



JOURNAL OF INFORMATION SCIENCE AND TECHNOLOGY (JIST)

ISSN 2651-1053 (Online), ISSN 1906-9553 (Print)

The Journal of Information Science and Technology (JIST) is an academic journal established by the collaboration of 21 faculties that conduct courses related to Information Technology, namely, Bangkok University, Dhurakij Pundit University, King Mongkut's Institute of Technology Ladkrabang, King Mongkut's University of Technology North Bangkok, King Mongkut's University of Technology Thonburi, Mae Fah Luang University, Mahanakorn University of Technology, Mahasarakham University, Mahidol University, Nakhon Phanom University, Panyapiwat Institute of Management, Prince of Songkla University, Rangsit University, Siam University, Silpakorn University, Sripatum University, Thai-Nichi Institute of Technology, Walailak University, Burapha University, Phayao University and Ubon Ratchathani Rajabhat University. According to the agreement of deans of all faculties (Council of IT Deans of Thailand (CITT)), Mahanakorn University of Technology was appointed as a coordinator of the journal.

The journal was established in 2010 and plans to publish 2 issues per year (JAN – JUNE and JULY – DECEMBER per year). The journal was established first print journal publication in 2010 (Vol 1. No.1) with ISSN 1906-9553 (Print) and plans to publish 2 issues per year on during January - June and July - December. Also the journal was established first online journal publication in 2010 (Vol 1. No.1) with ISSN 2651-1053 (Online) and plans to publish 2 issues per year on during January - June and July - December.

EDITOR IN CHIEF

Werasak Kurutach, Mahanakorn University of Technology, Thailand

ADVISORY BOARD

Ruttikorn Varakulsiripunth Thai-Nichi Institute of Technology, Thailand	Krisana Chinnasarn Burapha University, Thailand	Sunantha Sodsee King Mongkut's University of Technology North Bangkok, Thailand
Pisit Charnkeitkong Panyapiwat Institute of Management, Thailand	Panavy Pookaiyaudom Mahanakorn University of Technology, Thailand	Poonpong Boonbrahm Walailak University, Thailand
Pattanasak Mongkolwat Mahidol University, Thailand	Sasitorn Kaewman Mahasarakham University, Thailand	Chetneti Srisaan Rangsit University, Thailand
Siridech Boonsang King Mongkut's Institute of Technology Ladkrabang, Thailand	Thirapon Wongsardsakul Bangkok University, Thailand	Teeravisit Laohapensaeng Mae Fah Luang University, Thailand
Thana Sukvaree Sripatum University, Thailand	Dechanuchit Katanyutaveetip Siam University, Thailand	Chaiyaporn Khemapatapan Dhurakij Pundit University, Thailand
Nathaporn Karnjnapoomi Silpakorn University, Thailand	Sinchai Kamolphiwong Prince Of Songkla University, Thailand	Anong Rungsuk Nakhon Phanom University, Thailand
Pornthep Rojanavas University of Phayao, Thailand	Kriengkrai Porkaew King Mongkut's University of Technology Thonburi, Thailand	Supawee Makdee Ubon Ratchathani Rajabhat University, Thailand

JOURNAL OF INFORMATION SCIENCE AND TECHNOLOGY (JIST)

ISSN 2651-1053 (Online), ISSN 1906-9553 (Print)

EDITORIAL BOARD

Werasak Kurutach
Mahanakorn University of
Technology, Thailand

Woraphon Lilakiatsakun
Mahanakorn University of
Technology, Thailand

Surakarn Duangphasuk
Mahanakorn University of
Technology, Thailand

Rattana Wetprasit
Prince Of Songkla University,
Thailand

Datchakorn Tancharoen
Panyapiwat Institute of Management,
Thailand

Wanvipa Wongvilaisakul
Panyapiwat Institute of
Management, Thailand

Pawitra Chiravirakul
Mahidol University, Thailand

Songsri Tangsripairoj
Mahidol University, Thailand

Suppat Rungraungsilp
Walailak University,
Thailand

Sarayut Nonsiri
Thai-Nichi Institute of
Technology, Thailand

Annop Monsakul
Thai-Nichi Institute of
Technology, Thailand

Sasitorn Kaewman
Mahasarakham University,
Thailand

Nutchanat Buasri
Mahasarakham University,
Thailand

Sayan Riwtong
Mahanakorn University of
Technology, Thailand

Kaboon Thongtha
Mahanakorn University of
Technology, Thailand

Mudarmeen Munlin
Mahanakorn University of
Technology, Thailand

Sureeluk Weerajong
Mahanakorn University of Technology,
Thailand

Benjamas Meansamut
Mahanakorn University of
Technology, Thailand

Chatchai Tubgri
Mahanakorn University of
Technology, Thailand

MANAGER

Kaboon Thongtha, Mahanakorn University of Technology, Thailand

CONTACT ADDRESS

Council of IT Deans of Thailand (CITT),
140 Moo 1, Cheum-Sampan Road, Nongchok,
Bangkok, Thailand 10530
Tel: +(662) 988-3655 ext 4115

CALL FOR PAPER

Journal of Information Science and Technology (JIST)

<https://ph02.tci-thaijo.org/index.php/JIST>

Editor in Chief

Werasak Kurutach (MUT)

Advisory Board

Panavy Pookaiyaudom (MUT)

Kriengkrai Porkaew (KMUTT)

Sasitorn Kaewman (MSU)

Pattanasak Mongkolwat (MU)

Poonpong Boonbrahm (WU)

Ruttikorn Varakulsiripunth (TNI)

Pisit Charnkeitkong (PIM)

Chetnati Srisaon (RSU)

Siridech Boonsang (KMUTL)

Thirapon wongsaardsakul (BU)

Sunantha Sodsee (KMUTNB)

Teeravisit Laohapensaeng (MFU)

Thana Sukvaree (SPU)

Dechanuchit Katanyutaveetip (SU)

Chaiyaporn Khemapatapan (DPU)

Nathaporn Kamjapoomi (SU)

Sinchai Kamolphiwong (PSU)

Anong Rungsuk (NPU)

Krisana Chinasarn (BUU)

Pomthep Rojanavasu (UP)

Supawee Makdee (UBRU)

Editorial Board

Mudameen Munlin (MUT)

Sureeluk Weerajong (MUT)

Surakarn Duangphasuk (MUT)

Rattana Wetprasit (PSU)

Datchakorn Tancharoen (PIM)

Wanvipa Wongvilaisakul (PIM)

Pawitral Chiravirakul (MU)

Songsri Tangsripairoj (MU)

Suppat Rungraungsilp (WU)

Sarayut Nonsiri (TNI)

Annop Monsakul (TNI)

Sasitorn Kaewman (MSU)

Nuchanat Buasri (MSU)

Sayan Riwtong (MUT)

Benjamas Meansamut (MUT)

Chatchai Tubgri (MUT)

Manager

Kaboon Thonghta (MUT)

Submissions

Online Journal Submission

<https://ph02.tci-thaijo.org/index.php/JIST/>

Publication Periodicity

Published 2 times a year in

June and December.

(JAN – JUN and JULY – DEC).

(Indexed by TCI tier 2)

JIST Office Address

Council of IT Deans of Thailand

(CITT), Faculty of Engineering and

Technology, Mahanakorn

University of Technology 140 Moo

1, Cheum-Sampan Rd., Nongchok,

Bangkok, Thailand 10530

Tel: +(662)988-3655 ext 4115,

Fax: +(662)988-4027

Peer Review Process

All submitted manuscripts must be reviewed by at least three expert reviewers via the double-blinded review system.

Manuscript Preparation Guide (Publication Charge: Free)

The template for writing the journal is available for downloading (jist_template_eng-2020.doc or jist_template_thai-2020.doc). Minimum length of the paper should be in between 8-12 pages of A4 size. Text should be typewritten or printed with double-spaced in 11-point of A4 white paper with margins of 1.5" for top, 1.25" for left, and 1" for bottom and right sides. All pages must be numbered sequentially. Here are some guidelines.

- Abstract with length 100-200 words in length.
- Introduction which discusses existing problem and related research
- Materials and methods used for experimentation, Results of the experiment, Discussion and conclusion, References Style is IEEE.

The **Journal of Information Science and Technology (JIST)** aims to be the forum through which researchers, faculties, and experts of the computer and information technology and others technological related fields share and discuss their high quality research work as well as innovation. Original research articles, practical applications and innovations in the broad area of computer and information technology are suitable for publication in JIST.

