

$\Theta_q(B) = 1 - \Theta_1 B - \Theta_2 B^s - \dots - \Theta_q B^{qs}$ แทนค่าเฉลี่ยเคลื่อนที่แบบมีฤดูกาลอันดับ Q (Seasonal Moving Average: SMA(Q))

B แทนตัวดำเนินการถอยหลัง (Backward Operator) โดยที่ $B^d Y_t = Y_{t-d}$

d และ D แทนลำดับที่ของการหาผลต่างและผลต่างฤดูกาลของอนุกรมเวลา ตามลำดับ การสร้างตัวแบบพยากรณ์โดยวิธี บอซ-เจนกินส์ มีขั้นตอนดังนี้

การตรวจสอบข้อมูล เพื่อพิจารณาว่าอนุกรมเวลาคงที่หรือไม่ โดยพิจารณาจากกราฟอนุกรมเวลา หรือพิจารณาจากกราฟ ACF และ PACF หากพบว่าอนุกรมเวลาไม่คงที่ (Non-stationary) ต้องแปลงอนุกรมเวลาให้คงที่ก่อนขั้นตอนต่อไป เช่นการแปลงข้อมูลด้วยการหาผลต่างหรือผลต่างฤดูกาล (Difference or Seasonal Difference)

1) การกำหนดตัวแบบพยากรณ์ที่เป็นไปได้จากกราฟ ACF และ PACF ของอนุกรมเวลาใหม่ ที่อนุกรมเวลาแบบคงที่ นั่นคือ กำหนดค่า p, q, P และ Q พร้อมทั้ง ประมาณค่าพารามิเตอร์ของตัวแบบด้วยวิธีกำลังสองน้อยที่สุด (Ordinary Least Squares Method)

2) การตรวจสอบความเหมาะสมของตัวแบบพยากรณ์ที่กำหนด โดยการพิจารณาคัดเลือกตัวแบบพยากรณ์ที่มีค่าเกณฑ์ข้อสนเทศของเบส์ (Bayesian Information Criteria: BIC) (Sawa, 1978) ที่ต่ำที่สุด มีค่าสถิติทดสอบ Ljung-Box Q ที่ไม่มีนัยสำคัญ (Ruppert, 2011) และ ตรวจสอบข้อสมมติเกี่ยวกับความคลาดเคลื่อนจากการพยากรณ์ (e_t) ที่ต้องมีคุณลักษณะความเป็นอิสระกันและมีการแจกแจงปกติด้วยค่าเฉลี่ยเท่ากับศูนย์และมีความแปรปรวนคงที่ นั่นคือ

- ความเป็นอิสระกัน พิจารณาจากกราฟ ACF และ PACF ของค่าความคลาดเคลื่อน
- การแจกแจงปกติ โดยสถิติทดสอบ Anderson-Darling (Anderson & Darling, 1954)
- ค่าเฉลี่ยเท่ากับศูนย์ โดยสถิติทดสอบที (t-test)
- ความแปรปรวนคงที่ โดยสถิติทดสอบ Bartlett (Glaser, 1976)

3) พยากรณ์อนุกรมเวลาโดยใช้ตัวแบบพยากรณ์ที่เหมาะสมที่สุดจากขั้นตอนที่ 3

3.4 การพยากรณ์โดยวิธีการพยากรณ์รวม (Combined Forecasting Method) เป็นวิธีการพยากรณ์ที่ได้จากนำค่าพยากรณ์เดี่ยว (Individual Forecast) จากแต่ละวิธีมาพิจารณาร่วมกัน โดยหาตัวถ่วงน้ำหนักที่เหมาะสมสำหรับแต่ละวิธี ซึ่งค่าพยากรณ์ใหม่ที่ได้นั้นจะให้ความคลาดเคลื่อนต่ำสุด (Sukparungsee, 2003) ซึ่งมีวิธีการพยากรณ์รวมที่เหมาะสมหลายวิธี (Manmin, 2006) ในงานวิจัยนี้จะพิจารณาจากค่าพยากรณ์เดี่ยวจาก 3 วิธี คือ วิธีการแยกส่วนประกอบ วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์ และวิธีบอซ-เจนกินส์ โดยมีวิธีการที่ใช้ในการถ่วงน้ำหนัก 3 วิธีดังนี้

1) วิธีการถ่วงน้ำหนักด้วยสัมประสิทธิ์การถดถอย (Regression Analysis Method: Reg) (Keerativibool, 2013)

$$\hat{Y}_{C_t} = \hat{\beta}_0 + \hat{\beta}_1 \hat{Y}_{1t} + \hat{\beta}_2 \hat{Y}_{2t} + \hat{\beta}_3 \hat{Y}_{3t} \quad (8)$$

เมื่อ $\hat{Y}_{C_t}$ แทนค่าพยากรณ์รวมในเวลา t

$\hat{Y}_{1t}$, $\hat{Y}_{2t}$ และ $\hat{Y}_{3t}$ แทนค่าพยากรณ์เดี่ยวในเวลา t จากวิธีการวิธีการแยกส่วนประกอบ วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์ และวิธีบอซ-เจนกินส์ ตามลำดับ

$\hat{\beta}_0$ แทนตัวประมาณค่าคงที่ของตัวแบบการวิเคราะห์ถดถอย

$\hat{\beta}_1, \hat{\beta}_2$ และ $\hat{\beta}_3$ แทนตัวประมาณค่าถ่วงน้ำหนักของแต่ละวิธีการพยากรณ์เดี่ยวด้วยวิธีกำลังสองน้อยที่สุด (Least Squares Method) (Montgomery et al., 2006)

2) วิธีการถ่วงน้ำหนักค่าสัมบูรณ์ต่ำสุด (Least Absolute Value: LAV) เป็นวิธีการคำนวณหาค่าเฉลี่ยถ่วงน้ำหนักไดนามิกเทคนิคการโปรแกรมเชิงเส้นตรง (Linear Programming Technique) โดยมีสมการเป้าหมายทำให้ผลรวมของค่าสัมบูรณ์ของค่าความคลาดเคลื่อนต่ำสุด (Sukparungsee, 2003)

$$\hat{Y}_{C_t} = w_1 \hat{Y}_{1t} + w_2 \hat{Y}_{2t} + w_3 \hat{Y}_{3t} \quad \text{โดยที่} \quad \sum_{i=1}^3 w_i = 1 \quad (9)$$

เมื่อ w_i แทนน้ำหนักของวิธีการพยากรณ์ที่ i ; $i = 1, 2, 3$ ซึ่งหาได้จากการให้ผลรวมของค่าสัมบูรณ์ของค่าความคลาดเคลื่อนต่ำสุด โดยใช้คำสั่ง Solver Simplex LP ในโปรแกรม Microsoft Excel

3) วิธีการถ่วงน้ำหนักความคลาดเคลื่อนกำลังสองเฉลี่ยผกผัน (Inverse of Mean Squares Error Method: Inv) เป็นการนำค่าความคลาดเคลื่อนกำลังสองเฉลี่ยที่คำนวณได้จากแต่ละเทคนิคการพยากรณ์มาแปลงเป็นค่าถ่วงน้ำหนัก (Weiss et al., 2018)

$$\hat{Y}_{C_t} = w_1 \hat{Y}_{1t} + w_2 \hat{Y}_{2t} + w_3 \hat{Y}_{3t} \quad (10)$$

$$\text{โดยที่} \quad w_i = \frac{(1 / \text{MSE}_i)}{\sum_{i=1}^3 (1 / \text{MSE}_i)} ; i = 1, 2, 3 \quad (11)$$

$$\text{MSE} = \frac{\sum_{t=1}^{n_2} e_t^2}{n_2} \quad (12)$$

เมื่อ $e_t = Y_t - \hat{Y}_t$ แทนความคลาดเคลื่อนจากการพยากรณ์ ณ เวลา t และ

n_2 แทนจำนวนข้อมูลในอนุกรมเวลาชุดข้อมูลทดสอบ ($n_2 = 15$)

4. **เปรียบเทียบวิธีการพยากรณ์** การเปรียบเทียบประสิทธิภาพของตัวแบบพยากรณ์ โดยการเปรียบเทียบปริมาณการส่งออกถุงยางอนามัยของชุดข้อมูลทดสอบ ตั้งแต่เดือนมกราคม พ.ศ. 2563 ถึงเดือนมีนาคม พ.ศ. 2564 กับค่าพยากรณ์จากวิธีการสร้างตัวแบบทั้ง 4 วิธี ในการวิจัยครั้งนี้ได้พิจารณาใช้เกณฑ์ค่า RMSE ซึ่งเป็นเกณฑ์ที่เหมาะสมในการเปรียบเทียบวิธีการพยากรณ์หลายวิธีเมื่อใช้อนุกรมเวลาชุดเดียวกัน ในขณะที่เกณฑ์ MAPE ได้ถูกนำมาพิจารณาเนื่องจากเป็นค่าที่ไม่มีหน่วยและแสดงอยู่ในคำร้อยละ ดังสมการที่ (13) และ (14) ตามลำดับ

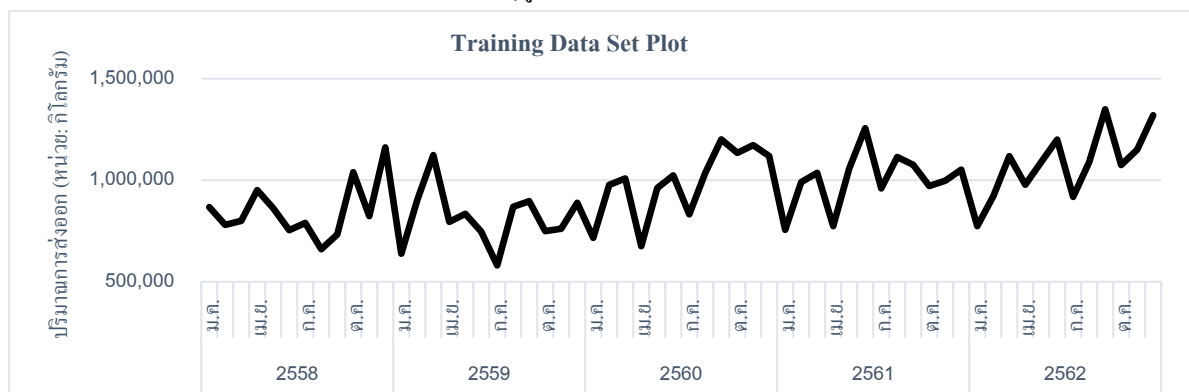
$$RMSE = \sqrt{\frac{\sum_{t=1}^{n_2} e_t^2}{n_2}} \quad (13)$$

$$MAPE = \frac{100}{n_2} \sum_{t=1}^{n_2} \left| \frac{e_t}{Y_t} \right| \quad (14)$$

โดยตัวแบบพยากรณ์ที่ให้ค่า RMSE และ MAPE ต่ำที่สุด จะถูกเลือกให้เป็นตัวแบบพยากรณ์ที่เหมาะสมที่สุดสำหรับพยากรณ์ปริมาณการส่งออกยางอนามัยของประเทศไทย ตั้งแต่เดือน เมษายน พ.ศ. 2564 ถึง เดือน มีนาคม พ.ศ. 2565 ต่อไป

ผลการวิจัย

จากการพิจารณาลักษณะการเคลื่อนไหวของอนุกรมเวลาปริมาณการส่งออกยางอนามัยของชุดข้อมูลฝึกฝนในภาพที่ 2 ตั้งแต่เดือนมกราคม พ.ศ. 2558 ถึง เดือนธันวาคม พ.ศ. 2562 จำนวน 60 ค่า พบว่ามีแนวโน้มการเปลี่ยนแปลงค่อนข้างจะคงที่ในช่วงปี 2558 – 2560 จากนั้นแนวโน้มปริมาณการส่งออกเริ่มมีการเพิ่มขึ้นต่อเนื่องในปี 2561 - 2562 และคาดว่าจะมีความผันแปรตามฤดูกาล



รูปที่ 2 ลักษณะการเคลื่อนไหวของอนุกรมเวลาปริมาณการส่งออกยางอนามัย ตั้งแต่ มกราคม พ.ศ. 2558 ถึง ธันวาคม พ.ศ. 2562

ในการนำเสนอผลการวิเคราะห์ด้วยเทคนิควิธีการพยากรณ์ 4 วิธี มีรายละเอียดดังนี้

1. ผลการสร้างตัวแบบพยากรณ์โดยวิธีการแยกส่วนประกอบ

จากการตรวจสอบส่วนประกอบของข้อมูลอนุกรมเวลาชุดข้อมูลฝึกฝนที่ระดับนัยสำคัญ 0.05 พบว่า อนุกรมมีส่วนประกอบของแนวโน้มจากสถิติทดสอบ Runs Test ($Z = -2.344$, $p\text{-value} = 0.019$) (Taesombat, 2006) และส่วนประกอบอิทธิพลของฤดูกาลจากสถิติทดสอบ Kruskal-Wallis Test ($\chi^2 = 24.60$, $p\text{-value} = 0.010$)