• **Soft Computing:**

Artificial Intelligence, Fuzzy and Neural Computing, Genetic Algorithms, Intelligent Agents, Neural Networks, Natural Language Processing, Robotics and Automation, Image Processing

• **Human-Computer Interaction:**

Graphical User Interfaces, Multimedia, Interactive Systems, Computer-Supported Cooperative Work, Human Cognitive Skills, Visualization, and Computer Simulation

• **Information Assurance and Security:**

Cryptography, Forensics and Incident Response, Biometrics, Security Policies and Procedures

• **Information Systems:**

Databases, Information Retrieval, Transaction Processing, Distributed and Object Databases, Data Warehousing, Knowledge Management, Expert Systems, Multimedia Information Systems, Digital Libraries, Geographical Information Systems

• **Networking:**

Advanced Computer Networks, Internet Technology and Applications, Distributed Systems, Wireless and Mobile Computing, Data Compression, Network Security, Enterprise Networking, Digital Communications

• **Programming:**

Object-Oriented Programming, Event-Driven Programming, Functional Programming, Logic Programming, XML and other Extensible Languages, Parallel Programming

• **Platform Technologies:**

Advanced Computer Architecture, Parallel Architectures, Cluster / Grid Computing, RFID, Embedded Systems

• **System Integration and Architecture:**

Software Acquisition and Implementation, System Needs Assessment, Software Economics, Enterprise Systems, Computing Economics

• **Social and Professional Issues:**

Professional Practice, Social Context of Computing, Computers Ethics, IT Law, Intellectual Property, Privacy and Civil Rights

• **Web Systems and Technologies:**

Programming for the WWW, E-commerce, E-Learning, Content Management Systems, E-Government and E-Service

• **Multidisciplinary:**

Bioinformatics, Biomedical Information Systems, Nano Technology

• **e-Business and m-Business:**

e-Business Process Aspects, Intelligent e-Business System, m-Business, e-Supply Chain Management & e-Logistics, e-Customer Relationship Management, e-Negotiation, e-Payment, e-Business Design and Developments

• **Business and Information System:**

Information management for business applications, Enterprise systems and architecture, Business systems infrastructure design for information integration, Service-oriented architecture, Service-component architecture

• **Social and Business Aspects of Convergence IT and Ubiquitous Computing:**

Economics of emerging technologies, Home and community computing, Economic and social complexities in CIT and UC, New forms of media and communications

• **Internet of Things (IoT):**

IoT system architecture, IoT enabling technologies, IoT communication and networking protocols such as network coding, and IoT services and applications and technology development in different standard development organizations (SDO).

• **Cloud Computing:**

Auditing, monitoring, scheduling, resource registration/discovery, Automatic reconfiguration, self healing/monitoring, Service level agreements, integration, management.

• **Fintech and blockchain:**

Development of Blockchain Security Enhancement Technologies, Platform Technology that Supports Safe Transactions, Financial and Economic Structures, Blockchain's Impacts on Workstyles, Fast Fintech for Smartphone Application Service Platform.

Prevention of Cryptocurrency Loss from Wrong Receiver Address

Nanta Janpitak^{1*} and Chanboon Sathitwiriawong²

¹ Faculty of Engineering at Sriracha, Kasetsart University Sri Racha Campus

² Faculty of Information Technology, King Mongkut's Institute of Technology Ladkrabang

Received: March 19, 2021; Revised: June 04, 2021; Accepted: July 21, 2021; Published: July 27, 2021

ABSTRACT – Digital money such as Bitcoin, Ethereum and Litecoin are blockchain-based cryptocurrencies. In today's financial age, many people are holding and trading cryptocurrencies instead of fiat money or other assets. To transfer money in the centralized banking application, the valid receiver will be checked before enable transaction to be process. To transfer cryptocurrency in blockchain network, there are steps to verify the transactions by the blockchain nodes instead of the central authority. However, most of the verification processes are based on sender's properties such as sender's digital signature and sender's account balance compared to the sending amount, etc. There is no process to verify whether the receiver address is valid or not. From this vulnerability, there is a risk that the sender may input the wrong receiver address. If this mistake happens and the transaction has been confirmed, the sender will lose his cryptocurrency without recovery option. The amount will be transfer to that receiving address, but we cannot know that it is belong to whom or may not belong to anybody. There are many topics in cryptocurrency forum and news that the owner cried about losing of their cryptocurrency. Many of them had mistaken input in receiver address. Up until now, there is no mechanism to protect this mistake. In this paper, we propose a process to verify the receiver address using the blockchain API before constructing and submitting the transaction in blockchain application. If the receiver address does not exist in the same blockchain network, there is a process to notify the sender. This process can prevent the loss of cryptocurrency not only from wrong receiver addresses but also from different network addresses.

KEYWORDS: Blockchain, Bitcoin, Ethereum, Cryptocurrency, API, Cryptocurrency Address

*Corresponding Author: njanpita@hotmail.com

1. Introduction

A blockchain [1] technology was invented by Satoshi Nakamoto in October 2008 via paper titled “Bitcoin: A Peer-to-Peer Electronic Cash System. Blockchain is running based on a peer-to-peer network in which transactions are sent from one party to another without a centralized administration. This characteristic of blockchain can improve system availability because the system does not depend on any centralized server. Every full node that is online can serve as a server. The node can leave and come back to join anytime at will. The blockchain will record the hash of transactions as a reference so that it is not possible to edit any transaction and make the same hash to the original transaction. This characteristic of blockchain can ensure that the transactions on blockchain are immutable records.

Bitcoin is the first cryptocurrency that was implemented successfully using blockchain technology. By using blockchain technology, Bitcoin can be sent from one account to other accounts without mint or central authority. Satoshi has defined the transactions of Bitcoin as a chain of digital signatures. The hash of the previous transaction and the public key of the receiver is signed by the sender’s private key and attached to the end of the coin prior to the transfer of the coin. A blockchain node can verify the signatures to verify the chain of ownership.

Ethereum [2], Litecoin [3], Monero [4], and Dash [5] are other cryptocurrencies that are implemented based on blockchain. These cryptocurrencies were called altcoin because they are alternate to Bitcoin. Bitcoin and altcoin are increasingly popular because users have believed in their several advantages over fiat money such as low fees and irreversible transactions.

The next section will explain more details about blockchain node and verification processes.

2. Blockchain Node and Verification Processes

According to the definition of a node in [6], there are three main types of blockchain nodes as follows.

1. Full-node client: A full client, or “full node”, is a client that stores the entire history of Bitcoin transactions, manages users’ wallets, and can initiate transactions directly on the Bitcoin network. Full nodes act as a server in a decentralized network. They handle all aspects of the protocol and can independently validate the entire blockchain and any transaction. A full-node client consumes substantial computer resources (e.g., more than 125 GB of disk, 2 GB of RAM) but offers complete autonomy and independent transaction verification.

2. Lightweight client: A lightweight client, also known as a simple-payment-verification (SPV) client. These types of nodes communicate with the blockchain while relying on full nodes to provide them with the necessary information. As they don’t store a copy of the chain, they only query the current status for which block is last, and broadcast transactions for processing.

3. Third-party API client: A third-party API client is one that interacts with Bitcoin through a third-party system of application programming interfaces (APIs), rather than by connecting to the Bitcoin network directly. The wallet may be stored by the user or by third-party servers, but all transactions go through a third party.

There are four main verification processes (or consensus processes) before a single raw transaction can be added successfully in a blockchain network as follows.

1. Every blockchain validating node that receives a transaction (in this state, called “unconfirmed” transaction) will independently verify the transaction to ensure that only valid transactions are propagated across the network.

2. Blockchain mining node will aggregate transactions into a candidate block and then solve the Proof-of-Work solution in order to get a new mining block.

3. Every blockchain validating node that receives the new block will verify the new block to ensure that only valid blocks are propagated across the network.

4. The final step in blockchain’s decentralized consensus mechanism is the assembly of blocks into chains and the selection of the chain with the most Proof-of-Work.

From a long checklist of verification criteria such as:

- The transaction’s syntax and data structure must be correct.
- The transaction size in bytes is less than MAX_BLOCK_SIZE.
- For each input, the referenced output must exist and cannot already be spent.
- The block data structure is syntactically valid.
- The block header hash is less than the target (enforces the Proof-of-Work).
- The block timestamp is less than two hours in the future (allowing for time errors).
- The block size is within acceptable limits.
- Etc.

We found that none of the verification processes care about the validity of the receiver address. We then tested the transaction in the cryptocurrency wallet application and found that if the sender inputs the invalid receiver address (does not exist in the network), the cryptocurrency balance, known as

unspent transaction outputs or UTXO will be decreased from the sender address and increased in the invalid receiver address. However, no one can claim to be the owner of those balance because no one has the private key of that invalid address. Therefore, this situation is considered as the loss of cryptocurrency without recovery option.

To demonstrate our test result, we use the Ropsten [7] which is a test network for Ethereum developer to develop and to do testing. We use the Metamask [8] as Ethereum wallet to connect to the Ropsten test network and the main network.

Our scenario is that Alice bought a 1 ETH watch from Bob's watch shop. The invalid receiver address can happen in many cases as the following:

1. Bob informs Alice an invalid address.
2. Bob's wallet is connected to a different network, then the valid address of the different network is informed to Alice.
3. Alice inputs the wrong address by herself.
4. Alice creates a fake evidence to deceive Bob that she already sent him 1 ETH using the Ropsten network and then gone with the watch.