(Taesombat, 2006) รวมทั้งผลจากสถิติทดสอบความผันแปรของฤดูกาลโดย Levene's Test = 0.98 (p-value = 0.476)

พบว่าฤดูกาลไม่มีการเปลี่ยนแปลงเมื่อเวลาที่มีการเปลี่ยนแปลง ดังนั้นจึงใช้การวิเคราะห์โดยวิธีแยกส่วนประกอบรูปแบบบวก

ค่าประมาณจากโปรแกรม Minitab ได้ว่า $b_0 = 783,442$ และ $b_1 = 5,404$ ในขณะที่ค่าประมาณความผันแปรของฤดูกาล $\hat{S}_t$ แสดงดังตาราง 1

ตาราง 1 ค่าประมาณความผันแปรของฤดูกาลด้วยวิธีการแยกส่วนประกอบ

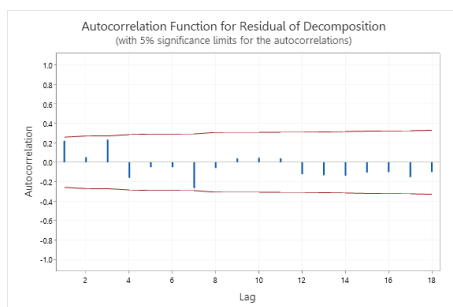
เดือน	$\hat{S}_t$	เดือน	$\hat{S}_t$	เดือน	$\hat{S}_t$	เดือน	$\hat{S}_t$
ม.ค.	-221,700.9	เม.ย.	- 149,259.1	ก.ค.	-100,391.7	ต.ค.	51,291.3
ก.พ.	6,728.0	พ.ค.	27,187.2	ส.ค.	49,075.2	พ.ย.	- 25,058.0
มี.ค.	109,800.0	มิ.ย.	91,990.1	ก.ย.	77,831.2	ธ.ค.	82,506.8

จากตาราง 1 จะเห็นว่าปริมาณการส่งออกในช่วงเดือนที่ 3 ของแต่ละไตรมาสจะมีค่ามากที่สุด และลดลงในเดือนต่อไป ซึ่งก็คือเดือนแรกของในไตรมาสถัดไป

ดังนั้นตัวแบบพยากรณ์ที่ได้มีความเหมาะสม โดยสามารถแสดงตัวแบบพยากรณ์ได้สมการ (15) ดังนี้

$$\hat{Y}_t = 783,442 + 5,404t + \hat{S}_t \quad (15)$$

จากนั้นได้ทำการตรวจสอบคุณลักษณะของความคลาดเคลื่อนจากการพยากรณ์ พบว่า ความคลาดเคลื่อนมีการแจกแจงปกติจากสถิติทดสอบ Anderson-Darling (AD = 0.430, p-value = 0.299) การเคลื่อนไหวนั้นเป็นอิสระกัน ตรวจสอบโดยพิจารณาจากกราฟ ACF ของความคลาดเคลื่อนจากการพยากรณ์ ดังภาพ 3 พบว่า ค่าสัมประสิทธิ์สหสัมพันธ์ในตัวของความคลาดเคลื่อนตกอยู่ในขอบเขตความเชื่อมั่นร้อยละ 95 ส่วนการทดสอบค่าเฉลี่ยความคลาดเคลื่อนเท่ากับศูนย์ ($t=0.00$, p-value = 1.00) และมีความแปรปรวนคงที่ทุกช่วงเวลาโดยสถิติทดสอบ Bartlett ($\chi^2 = 9.28$, p-value = 0.596)



รูปที่ 3 กราฟ ACF ของความคลาดเคลื่อนจากการพยากรณ์ โดยวิธีแยกส่วนประกอบ

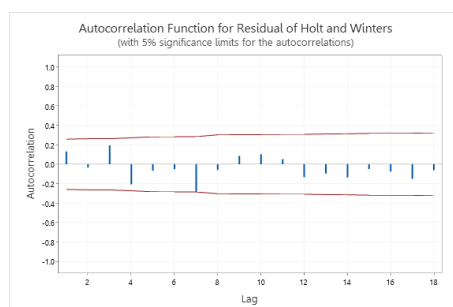
2. ผลการสร้างตัวแบบพยากรณ์โดยวิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์

จากการสร้างตัวแบบพยากรณ์โดยวิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์ รูปแบบบวก ใช้คำสั่ง Solver ในโปรแกรม Microsoft Excel ที่ทำให้ได้ค่า RMSE ต่ำที่สุด ได้ค่าทำให้เรียบที่เหมาะสม คือ

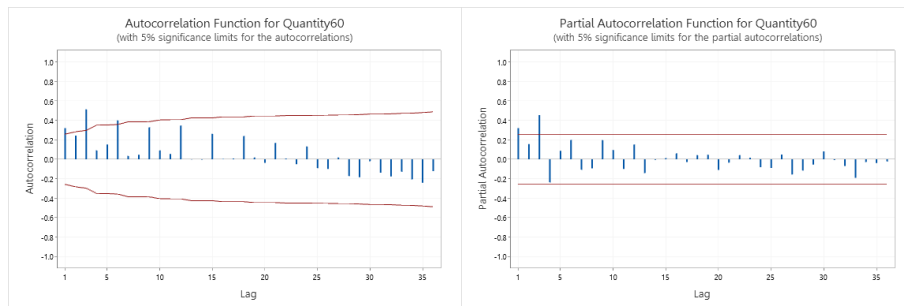
$\alpha = 0.2, \gamma = 0.00001013$ และ $\delta = 0.00005127$ เมื่อการตรวจสอบคุณลักษณะของความคลาดเคลื่อนจากการพยากรณ์ พบว่า ความคลาดเคลื่อนมีการแจกแจงปกติจากสถิติทดสอบ Anderson-Darling (AD = 0.416, p-value = 0.323) การเคลื่อนไหวเป็นอิสระกันตรวจสอบโดยพิจารณาจากกราฟ ACF ของความคลาดเคลื่อนจากการพยากรณ์ ดังภาพ 4 ซึ่งแสดงให้เห็นว่า ค่าสัมประสิทธิ์สหสัมพันธ์ในตัวของความคลาดเคลื่อนตกอยู่ในขอบเขตความเชื่อมั่นร้อยละ 95 ส่วนการทดสอบค่าเฉลี่ยความคลาดเคลื่อนเท่ากับศูนย์ ($t = -2.00$, p-value = 0.05) และมีความแปรปรวนคงที่ทุกช่วงเวลาโดยสถิติทดสอบ Bartlett ($\chi^2 = 7.02$, p-value = 0.797)

3. ผลการสร้างตัวแบบพยากรณ์โดยวิธีบ็อกซ์-เจนกินส์

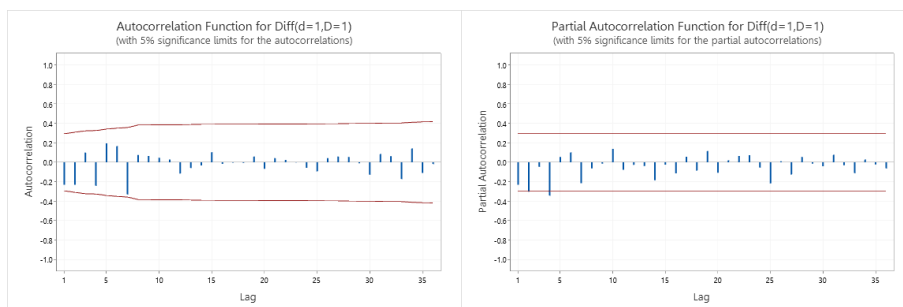
จากการพิจารณากราฟ ACF และ PACF ของปริมาณการส่งออกรายเดือนของยางพาราของประเทศไทย ตั้งแต่เดือนมกราคม พ.ศ. 2558 ถึงเดือนธันวาคม พ.ศ. 2562 จำนวน 60 ค่า ดังภาพที่ 5 ซึ่งแสดงให้เห็นว่าอนุกรมเวลาแบบไม่คงที่ พบว่ากราฟ ACF ในด้านซ้ายมือ ค่าสัมประสิทธิ์สหสัมพันธ์ในตัว ที่ช่วงเวลาที่ช้ากว่ากัน (lag) ที่ 1, 2, 3, ... มีลักษณะการเคลื่อนไหวลดลง (Die Down) ก่อนข้างช้า เช่นเดียวกับ ที่ lag 12, 24 และ 36 กราฟมีลักษณะเคลื่อนไหวลดลง ซึ่งคาดว่าข้อมูลน่าจะประกอบด้วยอิทธิพลของแนวโน้มและฤดูกาล แต่เมื่อทำการทดสอบสมมติฐานที่ lag 12 จะเห็นว่าอยู่ในช่วง 5% Significance limits for the autocorrelations ซึ่งสรุปว่าข้อมูลไม่มีอิทธิพลของฤดูกาล ดังนั้นผู้วิจัยจึงได้พิจารณาตัวแบบทั้ง 2 กรณี นั่นคือ กรณีแรกข้อมูลอนุกรมเวลาประกอบด้วยอิทธิพลของแนวโน้มและฤดูกาล ซึ่งเป็นอนุกรมเวลาแบบไม่คงที่ (Non-Stationary Time Series) ดังนั้นจึงทำการแปลงข้อมูลด้วยการหาผลต่างลำดับที่ 1 และหาผลต่างฤดูกาลลำดับที่ 1 นั่นคือ $d = 1$ และ $D = 1$ โดยกราฟ ACF และ PACF ของอนุกรมเวลาที่แปลงข้อมูลแล้ว แสดงดังภาพที่ 6 ส่วนกรณีที่ 2 ข้อมูลอนุกรมเวลาประกอบด้วยอิทธิพลของแนวโน้มเท่านั้น ดังนั้นจึงทำการแปลงข้อมูลด้วยการหาผลต่างลำดับที่ 1 นั่นคือ $d = 1$ โดยกราฟ ACF และ PACF ของอนุกรมเวลาที่แปลงข้อมูลแล้ว แสดงดังภาพที่ 7



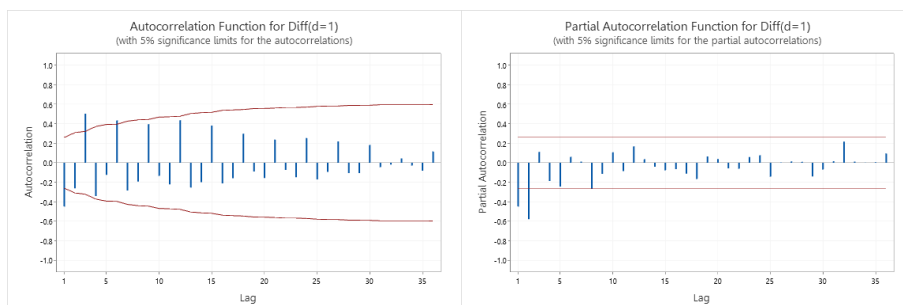
รูปที่ 4 กราฟ ACF ของความคลาดเคลื่อนจากการพยากรณ์ โดยวิธีการปรับเรียบ ด้วยเส้นโค้งเลขชี้กำลังของโฮลต์ และวินเทอร์ รูปแบบบวก



รูปที่ 5 กราฟ ACF และ PACF ของอนุกรมเวลาปริมาณการส่งออกของยานยนต์ ของประเทศไทย



รูปที่ 6 กราฟ ACF และ PACF ของอนุกรมเวลาปริมาณการส่งออกของยานยนต์ เมื่อแปลงข้อมูลด้วยการหาผลต่าง ลำดับที่ 1 และหาผลต่างฤดูกาลลำดับที่ 1 ($d = 1$ และ $D = 1$)



รูปที่ 7 กราฟ ACF และ PACF ของอนุกรมเวลาปริมาณการส่งออกของยานยนต์ เมื่อแปลงข้อมูลด้วยการหาผลต่าง ลำดับที่ 1 ($d = 1$)

จากรูปที่ 6 พบว่าการแปลงข้อมูลด้วยการหาผลต่างลำดับที่ 1 และหาผลต่างฤดูกาลลำดับที่ 1 ทำให้อนุกรมเวลามีลักษณะคงที่ ผู้วิจัยจึงได้กำหนด 2 ตัวแบบพยากรณ์เริ่มต้นที่คาดว่าเหมาะสมกับอนุกรมเวลา ได้แก่ SARIMA(4,1,0)(0,1,0)₁₂ และ SARIMA(2,1,0)(0,1,0)₁₂ พร้อมกับประมาณค่าพารามิเตอร์ ผลดังตาราง 2 เช่นเดียวกับภาพที่ 7 ผลจากการแปลงข้อมูลด้วยการหาผลต่างลำดับที่ 1 จะเห็นว่ากราฟ ACF ในด้านซ้ายมือค่าสัมประสิทธิ์สหสัมพันธ์ในตัว ที่ช่วงเวลาที่ 12, 24, 36 มีลักษณะการเคลื่อนไหวลดลงเร็ว และค่าสัมประสิทธิ์