The demonstration of this scenario is shown and explained in Figure 1-7.

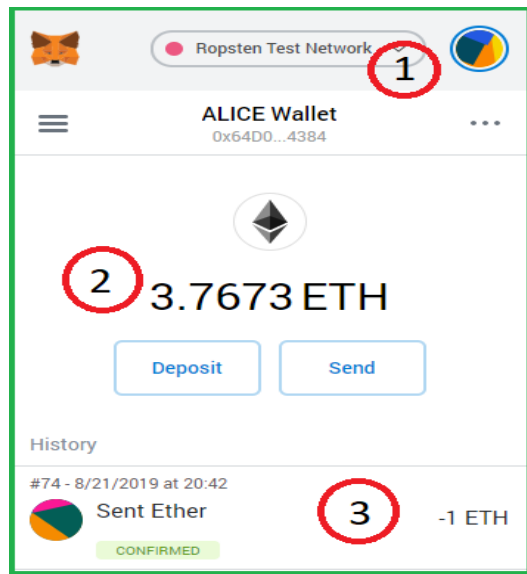


Figure 1. Alice's Ethereum Wallet on Ropsten Test Network using Metamask.

Figure 1 is Alice's Ethereum Wallet that has the following details:

1. Metamask wallet is connected to Ropsten test network.
2. Current balance or UTXO is 3.7673 ETH.
3. History of Alice transactions.

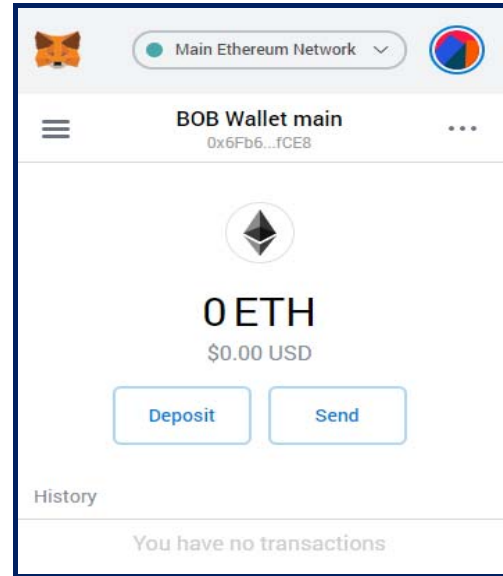


Figure 2. Bob's Ethereum Wallet on main network using Metamask.

Figure 2 shows Bob's Ethereum Wallet that has the following details:

1. Metamask wallet is connected to the main network.
2. Current balance or UTXO is 0 ETH.

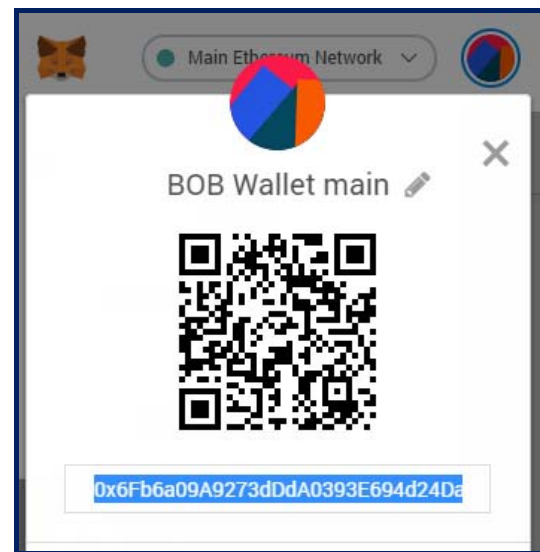


Figure 3. Bob's Ethereum address.

Figure 3 shows Bob's address to which Alice must pay. It is "0x6Fb6a09A9273dDdA0393E694d24Da9B28981fCE8".

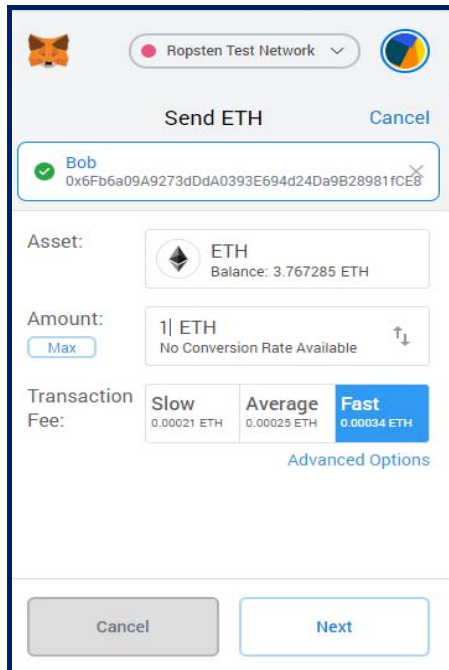


Figure 4. Alice is going to send 1 ETH to Bob's address.

Figure 4 shows the screen that Alice input 1 ETH and then she must push "Next" button for the next process.

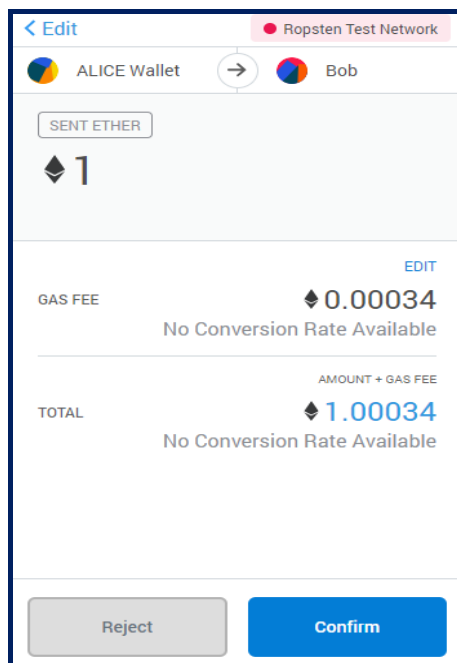


Figure 5. Alice is going to send 1 ETH to Bob's address.

Figure 5 is the screen that shows the transfer of information to Alice and waiting for her to push "Confirm" button to proceed to the next process.

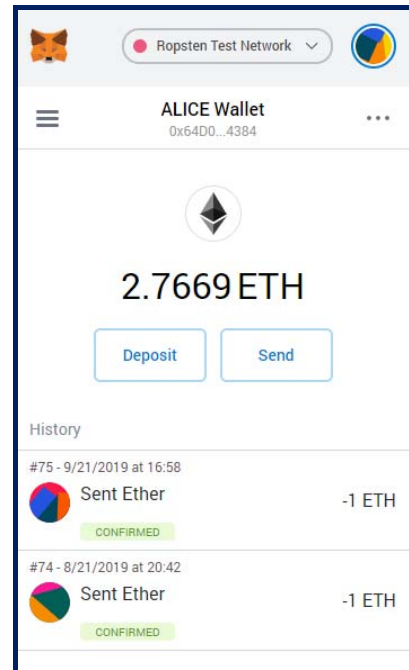


Figure 6. Alice has successfully sent 1 ETH to Bob's address.

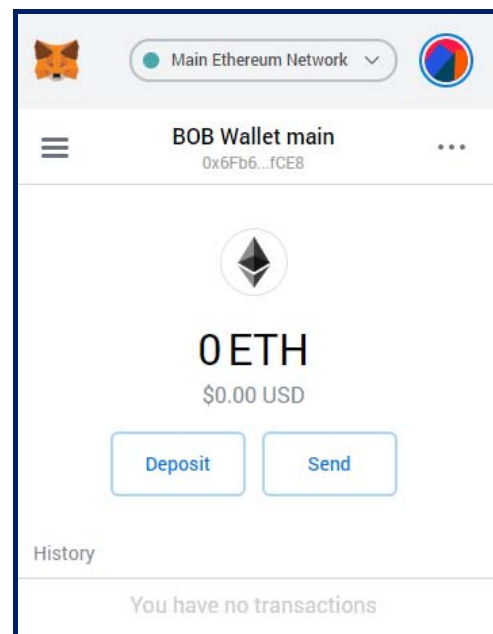


Figure 7. Bob's Ethereum Wallet after Alice has sent 1 ETH.

Figure 6 shows that the balance of Alice wallet has decreased from 3.7673 ETH to 2.7669 ETH while the balance of Bob in Figure 7 has not increased anyhow because the address is in a different network.

From this scenario, there is the cases that may happen,

1. Alice has attention to cheat, and she made the fake evidence shows to Bob that she

already paid. Bob may believe because he saw the figure 5. Alice got the watch for free. Bob loss his watch.

2. Alice has no attention to cheat, and she paid to Bob's valid address but difference network. If Bob waits for the transaction to be completed before giving away the watch, Alice must pay again. So, Alice loss her money.

To prevent a similar problem, we propose a process to verify the receiver address using the blockchain API before constructing and submitting the transaction across the network. If the receiver address does not exist in the same blockchain network, there is the process to notify the sender. This process can prevent the loss of cryptocurrency not only from invalid receiver addresses but also from different network addresses as well. This process needs to be developed in a wallet application.

3. Related Research

The anonymity is one core property of Bitcoin and other alternate coins. The purpose of anonymity is preserving the privacy of users by hiding user's identity from user's transactions. However, anonymity may affect the process of sending Bitcoin related to the target receivers addresses as mentioned before.

Many researchers have identified several problems associated with the privacy and anonymity of Bitcoin and then provide the solutions to address these issues. Niluka Amarasinghe, Xavier Boyen and Matthew McKague [9] conducted a survey of those solutions and concluded that many solutions do not provide an acceptable level of anonymity.

Mauro Conti et al. [10] conducted a survey on security and privacy issues of Bitcoin and one of the issues is anonymity. Even anonymity may affect some processes, but it is still required in terms of cryptocurrency properties. So, it can be said that anonymity needs to be improved but not removed.

QingChun ShenTu and JianPing Yu [11] also studied anonymization and de-anonymization technologies. They discussed some anonymization methods such as Analysis of Transaction Chain (ATC), coin-mixing and transaction remote release (TRR) that were used to cover the relationship between Bitcoin address and the user. Finally, they made some concluded that:

- ATC could not reach the practical stage. Many stolen Bitcoins failed to be identified its owner.
- The security and practicability of the decentralized coin-mixing protocols have not been estimated adequately. More de-anonymization research is expected to attack decentralized coin-mixing protocols.

- Group signature, group blind signature, privacy sharing, homomorphic encryption, lattice cryptography and other algorithms should be applied on Bitcoin system.

Liu et al. [12] discussed that the transactions can be attacked after submitting from wallet by interception, modification, and rebroadcast of a transaction into the Bitcoin network. They proposed the mechanism to recheck the balance of Bitcoin address after spending to confirm the completion of a transaction in case of transaction malleability. However, this mechanism is for mistake detection but not for mistake prevention.

Sai, Buckley, and Le Gear [13] discussed the privacy and security of cryptocurrency mobile applications by referring to the OWASP mobile top 10. They performed the test and investigation on wallet application vulnerabilities but none of the result discussed about validation of receiver address.

In the Bitcoin talk forum [14], there is the discussion header "If I send my Bitcoins to a wrong address, what happens?". Until today, the most discussion still saying that the Bitcoin will be lost forever if it has been sent to the wrong address.

Coinbase Q&A [15] also mentioned that "Due to the irreversible nature of digital currency protocols, transactions can neither be canceled nor reversed once sent. In this scenario, it would be necessary to contact the receiving party and seek their cooperation in returning the funds. If you do not know the owner of the address, there are no possible actions you can take to retrieve the funds."

In Quora's discussion [16] header "What happens if I send my Bitcoin to an incorrect address?" also mention that the cross-chain deposits are rare but possible; many people have lost their money by sending Bitcoin to a Bitcoin cash or sending NEO to Bitcoin address for example.

4. Proposed Receiver Address Verification Process

Since most of the blockchain verification processes will take time around 10 minutes, so it does not make sense to reject the transaction in the later steps if invalid receiver address is found. It will incur more workload to the blockchain network. The invalid receiver address should be tracked at the earliest stage possible. From this point, we proposed to verify the receiver address before constructing and submitting the transaction to the network. This can be better done by the wallet application. However, most of the wallet applications especially mobile wallets are implemented based on lite nodes that do not have the full data of blockchain.

So, we would propose to include a verification script into the wallet application by using the blockchain API from the third party to verify the receiver address. If the receiver address is not found in the

block explorer, the wallet application should notify user to re-input receiver address before proceeding to the next process. User can decide to continue or re-input receiver address at will (just in case the wallet application is in the test environment).

Since we cannot modify the worldwide wallet application such as Metamask, so we use our own

developed wallet application developed by the Geth [17] command and Web3.js [18] interface which are used by most Ethereum wallet applications in the real market. The process demonstration can be explained by Figure 8-11.

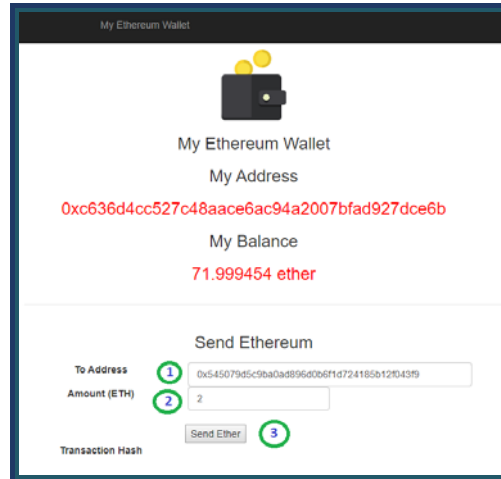


Figure 8. Simple Ethereum Wallet developed by Web3.

```
// Send ether
function sendEther(){
    var to = document.getElementById('send_to').value;
    var amount = document.getElementById('send_amount').value;

    var params = {
        from: web3.eth.coinbase,
        to: to,
        value: web3.toWei(amount, "ether"),
    };

    web3.eth.sendTransaction(params, function (error, result) {
        if(!error){
            document.getElementById('send_output').innerHTML = result;
        }else{
            document.getElementById('send_output').innerHTML = error;
        }
    });
}
```

Figure 9. JavaScript with Geth command to send Ether.

Figure 8 shows a simple Ethereum Wallet application. Once the user inputs the receiver address (1) and the amount (2) and clicks button “Send Ether” (3), the script is called to send Ether using the Geth command “web3.eth.sendTransaction()” as shown in Figure 9.

Command sendTransaction() required three important parameters which are: sender address (Figure 8: My Address), receiver address (Figure 8: To Address), and amount (Figure 8: Amount (ETH)). Application will retrieve sender address from coin base which is a primary wallet address of user who

logged in. The receiver address 0x545079d5c9ba0ad896d0b6f1d724185b12f043f9 from user input is a valid address and is in the same network with the sender address (0xc636d4cc527c48aace6ac94a2 007bfad927dce6b), so the command sendTransaction() can be executed correctly. Once the last digit of the receiver address was removed giving “0x545079d5c9ba0ad896d0b6f1d724185b12f 043f” and “Send Ether” was clicked, there was an error in Web3.js console as shown in Figure 10.

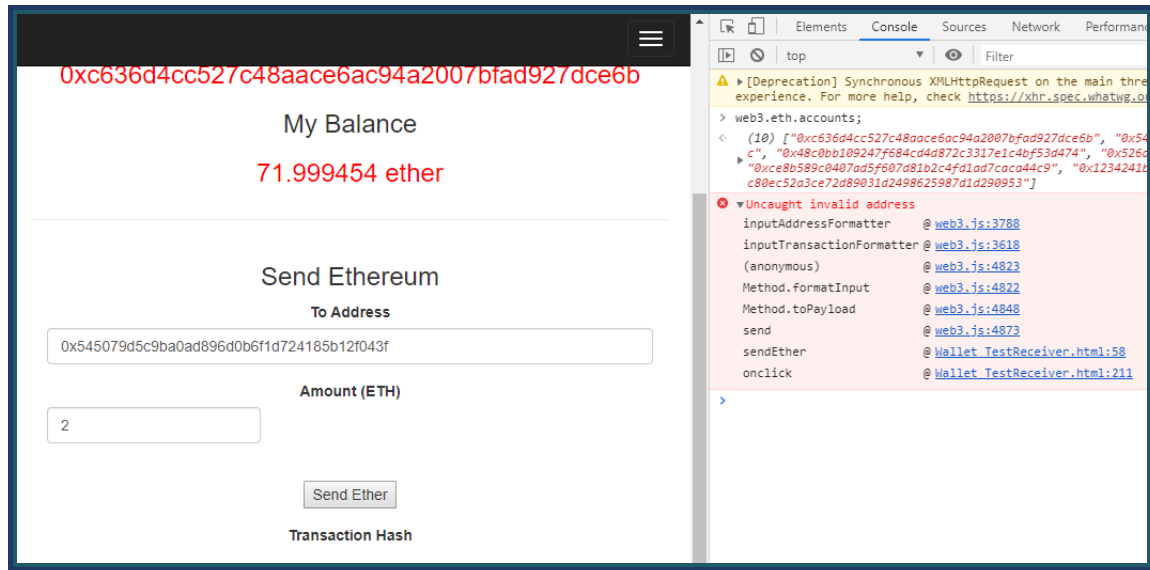


Figure 10. Web3.js error when inputting the wrong format of the receiver address.

We then tried again with the same address, but the last digit was changed from 9 to 8. As shown in Figure 11, we found that the command `sendTransaction()` can be executed normally. The transaction was confirmed, and the balance of “My Address: 0xc636d4cc527c48aace6ac94a2007bfad927dce6b” was decreased by the amount that has been sent. This

can be concluded that the command `sendTransaction()` can detect only the invalid address from the wrong format but cannot detect the invalid address in the right format. This problem also happens in Metamask which is the worldwide Ethereum wallet as we have already explained in section 2.