สหสัมพันธ์ในตัวบางส่วนในช่วงเวลาที่ช้ากว่ากันที่ lag 12 มีค่าเพิ่มจากช่วงเวลาที่ช้ากว่าข้างเคียง ทำให้คาดว่าข้อมูลของฤดูกาลมีผลต่อการพยากรณ์แบบถดถอยในตัวแบบมีฤดูกาลอันดับ 1 จึงได้กำหนด 4 ตัวแบบพยากรณ์เริ่มต้นที่คาดว่าเหมาะสมกับอนุกรมเวลา ได้แก่ SARIMA(2,1,0)(1,0,0)₁₂ SARIMA(0,1,3)(1,0,0)₁₂ SARIMA(0,1,2)(1,0,0)₁₂ และ SARIMA(0,1,1)(1,0,0)₁₂ พร้อมกับประมาณค่าพารามิเตอร์ ผลดังตาราง 2

จากตาราง 2 การตรวจสอบตัวแบบพยากรณ์ เมื่อพิจารณาที่ระดับนัยสำคัญ 0.05 พบว่าตัวแบบพยากรณ์ที่มีค่า BIC ต่ำที่สุด คือตัวแบบ SARIMA(0,1,1)(1,0,0)₁₂ ที่ไม่มีพจน์ค่าคงที่ เท่ากับ 23.948 อย่างไรก็ตามเมื่อทำการตรวจสอบลักษณะของความคลาดเคลื่อนที่ได้จากตัวแบบ SARIMA(0,1,1)(1,0,0)₁₂ จากสถิติทดสอบ Ljung-Box Q ณ lag 18 (Ljung-Box Q lag 18 = 27.508, p-value = 0.036) ซึ่งแสดงว่าค่าความคลาดเคลื่อนที่ได้จากตัวแบบพยากรณ์ไม่มีความเป็นอิสระกัน จึงทำการพิจารณาค่า BIC ต่ำสุดอันดับถัดไป ได้แก่ตัวแบบ SARIMA(2,1,0)(1,0,0)₁₂ ที่ไม่มีพจน์ค่าคงที่ เท่ากับ 23.988 เมื่อตรวจสอบลักษณะของความคลาดเคลื่อนที่ได้จากตัวแบบ SARIMA(2,1,0)(1,0,0)₁₂ ที่ไม่มีพจน์ค่าคงที่ จากสถิติทดสอบ Ljung-Box Q ณ lag 18 (Ljung-Box Q lag 18 = 18.156, p-value = 0.255) ซึ่งแสดงว่าค่าความคลาดเคลื่อนที่ได้จากตัวแบบพยากรณ์มีความเป็นอิสระกัน สอดคล้องกับค่า ACF ในภาพที่ 8 ที่มีค่าตกอยู่ภายในช่วงความเชื่อมั่น 95% ของค่า ACF แสดงให้เห็นว่าค่าความคลาดเคลื่อนที่ได้จากตัวแบบ SARIMA(2,1,0)(1,0,0)₁₂ ที่ไม่มีพจน์ค่าคงที่ เป็นอิสระกัน

นอกจากนั้น เมื่อตรวจสอบคุณลักษณะของความคลาดเคลื่อนจากการพยากรณ์ พบว่า ความคลาดเคลื่อนมีการแจกแจงปรกติจากสถิติทดสอบ Anderson-Darling (AD = 0.325, p-value = 0.515) ส่วนการทดสอบค่าเฉลี่ยความคลาดเคลื่อนเท่ากับศูนย์ (t = 0.60, p-value = 0.553) และมีความแปรปรวนคงที่ทุกช่วงเวลาโดยสถิติทดสอบ Bartlett ($\chi^2 = 8.62$, p-value = 0.657) ดังนั้นตัวแบบ SARIMA(2,1,0)(1,0,0)₁₂ ที่ไม่มีพจน์ค่าคงที่ มีความเหมาะสมซึ่งจากสมการที่ (7) สามารถเขียนเป็นตัวแบบได้ดังนี้

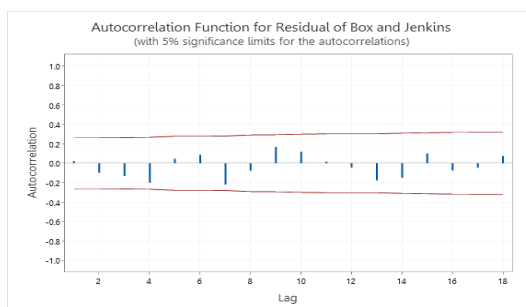
ตาราง 2 ค่าประมาณพารามิเตอร์ ค่า BIC และค่าสถิติ Ljung-Box Q ของตัวแบบ SARIMA(p,d,q)(P,D,Q)₁₂

ค่าประมาณพารามิเตอร์	SARIMA(p,d,q)(P,D,Q) ₁₂									
	SARIMA	SARIMA	SARIMA	SARIMA	SARIMA	SARIMA	SARIMA	SARIMA	SARIMA	SARIMA
	(4,1,0)	(2,1,0)	(2,1,0)	(0,1,3)	(0,1,2)	(0,1,1)	(4,1,0)	(2,1,0)	(2,1,0)	(0,1,1)
	(0,1,0) ₁₂	(0,1,0) ₁₂	(1,0,0) ₁₂	(1,0,0) ₁₂	(1,0,0) ₁₂	(1,0,0) ₁₂	(0,1,0) ₁₂	(0,1,0) ₁₂	(1,0,0) ₁₂	(1,0,0) ₁₂

ไม่มีพจน์ค่าคงที่ ไม่มีพจน์ค่าคงที่ ไม่มีพจน์ค่าคงที่ ไม่มีพจน์ค่าคงที่

ค่าคงที่	ค่าประมาณ	4,382.834	5,470.486	6,940.967	5,381.993	5,294.283	5,889.353		
	p-value	0.664 ^a	0.716 ^a	0.546 ^a	0.002	0.005	0.337 ^a		
AR(1)	ค่าประมาณ	-0.385	-0.357	-0.640				-0.384	-0.633
ϕ_1	p-value	0.009	0.018	0.000				0.008	0.000
AR(2)	ค่าประมาณ	-0.478	-0.358	-0.508				-0.478	-0.502
ϕ_2	p-value	0.004	0.015	0.000				0.003	0.000
AR(3)	ค่าประมาณ	-0.159						-0.155	
ϕ_3	p-value	0.288 ^a						0.295 ^a	
AR(4)	ค่าประมาณ	-0.423						-0.422	
ϕ_4	p-value	0.005						0.005	
MA(1)	ค่าประมาณ				0.660	0.678	0.787		0.738
θ_1	p-value				0.585 ^a	0.644 ^a	0.000		0.000
MA(2)	ค่าประมาณ				0.467	0.320			
θ_2	p-value				0.306 ^a	0.542 ^a			
MA(3)	ค่าประมาณ				-0.129				
θ_3	p-value				0.475 ^a				
SAR(1)	ค่าประมาณ			0.277	0.415	0.445	0.369	0.288	0.403
Φ_1	p-value			0.052 ^a	0.005	0.002	0.009	0.042	0.004
BIC*		24.474	24.394	24.068	24.119	24.046	24.017	24.374	24.293
Ljung-Box Q		10.068	14.339	18.159	24.252	27.876	28.792	10.199	14.354
p-value		0.757	0.573	0.254	0.043 ^a	0.022 ^a	0.025 ^a	0.747	0.572
								0.255	0.036 ^a

หมายเหตุ: ^aไม่ผ่านการทดสอบการประมาณค่าที่ระดับนัยสำคัญ 0.05



รูปที่ 8 กราฟ ACF ของความคลาดเคลื่อนของตัวแบบ SARIMA(2,1,0)(1,0,0)₁₂ ที่ไม่มีพจน์ค่าคงที่

$$\phi_2(B)\Phi_1(B^{12})(1-B)Y_t = \varepsilon_t$$

$$(1-\phi_1 B - \phi_2 B^2)(1-\Phi_1 B^{12})(1-B)Y_t = \varepsilon_t$$

$$Y_t = (\phi_1 + 1)Y_{t-1} + (\phi_2 - \phi_1)Y_{t-2} - \phi_2 Y_{t-3} + \Phi_1 Y_{t-12} - (\phi_1 + 1)\Phi_1 Y_{t-13} - (\phi_2 - \phi_1)\Phi_1 Y_{t-14} + \phi_2 \Phi_1 Y_{t-15} + \varepsilon_t$$

เมื่อแทนค่าประมาณพารามิเตอร์จากตาราง 2 จะได้ตัวแบบ พยากรณ์ดังนี้

$$\hat{Y}_t = 0.367490Y_{t-1} + 0.130781Y_{t-2} + 0.501729Y_{t-3} + 0.287854Y_{t-12} - 0.105784Y_{t-13} - 0.037646Y_{t-14} - 0.144425Y_{t-15} \quad (16)$$

4. ผลการสร้างตัวแบบการพยากรณ์โดยวิธีการพยากรณ์รวม

จากการใช้ตัวแบบพยากรณ์ 3 วิธีแรกมาสร้างตัวแบบการพยากรณ์รวม โดยแยกวิธีการได้มาซึ่งตัวถ่วงน้ำหนัก 3 วิธี ได้ผลลัพธ์ดังตารางที่ 3

- 1) วิธีการถ่วงน้ำหนักด้วยสัมประสิทธิ์การถดถอย

$$\hat{Y}_{C_t} = 9,184.535 + 0.570587\hat{Y}_{1t} + 0.127624\hat{Y}_{2t} + 0.238245\hat{Y}_{3t} \quad (17)$$

- 2) วิธีการถ่วงน้ำหนักค่าสัมบูรณ์ต่ำสุด

$$\hat{Y}_{C_t} = 1.00\hat{Y}_{1t} + 0.00\hat{Y}_{2t} + 0.00\hat{Y}_{3t} \quad (18)$$

- 3) วิธีการถ่วงน้ำหนักความคลาดเคลื่อนกำลังสองเฉลี่ยผกผัน

$$\hat{Y}_{C_t} = 0.4566\hat{Y}_{1t} + 0.2511\hat{Y}_{2t} + 0.2924\hat{Y}_{3t} \quad (19)$$

โดย t ของสมการ (17) (18) และ (19) คือ $t = 1, 2, 3, 4, \dots, 15$ เมื่อ $t = 1$ คือ เดือนมกราคม พ.ศ. 2563

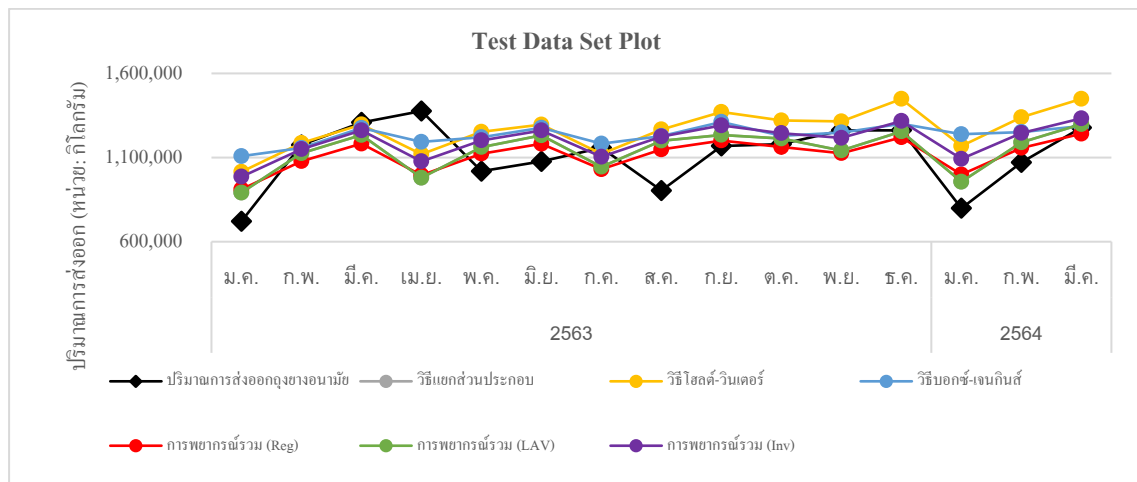
5. ผลการเปรียบเทียบวิธีการพยากรณ์

จากการใช้ตัวแบบพยากรณ์ 4 วิธี ได้ค่าพยากรณ์ปริมาณการส่งออกกุ้งยางอนามัยของประเทศไทยรายเดือน ตั้งแต่ มกราคม พ.ศ. 2563 ถึง มีนาคม พ.ศ. 2564 ในแต่ละวิธีดังตาราง 3 และภาพที่ 9 รวมทั้งค่า RMSE และ MAPE ของวิธีแยกส่วนประกอบ รูปแบบบวก วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลต์และวินเตอร์ วิธีบอกซ์-เจนกินส์ และวิธีการพยากรณ์รวม โดยวิธีการพยากรณ์รวมถ่วงน้ำหนักด้วยสัมประสิทธิ์การถดถอย ให้ค่า RMSE และ MAPE ต่ำที่สุดสอดคล้องกัน ดังนั้นในการพยากรณ์ปริมาณการส่งออกกุ้งยางอนามัยรายเดือนของประเทศไทย จะเลือกใช้วิธีการพยากรณ์รวมถ่วงน้ำหนักด้วยสัมประสิทธิ์การถดถอย

ตาราง 3 ค่าจริงและค่าพยากรณ์ปริมาณการส่งออกกุ้งยางอนามัย (กิโลกรัม) รายเดือนของประเทศไทย ตั้งแต่ เดือน มกราคม พ.ศ. 2563 ถึง เดือนมีนาคม พ.ศ. 2564

เดือน	ปริมาณการส่งออก	ปริมาณการส่งออก (หน่วย: กิโลกรัม)					
		แยกส่วนประกอบ	โฮลต์และวินเตอร์	บอกซ์-เจนกินส์	การพยากรณ์รวม (Reg)	การพยากรณ์รวม (LAV)	การพยากรณ์รวม (Inv)
มกราคม-63	721,443	891,359	1,015,238	1,108,334	911,406	891,359	985,897
กุมภาพันธ์-63	1,174,934	1,125,191	1,186,178	1,156,848	1,078,202	1,125,191	1,149,758
มีนาคม-63	1,308,341	1,233,667	1,295,969	1,275,571	1,182,394	1,233,667	1,261,560

เมษายน-63	1,376,297	980,011	1,120,083	1,193,009	995,544	980,011	1,077,451
พฤษภาคม-63	1,018,962	1,161,861	1,253,027	1,220,512	1,122,824	1,161,861	1,201,897
มิถุนายน-63	1,077,018	1,232,068	1,295,477	1,276,502	1,181,640	1,232,068	1,260,978
กรกฎาคม-63	1,157,467	1,045,089	1,122,022	1,182,236	1,030,357	1,045,089	1,104,501
สิงหาคม-63	903,809	1,199,960	1,266,816	1,227,786	1,148,056	1,199,960	1,224,880
กันยายน-63	1,168,557	1,234,119	1,370,735	1,311,552	1,200,766	1,234,119	1,291,057
ตุลาคม-63	1,180,216	1,212,983	1,320,481	1,228,829	1,162,584	1,212,983	1,244,604
พฤศจิกายน-63	1,261,848	1,142,037	1,314,451	1,248,555	1,126,033	1,142,037	1,216,465
ธันวาคม-63	1,261,848	1,255,006	1,448,401	1,300,488	1,219,960	1,255,006	1,316,857
มกราคม-64	798,493	956,202	1,168,180	1,238,615	998,962	956,202	1,091,989
กุมภาพันธ์-64	1,071,237	1,190,034	1,339,120	1,251,580	1,157,289	1,190,034	1,245,457
มีนาคม-64	1,277,441	1,298,510	1,448,911	1,286,874	1,241,604	1,298,510	1,332,867
RMSE		162,883	219,656	203,547	157,956	162,883	178,712
MAPE		12.27314	18.706	16.02937	12.26972	12.27314	14.69836



รูปที่ 9 การเปรียบเทียบค่าพยากรณ์ปริมาณการส่งออกของอนามัยจากวิธีการทางสถิติ 6 วิธี

6. การพยากรณ์ปริมาณการส่งออกของอนามัย

ในการพยากรณ์ปริมาณการส่งออกของอนามัยรายเดือนในเดือนเมษายน พ.ศ. 2564 ถึงเดือนมีนาคม พ.ศ. 2565 จะเลือกใช้วิธีการพยากรณ์รวมโดยถ่วงน้ำหนักด้วยสัมประสิทธิ์การถดถอย กับข้อมูลอนุกรมเวลาทั้งหมดจำนวน 75 ค่า ตั้งแต่เดือนมกราคม พ.ศ. 2558 ถึงเดือนมีนาคม พ.ศ. 2564 ด้วยตัวแบบที่สร้างจากวิธีการพยากรณ์

รวมโดยถ่วงน้ำหนักด้วยสัมประสิทธิ์การถดถอย $\hat{Y}_{C_t} = -10,200 + 0.584\hat{Y}_{1t} + 0.597\hat{Y}_{2t} - 0.198\hat{Y}_{3t}$ รายละเอียด

ดังตาราง 4

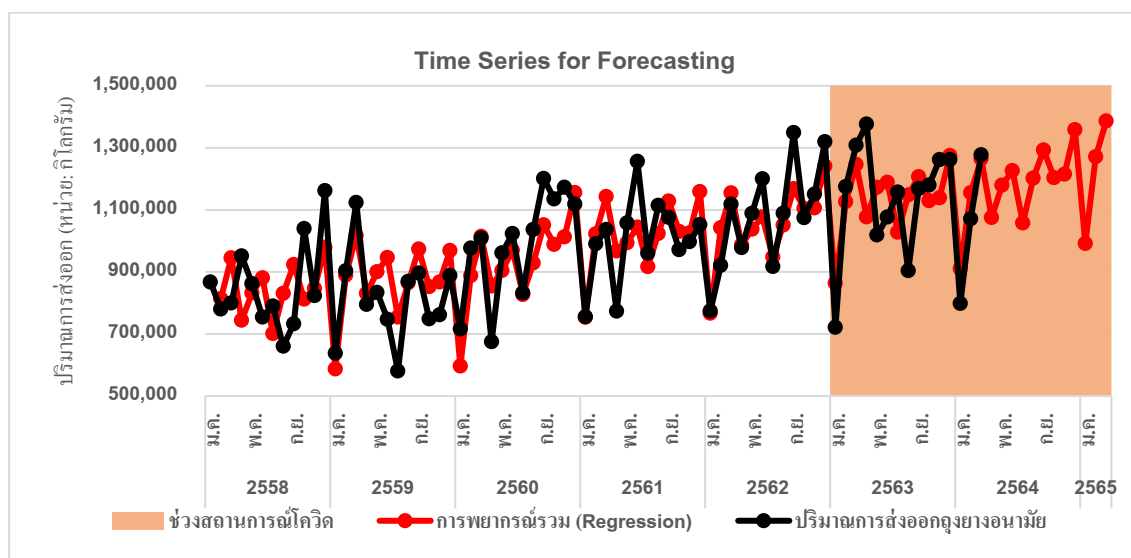
ตาราง 4 ค่าพยากรณ์ปริมาณการส่งออกกุ้งยางอนามัย (กิโลกรัม) ตั้งแต่ เดือนเมษายน พ.ศ. 2564 ถึง เดือนมีนาคม พ.ศ. 2565

ช่วงเวลา	ค่าพยากรณ์	ช่วงเวลา	ค่าพยากรณ์
เมษายน-64	1,075,249	ตุลาคม-64	1,203,816
พฤษภาคม-64	1,179,697	พฤศจิกายน-64	1,215,225
มิถุนายน-64	1,226,588	ธันวาคม-64	1,358,239
กรกฎาคม-64	1,057,790	มกราคม-65	991,777
สิงหาคม-64	1,202,204	กุมภาพันธ์-65	1,270,996
กันยายน-64	1,292,766	มีนาคม-65	1,386,161

วิจารณ์ผลการวิจัย

จากสถานการณ์การแพร่ระบาดของโควิด-19 ในประเทศไทยที่เริ่มตั้งแต่ช่วงต้นปี พ.ศ. 2563 จะเห็นได้ว่า ข้อมูลอนุกรมเวลาปริมาณการส่งออกกุ้งยางอนามัยของประเทศไทยยังคงไม่ได้รับผลกระทบจากสถานการณ์การแพร่ระบาดดังกล่าวดังภาพที่ 10 สอดคล้องกับ Thairath online (2020) ดังนั้นผลการวิเคราะห์ที่จากการสร้างตัวแบบเพื่อการพยากรณ์ข้อมูลปริมาณส่งออกกุ้งยางอนามัยทั้ง 6 วิธี จึงสามารถนำไปใช้พยากรณ์หรือทำนายปริมาณการส่งออกกุ้งยางอนามัยได้มีความแม่นยำระดับหนึ่ง เนื่องจากการสร้างตัวแบบพยากรณ์ อยู่ภายใต้ข้อสมมุติที่ว่า “รูปแบบพยากรณ์ได้จากข้อมูลในอดีต เหมือนเดิมหรือไม่มีการเปลี่ยนแปลงในอนาคต” (Rangkakulnuwat, 2019) ซึ่งผลจากการวิจัยครั้งนี้ ตัวแบบพยากรณ์ด้วยวิธีการพยากรณ์รวมถ่วงน้ำหนักการวิเคราะห์ถดถอย ให้ค่า RMSE และ MAPE ต่ำที่สุด เมื่อเปรียบเทียบกับค่าจริง จึงเป็นตัวแบบที่เหมาะสมที่สุดในการใช้ทำนายข้อมูลปริมาณการส่งออกกุ้งยางอนามัยในช่วง เมษายน พ.ศ. 2564 ถึง มีนาคม พ.ศ. 2565 ผลที่ได้แสดงดังภาพที่ 10 ซึ่ง DeLurgio (1998) เคยได้ให้ข้อสังเกตว่าการที่รวมเทคนิคการพยากรณ์เดี่ยวเข้าด้วยกัน อาจเพิ่มความแม่นยำในการพยากรณ์เมื่อประสิทธิภาพของเทคนิคการพยากรณ์เดี่ยวแต่ละวิธีไม่ได้โดดเด่นไปกว่ากัน อย่างไรก็ตามค่าพยากรณ์ที่ได้จากตัวแบบพยากรณ์เดี่ยวโดยวิธีการแยกส่วนประกอบก็เป็นอีกวิธีที่มีความน่าเชื่อถือ เนื่องจากถ้าพิจารณาจากภาพที่ 9 และ ตารางที่ 3 พบว่าค่าพยากรณ์ของตัวแบบพยากรณ์เดี่ยวด้วยวิธีการแยกส่วนประกอบ มีความแตกต่างจากวิธีการพยากรณ์รวมโดยถ่วงน้ำหนักด้วยสัมประสิทธิ์การถดถอยไม่มาก โดยมีค่า RMSE และ MAPE ใกล้เคียงกันมาก อีกทั้งมีค่าพยากรณ์ไปในทิศทางเดียวกัน และค่าพยากรณ์ของตัวแบบพยากรณ์เดี่ยวด้วยวิธีการแยกส่วนประกอบไม่มีความแตกต่างจากวิธีการพยากรณ์รวมโดยวิธีการถ่วงน้ำหนักค่าสัมบูรณ์ต่ำสุด โดยมีตัวแบบสมการพยากรณ์

จากวิธีแยกตัวประกอบดังนี้ $\hat{Y}_t = 790,781 + 5,059t + \hat{S}_t$ ในขณะที่ค่าประมาณความผันแปรของฤดูกาล $\hat{S}_t$ จากโปรแกรม Minitab แสดงดังตาราง 5



* หมายเหตุ สถานการณ์การแพร่ระบาดของโควิด-19 ในประเทศไทยเริ่มตั้งแต่ปี พ.ศ. 2563

รูปที่ 10 การเปรียบเทียบอนุกรมเวลาปริมาณการส่งออกยางอนามัย และค่าพยากรณ์จากวิธีการทางสถิติ โดยวิธีการพยากรณ์รวมถ่วงน้ำหนักการวิเคราะห์ถดถอย

ตาราง 5 ค่าประมาณความผันแปรของฤดูกาลด้วยวิธีการแยกตัวประกอบ

เดือน	$\hat{S}_t$	เดือน	$\hat{S}_t$	เดือน	$\hat{S}_t$	เดือน	$\hat{S}_t$
ม.ค.	-238,127	เม.ย.	-57,028	ก.ค.	-97,894	ต.ค.	-35,536
ก.พ.	38,325	พ.ค.	16,418	ส.ค.	29,901	พ.ย.	-19,925
มี.ค.	126,203	มิ.ย.	52,397	ก.ย.	80,329	ธ.ค.	104,937

สรุป

การวิจัยครั้งนี้ได้นำเสนอการเลือกตัวแบบพยากรณ์ที่เหมาะสมกับอนุกรมเวลาปริมาณการส่งออกยางอนามัยรายเดือนของประเทศไทย ช่วงสถานการณ์ โควิด-19 จากข้อมูลจากกรมศุลกากร พบว่าการพยากรณ์ปริมาณการส่งออกยางอนามัยตั้งแต่เดือนเมษายน พ.ศ. 2564 ถึง เดือนมีนาคม พ.ศ. 2565 ด้วยวิธีการพยากรณ์รวมโดยถ่วง

นำหลักด้วยสัมประสิทธิ์การถดถอยเป็นตัวแทนที่เหมาะสมเนื่องจากให้ค่า RMSE และ MAPE ต่ำที่สุด โดยมีตัวแบบสมการพยากรณ์คือ $\hat{Y}_{C_t} = -10,200 + 0.584\hat{Y}_{1t} + 0.597\hat{Y}_{2t} - 0.198\hat{Y}_{3t}$

ข้อเสนอแนะ

สำหรับการศึกษาการสร้างตัวแบบพยากรณ์ปริมาณการส่งออกของยางอนามัยรายเดือนของประเทศนี้ เนื่องจากข้อมูลปริมาณการส่งออกของยางอนามัยไม่ได้ขึ้นอยู่กับปัจจัยทางด้านของเวลาอย่างเดียว ดังนั้นในการศึกษา

ครั้งต่อไป จึงควรมีการศึกษาพิจารณาปัจจัยอื่นๆ เช่นปริมาณน้ำยางธรรมชาติ หรือน้ำยางข้นที่ผลิตได้ รวมไปถึงปริมาณการส่งออกของประเทศคู่แข่งการส่งออกของยางอนามัย ในการพิจารณาสร้างตัวแบบพยากรณ์ด้วย

เอกสารอ้างอิง

Abraham, B., & Ledolter, J. (1983). *Statistical Method for Forecasting*. John Wiley & Sons.