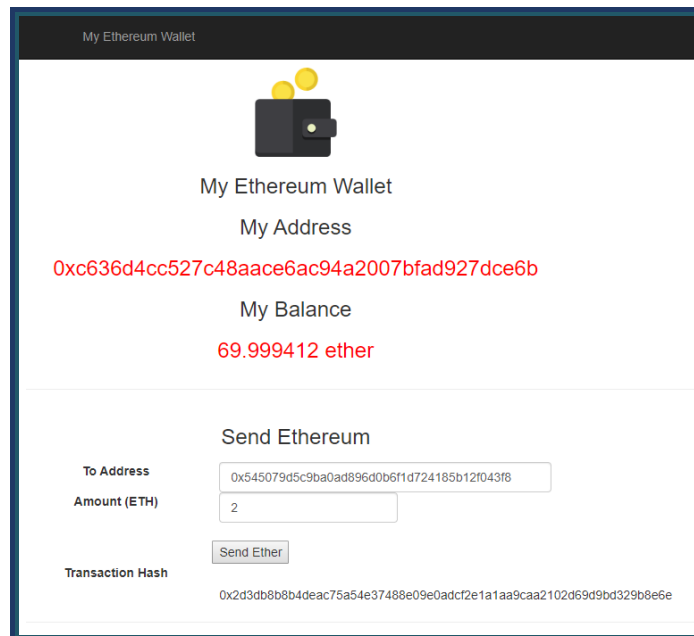


Figure 11. Command `sendTransaction()` can be executed normally on the invalid address.

To prevent the problem as shown in Figure 11, a script was developed to verify the receiver address

before sending it to command `sendTransaction()` as shown in Figure 12.

```

if(confirm("Send to : " + to + "\nAmount : " + amount + " ETH" + "\nDo you want to process?")){
    var xmlhttp = new XMLHttpRequest();
    var url = "https://api.blockcypher.com/v1/eth/main/addrs/" + to + "/balance";
    var valApi = {};
    var totalReceived = 0;
    xmlhttp.onload = function () {
        if (xmlhttp.readyState == 4 && xmlhttp.status == "200") {
            valApi = JSON.parse(this.responseText);
        }
    }
    xmlhttp.open("GET", url, true);
    xmlhttp.send();

    setTimeout(function(){
        totalReceived = valApi.total_received;
        if (totalReceived == 0){
            if(confirm("This address has never received so far. It may be wrong address!!! Do you want to process?")){
                webs.eth.sendTransaction(params, function (error, result) {
                    if(!error){
                        document.getElementById('send_output').innerHTML = result;
                    }else{
                        document.getElementById('send_output').innerHTML = error;
                    }
                });
            }else{
                webs3.eth.sendTransaction(params, function (error, result) {
                    if(!error){
                        document.getElementById('send_output').innerHTML = result;
                    }else{
                        document.getElementById('send_output').innerHTML = error;
                    }
                });
            }
        }
    }, 1000); //end function setTimeout
}

```

Figure 12. JavaScript to verify receiver address using blockchain API.

The verification script shown in Figure 12 uses the third-party API (1) to check and notify the user if the receiver address has never received any amount (2) so far in the same network. As an example, the result of using the direct API link shown in Figure 13 indicates that the total_received column of address

“0x64D0c83f7852A741381F16088b469135c6154384” is 0 which means that this address has never received any amount in the past. So, there will be the notification to the user, as shown in Figure 14. The user can continue to send to this address if he is sure that this address is newly created in the same network.

```

{
  "address": "0x64D0c83f7852A741381F16088b469135c6154384",
  "total_received": 0,
  "total_sent": 0,
  "balance": 0,
  "unconfirmed_balance": 0,
  "final_balance": 0,
  "n_tx": 0,
  "unconfirmed_n_tx": 0,
  "final_n_tx": 0
}

```

Figure 13. API checking result of address “0x64D0c83f7852A741381F16088b469135c6154384”.

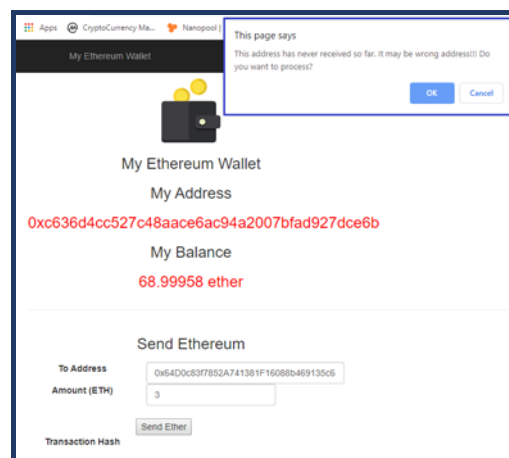


Figure 14. User notification if the receiver address has never been received so far.

This test result shows that the user's mistakes can be found and reduced before the construction and submission of the transaction to the blockchain network.

5. Conclusions and Discussions

Lack of KYC (Know Your Customer) is the main weaknesses of blockchain technology. There is no information helping users to verify whether the address is really belonging to the target receiver as the traditional centralized banking system. With the increasing use of cryptocurrency, new user has obtained new wallet every day without awareness of security. Some user never aware of any mistake that may happen until it happened. Whenever the mistaken input committed, user cannot claim to any centralized authority even the web exchange. They also cannot find out who is the owner of mistake address because there is no information like that in blockchain network. This mistake causes the cryptocurrency loss without recovery forever. To reduce the mistake transactions committed by users and reduce the forever loss of cryptocurrency in the market, we recommend adding our proposed verification process into the existing wallet applications.

References

- [1] S. Nakamoto, "Bitcoin: A peer-to-peer electronic cash system," 2008. [Online]. Available: <https://Bitcoin.org/Bitcoin.pdf>.
- [2] Ethereum. [Online]. Available: <https://www.ethereum.org>.
- [3] Litecoin - Open source P2P digital currency. [Online]. Available: <https://litecoin.org/>.
- [4] Monero - secure, private, untraceable. [Online]. Available: <https://www.getmonero.org/>.
- [5] Dash - Dash is Digital Cash You Can Spend Anywhere. [Online]. Available: <https://www.dash.org/>.
- [6] A. M. Antonopoulos, "Mastering Bitcoin: Unlocking Digital Cryptocurrencies," 1st ed. Sebastopol, CA, USA: O'Reilly Media, Inc., 2014.
- [7] Testnet Ropsten (ETH) Blockchain Explorer. [Online]. Available: <https://ropsten.etherscan.io/>.
- [8] Metamask, Brings Ethereum to your browser [Online]. Available: <https://metamask.io>.
- [9] N. Amarasinghe, X. Boyen, and M. McKague. "A Survey of Anonymity of Cryptocurrencies,". In Proceedings of the Australasian Computer Science Week Multiconference (ACSW 2019). ACM, New York, NY, USA.
- [10] M. Conti, E. Sandeep Kumar, C. Lal and S. Ruj, "A Survey on Security and Privacy Issues of Bitcoin," in IEEE Communications Surveys & Tutorials, vol. 20, no. 4, pp. 3416-3452, Fourth Quarter 2018.
- [11] Q. ShenTu and J. Yu, "Research on anonymization and deanonymization in the Bitcoin system," arXiv preprint arXiv:1510.07782, Oct. 2015. [Online]. Available: <https://arxiv.org/abs/1510.07782>.
- [12] Y. Liu, X. Liu, L. Zhang, C. Tang, and H. Kang, "An efficient strategy to eliminate malleability of Bitcoin transaction," 2017 4th International Conference on Systems and Informatics (ICSAI), Hangzhou, 2017, pp. 960-964.
- [13] A. R. Sai, J. Buckley and A. Le Gear, "Privacy and Security Analysis of Cryptocurrency Mobile Applications," 2019 Fifth Conference on Mobile and Secure Services (MobiSecServ), Miami Beach, FL, USA, 2019, pp. 1-6.
- [14] Bitcoin talk. If i send my Bitcoins to a wrong address, what happens? [Online]. Available: <https://Bitcointalk.org/index.php?topic=21127.0>.
- [15] Coinbase Q&A. "I sent funds to the wrong address. How do I get them back?" [Online]. Available: <https://support.coinbase.com/customer/en/portal/articles/1404411-i-sent-funds-to-the-wrong-address-how-do-i-get-them-back->.
- [16] Quora. "What happens if I send my Bitcoin to an incorrect address?" [Online]. Available: <https://www.quora.com/What-happens-if-I-send-my-Bitcoin-to-an-incorrect-address>.
- [17] Official Go implementation of the Ethereum protocol. [Online]. Available: <https://geth.ethereum.org/>.
- [18] Ethereum JavaScript API. [Online]. Available: <https://github.com/ethereum/web3.js/>.