Anderson, T. W., & Darling, D. A. (1954). A test of goodness of fit. *Journal of the American Science Association*, 49, 765-769.

Bowman, B. L., & O'Connell, R. T. (1993). *Forecasting and Time Series: An Applied Approach* (3rd ed.). Duxbury Press, Belmont, California.

Box, G. E. P., Jenkins, G. M., & Reinsel, G. C. (1994). *Time Series Analysis: Forecasting and Control* (3rd ed.). Prentice Hall.

Brown, M. B., & Forsythe, A. B. (1974). Robust tests for the equality of variance. *Journal of the American Statistical Association*, 69, 364-367.

DeLurgio, S. A. (1998). *Forecasting principles and applications*. McGraw-Hill.

Department of International Trade Promotion. (2020). *Condoms are in short world supply, competition with face masks*. <https://gestyy.com/eu5JdF> (in Thai)

Glaser, R. E. (1976). Exact Critical Values for Bartlett's Test for Homogeneity of Variances. *Journal of the American Statistical Association*, 71(354), 488-490.

Keerativibool, W. (2013). A Comparison of Forecasting Methods between Box-Jenkins, Simple Seasonal Exponential Smoothing, and Combined Forecasting Methods for Predicting Monthly Mean Temperature. *Burapha Science Journal*, 18(2), 149-160.

Ket-iam, S. (2005). *Forecasting technique* (2nd ed.). Thaksin University. (in Thai)

Levene, H. (1960). *Contributions to Probability and Statistics*. Stanford University Press.

- Lorchirachoonkul, V., & Jitthavech, J. (2005). *Forecasting techniques* (3rd ed.). Bangkok: Project to promote academic documents, National Institute of Development Administration (NIDA). (in Thai)
- Manmin, Mookda. (2006). *Time series and forecasting*. Four Prining. (in Thai)
- Microsot. (2016). Microsoft Excel [Computer software]. <https://www.microsoft.com/en-us/microsoft-365/excel>
- Minilab. (2018). Minitab Statistical Software (Version 18) [Computer software]. <https://www.minitab.com/en-us/>
- Minsan, W. (2014). Seasonal index estimation by GRG2 decomposition method. In *the Proceedings of 10th Mahasarakham Research Conference* (pp. 243-249). Mahasarakham University, Mahasarakham. (in Thai)
- Montgomery, D. C., Peck, E. A., & Vining, G. G. (2006). *Introduction to Linear Regression Analysis* (4th ed.). Wiley.
- Panichkitkosolkul, W. (2009). Monthly Rainfall Amount Forecasting of Meteorological Stations and Agrometeorological Stations in Northeastern Thailand. *Thai Science and Technology Journal*, 17(2), 1-12. (in Thai)
- Phadungkarn, S. (2020). *Because of what other products, sales fell due to COVID-19, but the condom sales increase?* <https://www.brandage.com/article/18028/TNR> (in Thai)
- Rangkakulnuwat, P. (2019). *Time Series Analysis for Economics and Business*. <https://gestyy.com/eu6yX5> (in Thai)
- Riansut, W. (2020). Daily Wind Speed Forecast Model at an Altitude of 120 Meters, Pak Phanang District, Nakhon Si Thammarat Province. *The Journal of Applied Science*, 19(1), 95-109.
- Ruppert, D. (2011). *Statistics and Data Analysis for Financial Engineering*. Springer.
- Sawa, T. (1978). Information Criteria for Discriminating among Alternative Regression Model. *Journal of Econometrica*, 46, 1273-1282.
- Sukparungsee, S. (2003). Using the Combined Forecasting Technique by Weighted Average. *Journal of Technical Education Development*, 15(46), 1-7.
- Taesombat, S. (2006). *Quantitative Forecasting*. Kasetsart University Press. (in Thai)
- Thai Customs. (2020). *Statistic Report*. <https://bit.ly/2BXDNiK>
- Tirkes, G., Guray, C., & Celebi, N. E. (2017). Demand forecasting: a comparison between the Holt-Winters, trend analysis and decomposition models. *Tehnicki Vjesnik-Technical Gazette*. 24(2), 503-509.

- Thairath online. (2020). *TNR reveals condoms growing against economic trends in the first 9 months of the year, making a profit of 68 million, growing 70%*. <https://gestyy.com/eu6tuP> (in Thai)
- Weiss, C. E., Raviv, E., & Roetzer, G. (2018). Forecast Combinations in R using the ForecastComb Package. *The R Journal*, 10(2), 262-281.
- Yonar, H., Yonar, A., Tekindal, M. A., & Tekindal, M. (2020). Modeling and Forecasting for the number of cases of the COVID-19 pandemic with the Curve Estimation Models, the Box-Jenkins and Exponential Smoothing Methods. *Eurasian Journal of Medicine and Oncology*, 4(2), 160-165.

Research Article

การพยากรณ์ราคาข้าวเปลือกหอมมะลิ ด้วยเทคนิคการทำเหมืองข้อมูล

Forecast of the price of jasmine rice with Data Mining Techniques

เฉลิมวุฒิ คำเมือง^{1*} และ Wachirarak Orosram¹

Chalermwut Comemuang^{1*} and Wachirarak Orosram¹

¹สาขาวิชาคณิตศาสตร์ คณะวิทยาศาสตร์ มหาวิทยาลัยราชภัฏบุรีรัมย์ อ.เมือง จ.บุรีรัมย์ 31000 ประเทศไทย

¹Department of Mathematics, Faculty of Science, Buriram Rajabhat University, Muang, Buriram 31000, Thailand

*E-mail: chalermwut.cm@bru.ac.th

Received: 30/03/2021; Revised: 23/09/202021; Accepted: 22/09/2021

บทคัดย่อ

การวิจัยนี้มีวัตถุประสงค์เพื่อศึกษาแบบจำลองทางคณิตศาสตร์เพื่อทำนายราคาข้าวเปลือกหอมมะลิ ข้อมูลที่ใช้ในการศึกษาราคาข้าวเปลือกหอมมะลิเฉลี่ยรายเดือน ตั้งแต่ มกราคม พ.ศ. 2553 - ธันวาคม พ.ศ. 2563 จำนวน 132 ข้อมูล โดยแบ่งข้อมูลออกเป็น 2 ชุด ข้อมูลชุดที่ 1 ตั้งแต่เดือน มกราคม พ.ศ. 2553 - ธันวาคม พ.ศ. 2562 จำนวน 120 ข้อมูล สำหรับศึกษาตัวแบบพยากรณ์ โดยวิธีบ็อกซ์เจนกินส์ (Box-Jenkins) วิธีแบ็กกิ้ง (Bagging) และ วิธีแรนดอมฟอเรสต์ (Random Forest) ข้อมูลชุดที่ 2 ตั้งแต่เดือนมกราคม พ.ศ. 2563 - ธันวาคม พ.ศ. 2563 จำนวน 12 ค่า สำหรับการเปรียบเทียบความแม่นยำของค่าพยากรณ์ โดยใช้เกณฑ์รากที่สองของความคลาดเคลื่อนเฉลี่ย (Root Mean Square Error: RMSE) และเกณฑ์เปอร์เซ็นต์ความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (Mean Absolute Percentage Error: MAPE) ผลการศึกษาพบว่าจากวิธีการพยากรณ์ที่เหมาะสมที่สุดคือ วิธีแรนดอมฟอเรสต์

คำสำคัญ: พยากรณ์, บ็อกซ์เจนกินส์, แรนดอมฟอเรสต์, แบ็กกิ้ง

Abstract

The purpose of this research was to construct model for predicting the price of Jasmine Paddy price. Data used to study the average monthly price of jasmine rice from January 2010 to December 2020, 132 values, divided into two sets. The first set from January 2010 to December 2019, 120 values were used for the modeling by of box jenkins method, bagging and random forest. Another set from January 2020 to December 2020, 12 values were used for checking the accuracy of the forecasting models via the determination of the lowest root mean square error and mean absolute percentage error. The results showed that all forecasting methods studied, the random forest was the most appropriate method.

Keywords: forecast, box-jenkins, random forest, bagging

บทนำ

ข้าวเป็นธัญญาหารหลักของชาวโลก จัดเป็นพืชสายพันธุ์เดียวกับหญ้า ซึ่งนับได้ว่าเป็นหญ้าที่มีขนาดใหญ่ที่สุดในโลก และมีความหลากหลายทางชีวภาพ สามารถปลูกขึ้นได้ง่ายมีความทนทานต่อทุกสภาพภูมิประเทศในโลก ไม่ว่าจะเป็นถิ่นแห้งแล้งแบบทะเลทราย พื้นที่ราบลุ่มน้ำท่วมถึง หรือแม้กระทั่งบนเทือกเขาที่หนาวเย็น ข้าวก็ยังสามารถงอกงามขึ้นมาได้อย่างทรหดอดทน ในส่วนของประเทศไทยข้าวเป็นพืชที่เป็นอาหารประจำชาติและมีตำนานประวัติศาสตร์มายาวนาน ปรากฏเป็นร่องรอยพร้อมกับอารยธรรมไทยมาไม่น้อยกว่า 5,500 ปี และข้าวหอมมะลิเป็นพันธุ์ข้าวที่ได้ถือกำเนิดขึ้นในประเทศไทย โดยมีจุดเด่นคือเมล็ดมีคุณภาพดี สี และกลิ่นหอมคล้ายใบเตย เวลารับประทานจะมีความนุ่มและหอมอร่อยมากกว่าข้าวทั่วไป และประเทศไทยเป็นประเทศสิทธกรรม ประชาชนส่วนใหญ่เป็นกสิกร ทำการเพาะปลูกพืชสวน และพืชไร่ โดยที่ข้าวมีพื้นที่ปลูกมากกว่าพืชชนิดอื่น ๆ คิดเป็นพื้นที่ประมาณ 11.3% ของพื้นที่ทั่วประเทศและมีแนวโน้มการปลูกข้าวเพิ่มมากขึ้น ภาคกลาง และภาคตะวันออกเฉียงเหนือ มีพื้นที่ทำนามากที่สุด รองลงมา ได้แก่ ภาคเหนือ และภาคใต้ตามลำดับ

การส่งออกข้าวในช่วง มกราคม พ.ศ. 2563 – พฤศจิกายน พ.ศ. 2563 มีปริมาณ 5.24 ล้านตัน ลดลง 26.3% จากช่วงเดียวกันของปี 2562 และมีมูลค่า 106,656 ล้านบาท หรือ 3,418.8 ล้านดอลลาร์สหรัฐ ลดลง 12.1% จากช่วงเดียวกันของปีก่อน ซึ่งไทยถือเป็นผู้ส่งออกอันดับที่ 3 โดยประเทศอินเดียที่ส่งออกได้เป็นอันดับที่ 1 จำนวน 13.05 ล้านตัน และเวียดนามส่งออกได้เป็นอันดับที่ 2 จำนวน 5.79 ล้านตัน แต่ผลจากค่าเงินบาทแข็ง ทำให้ราคาข้าวไทยสูงกว่าคู่แข่ง 150 เหรียญสหรัฐต่อดัน โดยเมื่อวันที่ 23 ธันวาคม พ.ศ. 2563 ราคาข้าวขาว 5% ราคาตันละ 529 เหรียญสหรัฐ ขณะที่ราคาข้าวของเวียดนามราคาตันละ 498-502 เหรียญสหรัฐ และราคาข้าวของอินเดียราคาตันละ 368-372 เหรียญสหรัฐ (The Thai Rice Exporters Association, 2020)

ปัญหาของราคาข้าวเกิดจากหลายสาเหตุ ปัญหาด้านต้นทุนการผลิตและผลผลิตต่อไร่ซึ่งต้นทุนการผลิตต่อไร่สูงแต่ผลผลิตต่อไร่ต่ำ ระบบชลประทานในพื้นที่เพาะปลูกมีจำกัด ราคาส่งออกข้าวไทยสูงกว่าประเทศเพื่อนบ้าน ซึ่งในปัจจุบันรัฐบาลใช้วิธีการจํานาข้าวหอมมะลิในราคาสูง เพื่อช่วยเหลือเกษตรกรที่ปลูกข้าวหอมมะลิ ในขณะที่เดียวกันก็ทำให้เกิดปัญหาการปลอมปนข้าว เพราะมีผู้นำเข้าบางส่วนสั่งซื้อข้าวหอมมะลิจากไทยไป แล้วนำไปผสมในต่างประเทศ และอ้างว่าเป็นข้าวหอมมะลิจากไทย ซึ่งสร้างผลเสียต่อภาพลักษณ์การส่งออกข้าวหอมมะลิไทย ดังนั้นการพยากรณ์ราคาจึงสามารถช่วยทำให้ประเมินราคาลินค้าในอนาคตได้ การศึกษาแนวโน้มและพยากรณ์ราคาข้าวในประเทศไทยจึงมีความจำเป็นอย่างมาก เพื่อที่จะนำผลการศึกษาที่ได้จากการพยากรณ์มาเป็นแนวทางในการตัดสินใจสำหรับผู้ผลิตและผู้ส่งออก ในการวางแผนการผลิตและการส่งออกให้มีประสิทธิภาพดียิ่งขึ้นและสอดคล้องกับความต้องการของตลาด

Kodosueb & Boonlha (2016) ได้ศึกษาตัวแบบที่เหมาะสมและเพื่อเปรียบเทียบตัวแบบสำหรับการพยากรณ์ราคาข้าวหอมมะลิ 105 จากข้อมูลรายเดือนราคาข้าวหอมมะลิ 105 ตั้งแต่เดือนมกราคม พ.ศ.2540 ถึงเดือนธันวาคม พ.ศ.2558 รวม 18 ปี (หรือ 216 เดือน) ซึ่งเก็บรวบรวมโดยสำนักงานเศรษฐกิจการเกษตร ศึกษาตัวแบบการพยากรณ์โดยวิธีการปรับให้เรียบแบบเอกซ์โปเนนเชียลซ้ำสองครั้ง และวิธีบ็อกซ์และเจนกินส์ และเปรียบเทียบตัวแบบสำหรับการพยากรณ์โดยใช้ค่าคลาดเคลื่อนกำลังสองเฉลี่ย (MSE) น้อยที่สุด ผลการศึกษาพบว่า วิธีการพยากรณ์ที่เหมาะสมที่สุดคือ วิธีบ็อกซ์และเจนกินส์ ซึ่งได้ตัวแบบพยากรณ์ คือ ARIMA (1,1,2) Riansu (2016) การสร้างตัวแบบพยากรณ์ที่เหมาะสมกับอนุกรมเวลาราคาข้าวเปลือกเจ้าความชื้น 15% โดยใช้อนุกรมเวลาราคาข้าวเปลือกเจ้าความชื้น 15% จากเว็บไซต์ของสำนักงานเศรษฐกิจการเกษตร ตั้งแต่เดือนมกราคม 2548 ถึงเดือนกันยายน 2558 จำนวน 129 ค่าซึ่งข้อมูลถูกแบ่งออกเป็น 2 ชุด ข้อมูล ชุดที่ 1 ตั้งแต่เดือนมกราคม 2548 ถึงเดือนธันวาคม 2557 จำนวน 120 ค่า สำหรับการสร้างตัวแบบพยากรณ์ด้วยวิธี บ็อกซ์เจนกินส์วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังของโฮลด์ วิธีการปรับเรียบด้วยเส้นโค้งเลขชี้กำลังที่มีแนวโน้มแบบแฉก และวิธีการพยากรณ์รวม ข้อมูลชุดที่ 2 ตั้งแต่เดือนมกราคมถึงเดือนกันยายน 2558 จำนวน 9 ค่านำมาใช้สำหรับการเปรียบเทียบความแม่นยำของค่าพยากรณ์โดยใช้เกณฑ์เปอร์เซ็นต์ความคลาดเคลื่อนสัมบูรณ์เฉลี่ย และเกณฑ์รากที่สองของความคลาดเคลื่อนกำลังสองเฉลี่ยที่ต่ำที่สุด ผลการวิจัยพบว่า จากวิธีการพยากรณ์ทั้งหมดที่ได้ศึกษา วิธีที่มีความแม่นยำมากที่สุด คือวิธีการพยากรณ์รวม Yoyram et al. (2020) ได้สร้างแบบจำลองโดยใช้แรนดอมฟอเรสต์ และทำการเปรียบเทียบแบบจำลองโดยทำการคัดเลือกคุณลักษณะของข้อมูล 4 วิธี ได้แก่ IR IG CFS และ SU โดยใช้โปรแกรม WEKA ผลการศึกษาพบว่า โมเดลการพยากรณ์แบบแรนดอมฟอเรสต์ เมื่อ ใช้ร่วมกับเทคนิคการลดคุณลักษณะของข้อมูลแบบ SU ทำให้ประสิทธิภาพในการพยากรณ์เพิ่มขึ้น โดยวัดค่าความถูกต้องได้ค่าร้อยละ 98.74 Wanon et al. (2018) ได้ศึกษาเทคนิคพยากรณ์อาชีพสำหรับนักศึกษาระดับปริญญาตรีโดยใช้เทคนิคเหมืองข้อมูลที่เหมาะสม ทดลองวัดความแม่นยำด้วยเทคนิคการจำแนกข้อมูลด้วยวิธีต้นไม้ตัดสินใจ เทคนิคการจำแนกข้อมูลด้วยวิธีแรนดอมฟอเรสต์ และเทคนิคการจำแนกข้อมูลด้วยวิธีแบ็กกิง ผลการวิจัยพบว่า ความแม่นยำในการจำแนกประเภทข้อมูล จาก 3 เทคนิค 1) เทคนิคการจำแนกข้อมูล

ด้วยวิธีค้นไม่ตัดสินใจเท่ากับ 81.91% 2) เทคนิคการจำแนกข้อมูลด้วยวิธีเรนคอมฟอร์สต์ เท่ากับ 84.29% และ 3) เทคนิคการจำแนกข้อมูลด้วยวิธีแบ็กกิ้ง เท่ากับ 81.71% พบว่าเทคนิคเรนคอมฟอร์สต์ ให้ความถูกต้องในการจำแนกประเภทข้อมูลสูงที่สุด

เพื่อให้ได้ตัวแบบที่เหมาะสมสำหรับการพยากรณ์ราคาข้าวในประเทศไทยและสามารถนำแนวคิดในการหาตัวแบบของการพยากรณ์ไปใช้ในการพยากรณ์ราคาข้าวพันธุ์อื่นได้ เพื่อเป็นแนวทางในการตัดสินใจสำหรับผู้ผลิตและผู้ส่งออกข้าวในประเทศไทย ผู้วิจัยจึงนำข้อมูลราคาข้าวเปลือกหอมมะลิเฉลี่ยรายเดือน ตั้งแต่ มกราคม พ.ศ. 2553 - ธันวาคม พ.ศ. 2563 จำนวน 132 ข้อมูล โดยแบ่งข้อมูลออกเป็น 2 ชุด ข้อมูลชุดที่ 1 ตั้งแต่เดือน มกราคม พ.ศ. 2553 - ธันวาคม พ.ศ. 2562 จำนวน 120 ข้อมูล สำหรับศึกษาตัวแบบพยากรณ์ ข้อมูลชุดที่ 2 ตั้งแต่เดือนมกราคม พ.ศ. 2563 - ธันวาคม พ.ศ. 2563 จำนวน 12 ค่า สำหรับการเปรียบเทียบความแม่นยำของค่าพยากรณ์

จากข้อมูลดังกล่าวพบว่าราคาข้าวเปลือกหอมมะลิเฉลี่ยรายเดือนมีการเปลี่ยนแปลงตามฤดูกาลซึ่งเป็นส่วนประกอบหนึ่งของอนุกรมเวลา ผู้วิจัยจึงใช้วิธีบ็อกซ์-เจนกินส์ ซึ่งเป็นการพยากรณ์โดยอาศัยแนวคิดที่ว่าอนุกรมเวลาที่พิจารณาอาจมีลักษณะของสหสัมพันธ์ในตัวเอง (ACF) และสหสัมพันธ์ในตัวเองบางส่วน (PACF) ซึ่งเป็นวิธีที่นิยมในปัจจุบันแต่การพยากรณ์ในวิธีการดังกล่าวยังพบความคลาดเคลื่อนยังมีค่าสูง ดังนั้นผู้วิจัยจึงได้พยากรณ์ด้วยวิธีแบ็กกิ้ง วิธีเรนคอมฟอร์สต์ เนื่องจากเป็นการวิเคราะห์พยากรณ์ที่เป็นที่นิยมของการทำเหมืองข้อมูล (Data mining) ซึ่งเป็นวิธีเป็นการค้นหาหรือสกัดความรู้จากฐานข้อมูลขนาดใหญ่ เพื่อค้นหารูปแบบ แนวทาง และความสัมพันธ์ที่ซ่อนอยู่ในชุดข้อมูลนั้น โดยอาศัย หลักการทางคณิตศาสตร์ สถิติ การรู้จำ การเรียนรู้ ทางเครื่องจักร เพื่อนำความรู้ ที่ได้นั้นมาใช้ในการ แก้ปัญหา วางแผน หรือดำเนินกลยุทธ์ขององค์กรให้ ประสบความสำเร็จสูงสุด และสามารถดึงข้อมูลมาใช้ โดยการคัดเลือกข้อมูลออกมาใช้งานในส่วนที่ต้องการ

วิธีการทดลอง

1. วิธีบ็อกซ์-เจนกินส์ (Box-Jenkins method)

การพยากรณ์ด้วยวิธีของบ็อกซ์-เจนกินส์เป็นการพยากรณ์เชิงปริมาณวิธีหนึ่งที่มีแนวคิดว่าพฤติกรรมในอดีตของสิ่งที่ต้องการพยากรณ์นั้นเพียงพอที่จะพยากรณ์พฤติกรรมในอนาคตของตัวเองได้ โดยในการพยากรณ์ข้อมูลอนุกรมเวลาด้วยวิธีของบ็อกซ์-เจนกินส์นี้จะแตกต่างจากการพยากรณ์โดยวิธีอื่นซึ่งผู้ที่สร้างตัวแบบพยากรณ์นั้นต้องกำหนดรูปแบบของความสัมพันธ์ก่อนที่จะทำการวิเคราะห์ โดยเฉพาะเมื่ออนุกรมเวลาไม่มีแนวโน้มวัฏจักรหรือฤดูกาลที่ชัดเจน ทำให้ยากในการกำหนดรูปแบบหรือการวิเคราะห์การถดถอยที่เหมาะสมได้ ซึ่งจะต้องทำการกำหนดรูปแบบของความสัมพันธ์ระหว่างตัวแปรอิสระกับตัวแปรตามก่อน แต่วิธีพยากรณ์ของบ็อกซ์-เจนกินส์สามารถแก้ปัญหาดังกล่าวได้ เพราะวิธีพยากรณ์ของบ็อกซ์-เจนกินส์นั้นไม่มีการกำหนดรูปแบบที่ตายตัวขึ้นก่อนทำการวิเคราะห์ โดยในระหว่างการวิเคราะห์รูปแบบจะถูกกำหนดขึ้นมาเองซึ่งสามารถทำตามขั้นตอนของบ็อกซ์-เจนกินส์ได้ดังนี้

1. คำนวณค่าของฟังก์ชันสหสัมพันธ์ในตัวเอง (Autocorrelation Function: ACF) และฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน (Partial Autocorrelation Function: PACF)

2. กำหนดรูปแบบการพยากรณ์ $SARIMA(p,d,q)(P,D,Q)_s$ ที่เหมาะสมโดยพิจารณาจากค่า ACF และ PACF แสดงดังสมการที่ (1) (Jenkins & Reinsel, 1994)

$$f_p(B)F_p(B^s)(1-B)^d(1-B^s)^DY_t = d + q_q(B)Q_q(B^s)e_t \quad (1)$$

โดยที่ Y_t แทนอนุกรมเวลา ณ เวลา t

t แทนช่วงเวลา ซึ่งมีค่าตั้งแต่ 1 ถึง n

n แทนจำนวนข้อมูลในอนุกรมเวลาชุดที่ 1

s แทนจำนวนคาบของฤดูกาล

d และ D แทนลำดับที่ของการหาผลต่าง และผลต่างฤดูกาล ตามลำดับ

B แทนตัวดำเนินการถอยหลัง (Backward Operator) โดยที่ $B^s Y_t = Y_{t-s}$

e_t แทนอนุกรมเวลาของความคลาดเคลื่อนที่มีการแจกแจงปกติและเป็นอิสระกัน

ด้วยค่าเฉลี่ยเท่ากับศูนย์ และความแปรปรวนคงที่ทุกช่วงเวลา

$d = mf_p(B)F_p(B^s)$ แทนค่าคงตัว โดยที่ m แทนค่าเฉลี่ยของอนุกรมเวลาที่คงที่

$f_p(B) = 1 - f_1B - f_2B^2 - \dots - f_pB^p$ แทนตัวดำเนินการสหสัมพันธ์ในตัวเองแบบไม่มีฤดูกาลอันดับที่ p

(Non-Seasonal Autoregressive Operator of Order p : $AR(p)$)

$F_p(B^s) = 1 - F_1B^s - F_2B^{2s} - \dots - F_pB^{Ps}$ แทนตัวดำเนินการสหสัมพันธ์ในตัวเองแบบมีฤดูกาลอันดับที่ P