การวิเคราะห์ปัญหาสุขภาพจากภาพถ่ายเล็บด้วยเทคนิคการเรียนรู้เชิงลึก Health Problem Analysis from Nail Image using Deep Learning Technique

ณัฐวดี หงษ์บุญมี* และ จุฑามาศ กันยาประสิทธิ์

Nattavadee Hongboonmee and Jutamas Kanyaprasit*

ภาควิชาวิทยาการคอมพิวเตอร์และเทคโนโลยีสารสนเทศ คณะวิทยาศาสตร์ มหาวิทยาลัยนเรศวร

Department of Computer Science and Information Technology, Faculty of Science,

Naresuan University

Received: June 04, 2021; Revised: August 06, 2021; Accepted: October 06, 2021; Published: December 28, 2021

ABSTRACT – The objectives of this research are (1) to develop a nail image classification model for health problem analysis using deep learning technique; (2) to develop a system for analyzing health problems from nails image; and (3) to assess the system performance. This research was conducted by 400 sample images of five nail types such as psoriasis nails image, rheumatoid nails image, anemia nails image, melanoma nail images and normal nail image. The sample images were trained and generated classifier models using deep convolutional neural network. The results show that there is 90.20% of accuracy for the sample image set. Then develop the user interface through the application on the android operating system. The application will run the model through a series of instructions to analyze health problems from nail images. The results of the study showed that the applications that developed able to analyze health problems from nail images, the performance is good. It has an average accuracy of 78.00%. The results of the application satisfaction assessment by the users were at a good level. The average is 4.09. Therefore, the presented methods and applications can used as a tool to help analyze health problems that may arise from nail images.

KEYWORDS: Health, Nail, Image Classification, Deep Learning, Convolutional Neural Network

บทคัดย่อ -- งานวิจัยนี้มีวัตถุประสงค์ คือ (1) เพื่อพัฒนาโมเดลจำแนกภาพถ่ายเล็บสำหรับตรวจสอบปัญหาสุขภาพโดยใช้เทคนิคการเรียนรู้เชิงลึก (2) เพื่อพัฒนาระบบตรวจสอบปัญหาสุขภาพจากภาพถ่ายเล็บ และ (3) เพื่อประเมินประสิทธิภาพระบบ การดำเนินงานเริ่มจากการรวบรวมข้อมูลกลุ่มตัวอย่างภาพถ่ายเล็บจำนวน 5 ประเภท ประกอบด้วย ภาพถ่ายเล็บโรคสะเก็ดเงิน ภาพถ่ายเล็บโรคไขข้ออักเสบ ภาพถ่ายเล็บโรคโลหิตจาง ภาพถ่ายเล็บโรคเมรังคิวหนังและภาพถ่ายเล็บปกติ รวมทั้งหมด 400 ภาพ โดยการนำชุดข้อมูลภาพตัวอย่างมาฝึกสอนสร้างโมเดลจำแนกภาพถ่ายเล็บด้วยโครงข่ายประสาทเทียมเชิงลึกคอนโวลูชัน ผลการศึกษาพบว่าผลลัพธ์ของค่าความถูกต้องของโมเดลในการตรวจสอบชุดภาพตัวอย่าง มีค่าความถูกต้องเท่ากับ 90.20% จากนั้นพัฒนาส่วนติดต่อผู้ใช้งานผ่านแอปพลิเคชันบนระบบปฏิบัติการแอนดรอยด์ โดยแอปพลิเคชันจะเรียกใช้โมเดลผ่าน

*Corresponding Author: nattavadeeho@nu.ac.th

ชุดคำสั่งเพื่อตรวจสอบปัญหาสุขภาพจากภาพถ่ายเล็บ ผลการทดสอบพบว่าค่าความถูกต้องในการตรวจสอบปัญหาสุขภาพจากภาพถ่ายเล็บของแอปพลิเคชันประสิทธิภาพอยู่ในระดับดี มีค่าความถูกต้องเฉลี่ย 78.00% และผลการประเมินความพึงพอใจของระบบโดยผู้ใช้งานอยู่ในระดับดี ค่าเฉลี่ยเท่ากับ 4.09 แสดงให้เห็นว่าวิธีการและแอปพลิเคชันที่นำเสนอนี้สามารถนำไปใช้ประโยชน์ในการเป็นเครื่องมือช่วยตรวจสอบปัญหาสุขภาพที่อาจจะเกิดขึ้นจากภาพถ่ายเล็บได้

คำสำคัญ: สุขภาพ, เล็บ, การจำแนกภาพ, การเรียนรู้เชิงลึก, โครงข่ายประสาทเทียมคอนโวลูชัน

1. บทนำ

การวิเคราะห์ปัญหาสุขภาพเป็นการตรวจหาความผิดปกติในร่างกาย ซึ่งปัจจุบันในการตรวจสุขภาพของการแพทย์ใช้การพิจารณาความผิดปกติของร่างกายเพื่อวิเคราะห์ว่าสุขภาพเป็นอย่างไร และการพิจารณาความผิดปกติของเล็บก็เป็นวิธีหนึ่งที่สามารถวิเคราะห์ปัญหาสุขภาพได้ [1] โดยเล็บ คือ ส่วนหนึ่งของร่างกายลักษณะเป็นแผ่นปกคลุมอยู่บริเวณปลายนิ้วมือ นิ้วเท้า เล็บจึงเป็นอวัยวะหนึ่งที่มีความสำคัญ ลักษณะของเล็บสามารถบ่งบอกถึงสุขภาพภายในได้ [2] การตรวจลักษณะเล็บจึงเป็นแนวทางหนึ่งในการช่วยวิเคราะห์ปัญหาสุขภาพ หลักการในการตรวจสอบนั้น [3] จะเป็นการพิจารณาเล็บจากลักษณะรูปร่าง สีของเล็บ เพื่อวิเคราะห์ปัญหาสุขภาพ ซึ่งจะเป็นการตรวจสอบเบื้องต้นประเมินความเป็นไปได้ที่จะมีความเสี่ยงเป็นโรคอะไรได้บ้าง [4] ก่อนที่จะไปพบแพทย์เพื่อวินิจฉัยโดยละเอียดต่อไป แต่ปัญหาที่สำคัญคือประชาชนทั่วไปที่ไม่มีองค์ความรู้เรื่องการตรวจวินิจฉัยโรค ไม่สามารถรู้ได้ว่าเล็บที่มีลักษณะเช่นนี้เป็นโรคอะไร จึงทำให้เสียเวลาในการค้นหาจากหนังสือหรือแหล่งความรู้อื่นๆ ซึ่งไม่สามารถทำให้การตรวจวินิจฉัยโรคทำได้ทันเวลาเนื่องจากโรคบางประเภทสามารถรักษาให้หายขาดได้หากพบในระยะเริ่มแรก เช่น โรคมะเร็ง เป็นต้น

จากปัญหาดังกล่าว จะเห็นว่าการแก้ปัญหาคือการวิเคราะห์ปัญหาสุขภาพจากเล็บดังกล่าวโดยใช้เทคโนโลยีคอมพิวเตอร์ร่วมกับเทคโนโลยีการประมวลผลภาพจะทำให้มีความรวดเร็วและตอบสนองความต้องการของผู้ใช้ที่ต้องการตรวจสอบคัดกรองเบื้องต้นได้ว่ากำลังอยู่ในความเสี่ยงที่จะเป็นโรคอะไรได้ อย่างมีประสิทธิภาพ

ดังนั้นงานวิจัยนี้จึงนำเสนอการพัฒนาระบบเพื่อช่วยวิเคราะห์คัดกรองความเสี่ยงโรคเบื้องต้น โดยตรวจสอบจากภาพถ่ายเล็บผ่านแอปพลิเคชันบนสมาร์ตโฟน วัตถุประสงค์ของ

การวิจัย คือ (1) เพื่อพัฒนาโมเดลจำแนกภาพถ่ายเล็บสำหรับตรวจสอบปัญหาสุขภาพ (2) เพื่อพัฒนาระบบตรวจสอบปัญหาสุขภาพจากภาพถ่ายเล็บผ่านสมาร์ตโฟนระบบปฏิบัติการแอนดรอยด์ และ (3) เพื่อประเมินประสิทธิภาพแอปพลิเคชัน โดยประยุกต์ใช้เทคนิคการเรียนรู้เชิงลึก (Deep Learning) ซึ่งเป็นเทคนิคที่สามารถเรียนรู้คุณลักษณะพิเศษของรูปภาพและจำแนกประเภทข้อมูลแบบอัตโนมัติ [5] โมเดลที่ได้จากงานวิจัยนี้ ถูกนำมาประยุกต์ใช้บนแอปพลิเคชันจำแนกประเภทภาพถ่ายด้วยรูปภาพจากกล้องถ่ายรูปบนสมาร์ตโฟน เพื่อใช้เป็นเครื่องมือตรวจสอบสุขภาพเบื้องต้นให้ผู้ใช้งาน ตรวจสอบประเมินความเสี่ยงโรคด้วยตนเองเบื้องต้น ก่อนที่จะไปพบแพทย์เพื่อวินิจฉัยโดยละเอียดต่อไป

2. ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

2.1 ลักษณะเล็บบอกสุขภาพ

ปัญหาสุขภาพมักเกิดจากความผิดปกติของร่างกายส่วนใดส่วนหนึ่งและส่งผลกระทบถึงกันได้ ตัวอย่างเช่น เล็บที่สามารถบอกลักษณะผิดปกติของร่างกาย [1] ดังนั้นการที่เล็บมีความผิดปกติก็บ่งบอกถึงอาการของโรคที่อาจจะเกิดขึ้นได้ ในการศึกษาที่มีความสนใจที่จะจำแนกกลุ่มโรคที่สามารถบ่งบอกได้จากลักษณะของเล็บ 4 ประเภท ได้แก่ โรคไขข้ออักเสบ โรคสะเก็ดเงิน โรคมะเร็งผิวหนัง และโรคโลหิตจาง ซึ่งทั้ง 4 โรคนี้ บางครั้งมีลักษณะทางกายภาพที่คล้ายคลึงกันมาก ซึ่งถือเป็นความท้าทายที่จะทำให้ระบบคอมพิวเตอร์มีความสามารถที่จะจำแนกโรคต่างๆ จากลักษณะเล็บได้อย่างมีประสิทธิภาพ

โรคไขข้ออักเสบ (Rheumatoid) [1]-[4] เป็นภาวะที่ข้อต่อกระดูกเกิดการอักเสบ ผู้ป่วยจะมีอาการปวดข้อ บวมแดงเกิดขึ้นได้หลายส่วนในร่างกาย อาการของไขข้ออักเสบขนาดเล็กอย่างนิ้วมือ อาจทำให้สูญเสียแรงบีบมือ ซึ่งมีผลต่อการหยิบจับสิ่งต่างๆ นอกจากนี้ยังอาจพบความผิดปกติที่เล็บ เช่น ผิวของ

เล็บไม่เรียบเป็นหลุมเล็กๆ เล็บผิดรูปขรุขระ เล็บหนา มีขุยขาวใต้เล็บ หรือเล็บล่อนจากพื้นเล็บ เป็นต้น รูปที่ 1 แสดงตัวอย่างเล็บโรคไขข้ออักเสบ



รูปที่ 1. ตัวอย่างภาพเล็บโรคไขข้ออักเสบ

โรคสะเก็ดเงิน (Psoriasis) [1]–[4] ความผิดปกติของเล็บที่พบในโรคสะเก็ดเงิน อาจมีอาการเป็นเฉพาะที่เล็บหรือพบ รอยโรคที่ผิวหนัง สาเหตุเกิดจากการแบ่งตัวของเซลล์ผิวหนังที่ผิดปกติบริเวณใต้โคนเล็บ และบริเวณใต้แผ่นเล็บ ทำให้เกิดการอักเสบและมีรอยโรคเกิดขึ้นที่แผ่นเล็บ พบได้ใต้โคนเล็บ เล็บแยกชั้น เล็บขรุขระไม่เรียบหรืออาจพบจุดขาวที่เล็บ รูปที่ 2 แสดงตัวอย่างเล็บโรคสะเก็ดเงิน



รูปที่ 2. ตัวอย่างภาพเล็บโรคสะเก็ดเงิน

โรคมะเร็งเป็นโรคที่เกิดจากความผิดปกติของเซลล์ในอวัยวะต่างๆ ของร่างกาย ส่วนโรคมะเร็งผิวหนัง (Melanoma) [1]–[4] เป็นโรคมะเร็งชนิดหนึ่ง ซึ่งพบได้น้อยแต่มีความรุนแรงมากกว่ามะเร็งผิวหนังชนิดอื่น อาการของโรคคือ ตรวจพบการหนาตัวหรือก้อนที่ผิวหนัง พบรอยจุดที่ผิวหนังและมีรอยขีดสีน้ำตาลเข้มหรือดำอยู่บริเวณเล็บ หากตรวจพบตั้งแต่ในระยะแรกก็จะรักษาให้หายได้ แต่หากล่าช้าเซลล์มะเร็งอาจแพร่กระจายไปยังส่วนอื่นๆ ของร่างกายจนทำให้ยากต่อการรักษา รูปที่ 3 แสดงตัวอย่างเล็บโรคมะเร็งผิวหนัง



รูปที่ 3. ตัวอย่างภาพเล็บโรคมะเร็งผิวหนัง

โรคโลหิตจาง (Anemia) [1]–[4] หรือภาวะซีดเป็นภาวะที่ร่างกายมีปริมาณเม็ดเลือดแดงในเลือดน้อยกว่าปกติทำให้นำออกซิเจนไปยังเซลล์และเนื้อเยื่อในอวัยวะต่างๆ น้อยลงส่งผลให้ผู้ป่วยมีอาการ เช่น เหนื่อยล้า อ่อนเพลีย ผิวซีดหรือผิวเหลือง และมีอาการตัวซีดเล็บสีซีด เป็นต้น รูปที่ 4 แสดงตัวอย่างเล็บโรคโลหิตจาง



รูปที่ 4. ตัวอย่างภาพเล็บโรคโลหิตจาง

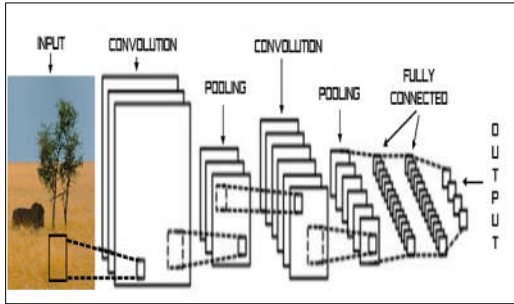
2.2 การเรียนรู้เชิงลึก

ในปัจจุบันมีเทคนิคที่ใช้ในการจำแนกภาพ (Image Classification) หลายรูปแบบ โดยเทคนิคการเรียนรู้เชิงลึกเป็นวิธีการที่ได้รับความนิยมมากในการนำมาใช้จำแนกรูปภาพ และกล่าวได้ว่าเป็นการจำแนกรูปภาพที่ทันสมัยที่สุดแบบหนึ่ง

การเรียนรู้เชิงลึก (Deep Learning) [6] คือ วิธีการเรียนรู้แบบอัตโนมัติที่จำลองมาจากโครงข่ายประสาทเทียม (Artificial Neuron Network : ANN) โดยนำระบบ ANN มาซ้อนกันหลายชั้น (Layer) และทำการเรียนรู้ข้อมูลตัวอย่างจากนั้นนำข้อมูลที่ได้ออกมาหาารูปแบบ (Pattern) หรือจัดหมวดหมู่ข้อมูล (Data Classification)

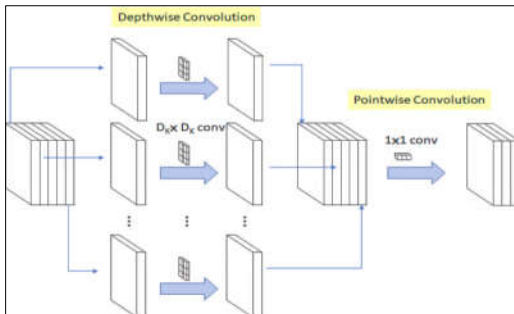
2.3 โครงข่ายประสาทเทียมคอนโวลูชัน

โครงข่ายประสาทเทียมคอนโวลูชัน (Convolution Neural Network: CNN) [7] เป็นรูปแบบหนึ่งของโครงข่ายประสาทเทียม (ANN) ซึ่งเป็นลักษณะโครงข่ายประสาทเทียมแบบเชิงลึก (Deep Learning Neural Network) โดยที่ CNN เป็นอัลกอริทึมที่เน้นใช้กับรูปภาพ โดยจะดึงจุดเด่นของภาพนั้นๆ ออกมา ส่วนการคอนโวลูชัน (Convolution) ภาพ หมายถึง การมองภาพเป็นเมตริกซ์ของพิกเซลเรียงต่อกันและจำลองการมองเห็นของมนุษย์ที่มองพื้นที่เป็นพื้นที่ย่อยๆ หรือนำมาจำแนก (Classification) และนำกลุ่มของพื้นที่ย่อยๆ มาผสานกันเพื่อดูว่าสิ่งที่เห็นอยู่เป็นสิ่งใด



รูปที่ 5. แสดงสถาปัตยกรรมของโครงข่ายประสาทเทียมคอนโวลูชัน

งานวิจัยนี้เลือกใช้โครงข่ายประสาทเทียมคอนโวลูชันแบบ MobileNet ซึ่งเป็นโครงข่ายการเรียนรู้ที่ถูกรออกแบบมาสำหรับงานที่มีทรัพยากรจำกัด ใช้พลังงานในการประมวลผลไม่มาก โมเดลมีขนาดเล็ก มีประสิทธิภาพในการจำแนกข้อมูลรูปภาพได้รวดเร็วเมื่อนำโมเดลไปใช้งานบนโทรศัพท์มือถือ ดังแสดงในงานวิจัยของ [8]