(Seasonal Autoregressive Operator of Order P : $SAR(P)$)

$q_q(B) = 1 - q_1B - q_2B^2 - \dots - q_pB^p$ แทนตัวดำเนินการเฉลี่ยเคลื่อนที่แบบไม่มีฤดูกาลอันดับที่ q

(Non-Seasonal Moving Average Operator of Order q : $MA(q)$)

$Q_q(B^s) = 1 - Q_1B^s - Q_2B^{2s} - \dots - Q_QB^{Qs}$ แทนตัวดำเนินการเฉลี่ยเคลื่อนที่แบบมีฤดูกาลอันดับที่ Q

(Seasonal Moving Average Operator of Order Q : $SMA(Q)$)

2. วิธีแบ็กกิ้ง (Bagging Method)

วิธีแบ็กกิ้งเป็นเทคนิคย่อยหนึ่งในเทคนิคการรวมกลุ่มตัวจำแนกประเภท (Ensemble Classify) โดยมีวิธีการสุ่มข้อมูลย่อยจากชุดข้อมูลฝึกสอนแบบสุ่มซ้ำได้ (Bootstraps Replacement) จำนวนหลายชุดย่อย นำมาสร้างตัวจำแนก ที่แตกต่างกันแล้วจึงนำผลการทำนายของแต่ละตัวจำแนกย่อยมาพิจารณาาร่วมกัน แล้วสร้างตัวแบบแต่ละตัวที่ใช้ชุดข้อมูลเรียนรู้ที่สุ่มมาจากชุดข้อมูลเรียนรู้ที่กำหนดไว้ในแบบก๊น (Sampling with Replacement) กล่าวคือเมื่อสุ่มข้อมูลขึ้นมาได้หนึ่งตัวจะคืนข้อมูลนั้นกลับเข้าไปยังชุดข้อมูลเดิม ทำให้ข้อมูลนั้นมีโอกาสถูกสุ่มซ้ำอีกในอนาคต

ข้อมูลแต่ละตัวในชุดข้อมูลเรียนรู้ จะมีโอกาสถูกสุ่มเพื่อใช้ในการสร้างตัวแบบเท่ากับ $1 - \frac{1}{n}$ โดยที่ n เป็นจำนวน ข้อมูลในชุดข้อมูลเรียนรู้ที่กำหนดให้ ถ้าทำการสุ่มข้อมูลด้วยวิธีนี้ n ครั้ง เพื่อสร้างชุดข้อมูล สำหรับสร้างตัวแบบ 1 ตัว จะได้ชุดข้อมูลที่เป็นเซตย่อย (subset) ของชุดข้อมูลเรียนรู้ที่กำหนดให้ ทั้งนี้เนื่องจากข้อมูลเดียวกันอาจถูกสุ่มซ้ำ ทำให้ชุดข้อมูลจากการสุ่มมีจำนวนข้อมูลไม่ถึง n ตัว โดยเฉลี่ยแล้วจะมีขนาดประมาณ 63.2% ของ n (Gorad et al., 2017)

3. วิธีเร้นดอมฟอเรสต์ (Random Forest Method)

เทคนิคที่ทำการสุ่มเลือกคุณสมบัติ ออกมาจากชุดข้อมูลหลาย ๆ ชุด จากนั้น นำเอาชุดของคุณสมบัติเหล่านั้นมาสร้างแบบจำลองด้วยเทคนิคต้นไม้ตัดสินใจหลาย ๆ ต้น ลักษณะของต้นไม้ที่อยู่ในป่าของเทคนิควิธีเร้นดอมฟอเรสต์ จะถูกควบคุมด้วย 3 ปัจจัยคือ

1. ต้นไม้แต่ละต้นจะถูกสอน (Train) โดยการใช้เซตย่อยจากข้อมูลตัวอย่าง
2. เมื่อต้นไม้โตขึ้น จะสามารถค้นหาโนด (Node) แต่ละ โหนดที่อยู่ในกิ่งที่ดีที่สุดของต้นไม้โดยใช้การสุ่มเลือกคุณสมบัติจาก N คุณสมบัติ
3. ต้นไม้แต่ละต้นจะไม่มีการตัดออก แต่จะปล่อยให้ต้นไม้โตขึ้นไปเรื่อย ๆ จนได้ผลลัพธ์ที่ดีที่สุดหลังจากการสร้างป่า แล้วทำการให้คะแนนโหนด โดยต้นไม้ภายในป่า หากต้นไม้ต้นใดได้คะแนน โหนดสูงสุดก็จะนำเอาต้นไม้ต้นนั้นออกมาสร้างเป็น โมเดล (Sujatha & Devi, 2013)

4. การวัดความคลาดเคลื่อนของการพยากรณ์ (Forecasting Error)

การวัดความคลาดเคลื่อนของการพยากรณ์จะอาศัยหลักการ การเปรียบเทียบระหว่างค่า พยากรณ์ที่คำนวณได้กับข้อมูลจริงในช่วงเวลา หากค่าพยากรณ์มีค่าคลาดเคลื่อนมากอาจหมายถึงวิธีการที่ใช้พยากรณ์นั้นไม่เหมาะสม หรืออาจจำเป็นต้องเปลี่ยนค่าพารามิเตอร์บางค่าให้เหมาะสม สำหรับงานวิจัยนี้จะใช้วิธีวัดความคลาดเคลื่อนของการพยากรณ์ 2 วิธีคือ

- 1) เกณฑ์รากที่สองของความคลาดเคลื่อนเฉลี่ย (RMSE)

$$RMSE = \sqrt{\frac{\sum_{i=1}^n e_i^2}{n}} \quad (\text{Chai \& Draxler, 2014})$$

- 2) เปอร์เซนต์ความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (MAPE)

$$MAPE = \frac{\sum_{i=1}^n |(e_i / Y_i)|}{n} \times 100 \quad (\text{Isserman, 1977})$$

โดยที่ Y_i แทนค่าของข้อมูลจริงที่เวลา i

F_i แทนค่าพยากรณ์ที่เวลา i

e_i แทนค่าความคลาดเคลื่อนที่เวลา i

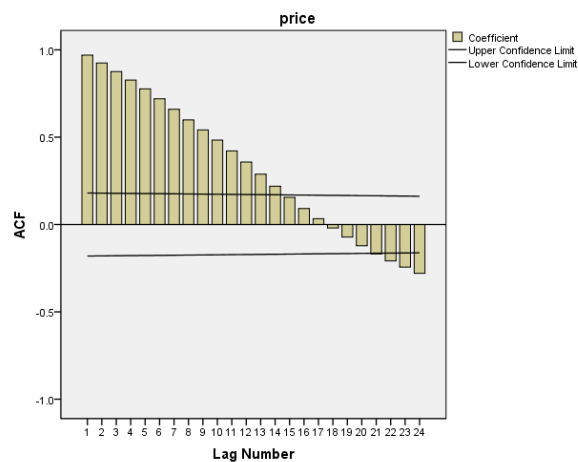
$$e_i = Y_i - F_i$$

ผลและวิจารณ์ผลการทดลอง

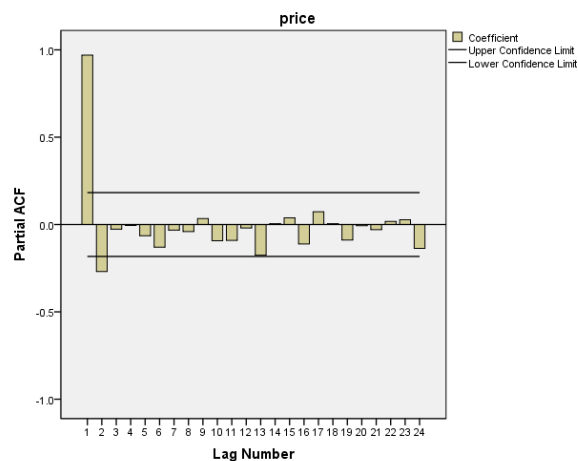
1. วิธีบอกซ์เจนกินส์ (Box-Jenkins method)

นำข้อมูลราคาข้าวเปลือกหอมมะลิรายเดือนไปวิเคราะห์โดยใช้โปรแกรม Statistics Package for Social Sciences: SPSS ตามขั้นตอนต่อไปนี้

ตรวจสอบข้อมูลจากกราฟของฟังก์ชันสหสัมพันธ์ในตัวเอง ACF และการฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน PACF ของราคาข้าวเปลือกหอมมะลิตั้งรูปที่ 1 และรูปที่ 2

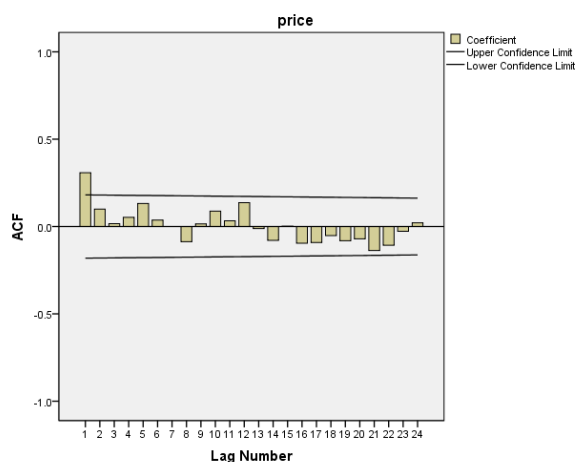


รูปที่ 1 กราฟแสดงค่าฟังก์ชันสหสัมพันธ์ในตัวเอง (ACF)

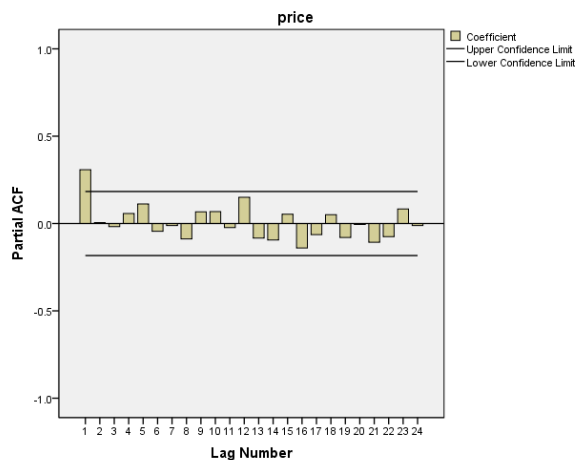


รูปที่ 2 กราฟแสดงค่าฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน (PACF)

จากรูปที่ 1 พบว่าการฟังก์ชันสหสัมพันธ์ในตัวเอง (ACF) ลดลงเข้าสู่ศูนย์อย่างช้า ๆ แสดงว่าอนุกรมเวลาไม่นิ่ง จึงทำการแปลงอนุกรมเวลาโดยการหาผลต่างอันดับที่ 1 ($d = 1$) ของราคาข้าวเปลือกหอมมะลิทั้ง 120 เดือน ได้กราฟ ACF และ PACF ดังรูปที่ 3 และรูปที่ 4



รูปที่ 3 กราฟแสดงค่าฟังก์ชันสหสัมพันธ์ในตัวเอง (ACF) ของผลต่างอันดับที่ 1 ($d = 1$)



รูปที่ 4 กราฟแสดงค่าฟังก์ชันสหสัมพันธ์ในตัวเองบางส่วน (PACF) ของผลต่างอันดับที่ 1 ($d = 1$)

จากรูปที่ 3 พบว่า ACF มีค่าแตกต่างไปจากศูนย์อย่างมีนัยสำคัญ และลดลงเข้าสู่ศูนย์อย่างรวดเร็วแสดงว่าอนุกรมเวลานิ่ง จึงกำหนดตัวแบบพยากรณ์ที่เป็นไปได้เริ่มต้นคือตัวแบบ SARIMA(1,1,0)(1,0,0)₁₂ แบบมีพจน์ค่าคงตัว พร้อมกับประมาณค่าพารามิเตอร์ โดยตัวแบบพยากรณ์มีพารามิเตอร์ทุกตัวมีนัยสำคัญทางสถิติที่ระดับ 0.05 มีค่า

BIC ต่ำที่สุด (BIC = 0.144) และมีค่าสถิติ Ljung-Box Q ไม่มีนัยสำคัญที่ระดับ 0.05 (Ljung-Box Q ณ lag 18 = 12.391 p-value = 0.873) เมื่อตรวจสอบคุณลักษณะของความคลาดเคลื่อนจากการพยากรณ์ที่ระดับนัยสำคัญ 0.05 พบว่าความคลาดเคลื่อนมีการแจกแจงปกติ มีการเคลื่อนไหวเป็นอิสระกัน และจากรูปที่ 5 พบว่าค่าสัมประสิทธิ์สหสัมพันธ์ในตัวเองและสัมประสิทธิ์สหสัมพันธ์ในตัวเองบางส่วนของความคลาดเคลื่อนตกอยู่ในขอบเขตความเชื่อมั่นร้อยละ 95 ดังนั้นตัวแบบดังกล่าวจึงมีความเหมาะสม จากสมการที่ (1) สามารถเขียนเป็นตัวแบบได้ดังนี้

$$(1 - f_1 B)(1 - F_1 B^{12})(1 - B)^4 Y_t = d + e_t$$

แทนค่าประมาณพารามิเตอร์ $d = 43.289$, $f_1 = 0.344$ และ $F_1 = 0.261$ จะได้

$$(-0.089784B^{14} + 0.350784B^{13} - 0.261B^{12} + 0.344B^2 - 1.344B + 1)Y_t = 43.289 + e_t$$

$$Y_t = 43.289 + 0.089784Y_{t-14} - 0.350784Y_{t-13} + 0.261Y_{t-12} - 0.344Y_{t-2} + 1.344Y_{t-1} - 1 + e_t$$

ดังนั้นตัวแบบพยากรณ์คือ

$$\hat{Y}_t = 43.289 + 0.089784Y_{t-14} - 0.350784Y_{t-13} + 0.261Y_{t-12} - 0.344Y_{t-2} + 1.344Y_{t-1} - 1 + e_t$$

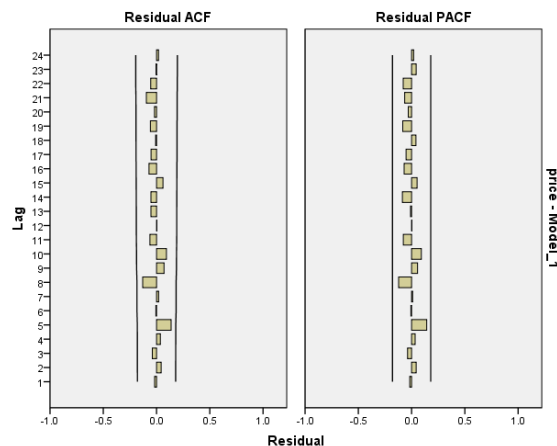
โดยที่ $\hat{Y}_t$ แทนค่าพยากรณ์ ณ เวลา t

Y_{t-j} แทนอนุกรมเวลา ณ เวลา $t - j$

e_t ค่าความคลาดเคลื่อนจากการพยากรณ์ ณ เวลา t

นำสมการดังกล่าวไปพยากรณ์ทำนาราคาข้าวเปลือกหอมมะลิช่วงหน้า 12 เดือน คือ เดือนมกราคม พ.ศ.

2563 - ธันวาคม พ.ศ. 2563 แสดงได้ดังตารางที่ 1



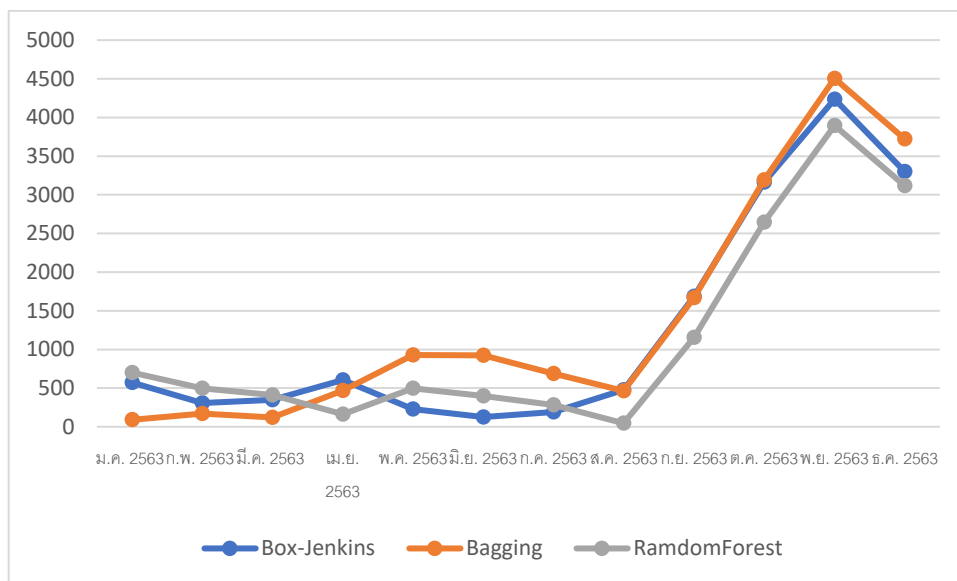
รูปที่ 5 กราฟ ACF และ PACF ของความคลาดเคลื่อนจากการพยากรณ์โดยวิธีบ็อกซ์-เจนกินส์

2. วิธีแบ็กกิ้ง (Bagging method) และ วิธีแรนดอมฟอเรสต์ (Random Forest method)

สำหรับการสร้างแบบจำลองโดยวิธีแบ็กกิ้งและวิธีแรนดอมฟอเรสต์นั้น ได้ใช้โปรแกรม Weka ในการคำนวณ โดยวิธีแบ็กกิ้ง ผู้วิจัยได้กำหนดค่าพารามิเตอร์ ดังนี้ Sample ratio มีค่าเท่ากับ 0.9, Iteration มีค่าเท่ากับ 7 และ Average Confidences มีค่าเป็น True และวิธีแรนดอมฟอเรสต์ ผู้วิจัยได้กำหนดค่าพารามิเตอร์ ดังนี้ Number of tree มีค่าเท่ากับ 50, Criterion มีค่าเป็น Information_gain, Maximal depth มีค่าเท่ากับ 9, Apply pruning มีค่าเป็น True และค่า Confidence มีค่าเท่ากับ 0.25 ซึ่งได้ผลค่าพยากรณ์ตามตารางที่ 1

ตารางที่ 1 ค่าจริงและค่าพยากรณ์ราคาข้าวเปลือกหอมมะลิ(บาท/เมตริกตัน) ตั้งแต่เดือนมกราคม พ.ศ. 2563 – เดือนธันวาคม พ.ศ. 2563

เดือน	ค่าจริง	ค่าพยากรณ์ล่วงหน้า		
		Box-Jenkins	Bagging	Ramdom Forest
มกราคม พ.ศ. 2563	13,756	13183	13846.38	14455.71
กุมภาพันธ์ พ.ศ. 2563	13,913	13607	14084.90	14408.86
มีนาคม พ.ศ. 2563	13,949	14298	14068.16	14357.77
เมษายน พ.ศ. 2563	14,471	13864	14003.06	14306.95
พฤษภาคม พ.ศ. 2563	14,871	14643	13942.96	14371.83
มิถุนายน พ.ศ. 2563	14,868	14994	13942.96	14471.17
กรกฎาคม พ.ศ. 2563	14,630	14820	13942.96	14349.66
สิงหาคม พ.ศ. 2563	14,409	14885	13942.96	14363.21
กันยายน พ.ศ. 2563	13,237	14921	14907.42	14393.36
ตุลาคม พ.ศ. 2563	11,715	14876	14907.42	14362.61
พฤศจิกายน พ.ศ. 2563	10,403	14640	14907.42	14298.41
ธันวาคม พ.ศ. 2563	11,185	14485	14907.42	14303.89
ความคลาดเคลื่อน	RMSE	1,891.05	2,037.26	1,699.64
ของการพยากรณ์	MAPE	10.83	11.90	9.80



รูปที่ 6 กราฟเปรียบเทียบค่าความคลาดเคลื่อน(บาท/เมตรกตัน) ของค่าพยากรณ์

จากผลการทำนายราคาข้าวเปลือกหอมมะลิ ผลการทดลองพบว่า วิธีบอกซ์เจนกินส์, วิธีแบ็กกิ้ง และ วิธี แรนดอมฟอเรสต์ ได้ค่า RMSE เป็น 1,891.05, 2,037.26 และ 1,699.64 ตามลำดับ และได้ค่า MAPE เป็น 10.83, 11.90 และ 9.80 ตามลำดับ ซึ่งค่าที่ได้มีความใกล้เคียงกัน แต่วิธีแรนดอมฟอเรสต์ เป็นวิธีที่ได้ค่าทั้งสองต่ำที่สุด เมื่อเทียบกับ วิธีบอกซ์เจนกินส์และวิธีแบ็กกิ้ง ดังรูปที่ 6

ผลการวิจัยนี้สอดคล้องกับงานวิจัยของ Poonsawad & Sanrach (2019) ซึ่งศึกษาเทคนิคพยากรณ์การได้รับ ปัจจัยพื้นฐานนักเรียนยากจนของนักเรียนโรงเรียนวัดพระขาว (ประชานุเคราะห์) ด้วยเทคนิคเหมืองข้อมูล โดยใช้ เทคนิคการจำแนกข้อมูล 3 เทคนิค ได้แก่ ต้นไม้ตัดสินใจ แรนดอมฟอเรสต์ และเทคนิคแบ็กกิ้ง ผลการวิจัย พบว่า เทคนิคการจำแนกข้อมูลด้วยเทคนิคแรนดอมฟอเรสต์ให้ค่าประสิทธิภาพมากที่สุด และสอดคล้องกับงานวิจัยของ Intarat & Sillaparat (2019) ได้ศึกษาการจำแนกพันธุ์ไม้ป่าชายเลนเขตร้อนด้วยวิธีแรนดอมฟอเรสต์ และข้อมูลภาพถ่ายจากดาวเทียมรายละเอียดสูง ผลการวิจัยพบว่า การใช้วิธีแรนดอมฟอเรสต์ ให้ประสิทธิภาพในการจำแนกที่สูง กว่าวิธีการจำแนกเชิงคุณภาพแบบความน่าจะเป็นสูงสุด โดยมีค่าความถูกต้องโดยรวมที่ร้อยละ 78.00 และค่าสถิติ แคลปปาที่ 0.72 ในขณะที่ผลการจำแนกด้วยความน่าจะเป็นสูงสุดให้ค่าความถูกต้องโดยรวมที่ร้อยละ 56.00 และ ค่าสถิติแคลปปาที่ 0.44 ผลจากการทดสอบค่าสถิติ Z ($Z = 3.68$) ช่วยยืนยันถึงความแตกต่างระหว่างทั้งสองวิธีการ จำแนก อย่างมีนัยสำคัญที่ค่าความเชื่อมั่นร้อยละ 95.00

สรุปผลการทดลอง

แบบจำลองทางคณิตศาสตร์ทำนายราคาข้าวเปลือกหอมมะลิ โดยใช้ตัวแบบพยากรณ์ วิธีบอกซ์เจนกินส์ วิธีเบ็กกิ้ง และ วิธีเรแนคอมโพเรสต์พบว่า วิธีเรแนคอมโพเรสต์ได้ทำนายราคาข้าวเปลือกหอมมะลิได้ใกล้เคียงกับค่าจริงที่สุด และเมื่อคำนวณค่า RMSE ได้ 1,699.64 และได้ค่า MAPE เป็น 9.80

เอกสารอ้างอิง

- Chai, T., & Draxler, R. R. (2014). Root mean square error (RMSE) or mean absolute error (MAE)? –Arguments against avoiding RMSE in the literature. *Geoscientific model development*, 7(3), 1247-1250.
- Gorad, N., Zalte, I., Nandi, A., & Nayak, D. (2017). Career counselling using data mining. *International Journal of Engineering Science and Computing*, 7(4), 10271-10274.
- Intarat, K., & Sillaparat, S. (2019). Tropical mangrove species classification using random forest algorithm and very high-resolution satellite imagery. *Burapha Science Journal*, 24(2), 742-753. (in Thai)
- Isserman, A. M. (1977). The accuracy of population projections for subcounty areas. *Journal of the American Institute for Planners*, 43(3), 247-259.
- Jenkins, G. M., & Reinsel, G. C. (1994). *Time series analysis; forecasting and control* (3rd ed.). Prentice Hall, Englewood Cliff, New Jersey.
- Kodosueb, S., & Boonlha, K. (2016). Construction of model for the price of Thai jasmine rice 105. *Science and technology Nakhon Sawan Rajabhat University Journal*, 8(8), 49-60. (in Thai)
- Poonsawad, A., & Sanrach, C. (2019). A study of techniques in predicting of a fund receiving for poor students of Phrakhao school students by using data mining technique. *The Sci of Phetchaburi Rajabhat University*, 16(2), 1-10. (in Thai)
- Riansu, W. (2016). Forecasting the prices of paddy rice at 15% moisture content. *The Sci J of Phetchaburi Rajabhat University*, 13(1), 14-25. (in Thai)
- Sujatha, M., & Devi, L. (2013). Feature selection techniques using for high dimensional data in machine learning. *International Journal of Engineering Research & Technology*, 2(9), 97-102.
- Thai Rice Exporters Association. (2020). the 53rd Annual General Assembly Meeting for the Year 2020. (in Thai)
- Wanon, S., Arreerard, T., & Sanrach, C. (2018). A study of techniques in predicting career counseling for undergraduate students of the computer program by using data mining technique. *Journal of Innovation Technology Management*, 5(1), 164-171. (in Thai)

Yoyram, J., Suntaralak, S., & Somchai, K. (2020). Analysis of consumer factors towards the purchase of silk pupa in Ban Hua Saphan, Ban Yang, Phutthaisong District, Buriram Province with data mining techniques. *The 6th National Conference on Technology and Innovation Management NCTIM 2020* (pp. 1-7). Rajabhat Maha Sarakham University, Maha Sarakham. (in Thai)