รูปที่ 6. แสดงสถาปัตยกรรมของ MobileNet

2.4 งานวิจัยที่เกี่ยวข้อง

ปัจจุบันเทคโนโลยีด้านการเรียนรู้เชิงลึกมีการพัฒนาไปอย่างมากจนทำให้คอมพิวเตอร์สามารถจำแนกภาพถ่ายได้ด้วยการเรียนรู้จากภาพ และในช่วงทศวรรษที่ผ่านมาได้มีนักวิจัยหลายท่านได้ทำการวิจัยเกี่ยวกับการจำแนกรูปภาพทางการแพทย์ ตัวอย่างเช่น งานวิจัยของ Orlando และคณะ [9] นำเสนอวิธีการตรวจหารอยแดงที่ตา ซึ่งอาจจะบ่งบอกถึงโรคเบาหวาน โดยอาศัยหลักการเรียนรู้เชิงลึกด้วยโครงข่ายประสาทเทียม CNN ร่วมกับหลักการ Domain Knowledge ผลลัพธ์ที่ได้พบว่าการรวมวิธีการทั้งสองเข้าด้วยกันให้ผลลัพธ์การตรวจหารอยแดงที่ตาได้ดีกว่าการแยกกัน

Metkarunchit และ Charoenpojvajan [10] ศึกษาเรื่อง การตรวจหา Covid-19 โดยใช้การเรียนรู้เชิงลึกกับภาพซีทีสแกนประยุกต์ใช้ Mask RCNN เพื่อทำนายส่วนบริเวณเนื้อเยื่อที่ได้รับผลกระทบจากไวรัสจากภาพถ่ายซีทีสแกนบริเวณทรวงอก ผลการทดสอบโมเดลนี้มีความถูกต้องเท่ากับ 89.00%

งานวิจัยของ Chin และคณะ [11] ได้นำเสนอระบบตรวจจับโรคหลอดเลือดสมองขาดเลือดในระยะแรกแบบอัตโนมัติ โดยใช้อัลกอริทึมการเรียนรู้เชิงลึกแบบ CNN ระบบจะประมวลผลภาพซีทีสแกนของสมอง ผลการทดลองพบว่าอัตราความถูกต้องของขั้นตอนการ Train คือ 97.65%

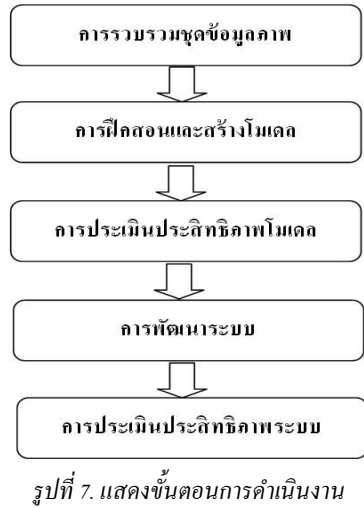
งานวิจัย Que และคณะ [12] ใช้การเรียนรู้เชิงลึกสร้างโมเดลที่เรียกว่า CardioXNet เพื่อวินิจฉัยโรคหัวใจจากภาพถ่ายรังสีทรวงอก ผลการศึกษาพบว่าโมเดลนี้มีความถูกต้องเท่ากับ 93.75%

Lamsamut และ Valuvanathorn [13] ศึกษาการประยุกต์ใช้โครงข่ายประสาทเทียมคอนโวลูชัน สำหรับการจำแนกภาพซีทีสแกนโรคหลอดเลือดสมอง โดยใช้ชุดข้อมูลภาพ 3 ประเภท ได้แก่ ภาพผู้ป่วยโรคหลอดเลือดสมองตีบ ภาพผู้ป่วยโรคหลอดเลือดสมองแตก และภาพหลอดเลือดสมองปกติ ผลการทดลองพบว่าได้ค่าความถูกต้อง 92.60%

3. วิธีการดำเนินการ

งานวิจัยนี้ดำเนินการวิจัยโดยมีวัตถุประสงค์เพื่อพัฒนาโมเดลจำแนกกลุ่มโรคที่สามารถบ่งบอกได้จากลักษณะของเส้นด้วยเทคนิคการเรียนรู้เชิงลึก และพัฒนาระบบตรวจสอบปัญหาสุขภาพจากภาพถ่ายเส้นผ่านสมาร์ทโฟน โดยวิธีการดำเนินการวิจัย ได้แบ่งขั้นตอนการทำงานออกเป็น 5 ขั้นตอน ได้แก่ (1) การเก็บรวบรวมชุดข้อมูลภาพ (2) การฝึกสอนและสร้างโมเดล (3) การประเมินประสิทธิภาพโมเดล (4) การพัฒนาระบบ และ (5) การประเมินประสิทธิภาพระบบ ดังรูปที่ 7 สามารถอธิบายรายละเอียดในแต่ละขั้นตอน ได้ดังนี้

การศึกษานี้มีความสนใจที่จะจำแนกกลุ่มโรคที่สามารถบ่งบอกได้จากลักษณะของเส้น 4 ประเภท ได้แก่ โรคไขข้ออักเสบ โรคสะเก็ดเงิน โรคเมะเร็งผิวหนัง และโรคโลหิตจาง เนื่องจากทั้ง 4 โรคนี้บางครั้งมีลักษณะทางกายภาพที่คล้ายคลึงกันมาก ซึ่งถือเป็นความท้าทายที่จะทำให้ระบบคอมพิวเตอร์ มีความสามารถที่จะจำแนกโรคต่างๆ จากลักษณะเส้นได้อย่างมีประสิทธิภาพ



3.1 การเก็บรวบรวมชุดข้อมูลภาพ

การเก็บรวบรวมข้อมูลและเตรียมข้อมูล (รูปภาพ) ข้อมูลส่วนหนึ่งเก็บรวบรวมในจังหวัดพิษณุโลกจากกล้องโทรศัพท์มือถือ แต่ข้อมูลภาพที่รวบรวมได้มีจำนวนน้อยไม่เพียงพอ จึงค้นหาเพิ่มเติมจากเว็บไซต์ที่เป็นแหล่งแจกจ่ายรูปที่อนุญาตให้ใช้ได้ โดยรวบรวมรูปภาพสำหรับจำแนกกลุ่มโรคที่สามารถบ่งบอกได้จากลักษณะของเล็บ 5 ประเภท ได้แก่ โรคไขข้ออักเสบ โรคสะเก็ดเงิน โรคมะเร็งผิวหนัง โรคโลหิตจาง และภาพเล็บปกติ ความละเอียดภาพ 224x224 พิกเซล จากนั้นให้ผู้เชี่ยวชาญทางการแพทย์ตรวจสอบความถูกต้องของรูปภาพให้ตรงกับประเภทโรคต่างๆ ดังตัวอย่างในรูปที่ 8 จำนวนรูปภาพที่ผ่านการตรวจสอบมีจำนวน 400 ภาพ ดังตารางที่ 1 จากนั้นนำมาแบ่งด้วยวิธีสุ่มเพื่อใช้เป็นชุดข้อมูลฝึกสอน และชุดข้อมูลทดสอบจำนวน 320 และ 80 ภาพ ตามลำดับ



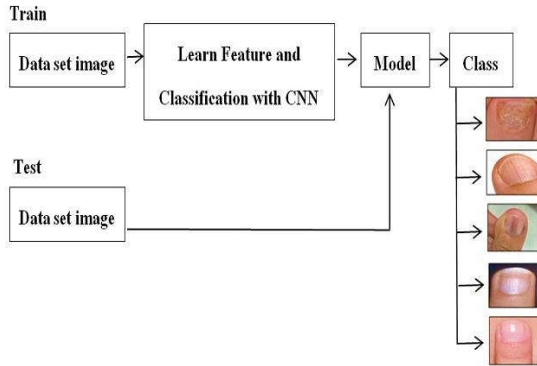
รูปที่ 8. ตัวอย่างชุดข้อมูลภาพในแต่ละประเภท

ตารางที่ 1. จำนวนชุดข้อมูลภาพเล็บแต่ละประเภท

ประเภทของ ภาพเล็บ	จำนวนรูปภาพ		รวม
	Train	Test	
ไขข้ออักเสบ (Rheumatoid)	64	16	80
สะเก็ดเงิน (Psoriasis)	64	16	80
มะเร็งผิวหนัง (Melanoma)	64	16	80
โลหิตจาง (Anemia)	64	16	80
ปกติ (Normal)	64	16	80
รวม	320	80	400

3.2 การฝึกสอนและสร้างโมเดล

การฝึกสอนและสร้างโมเดลจำแนกกลุ่มโรคประยุกต์ใช้เทคนิคการเรียนรู้เชิงลึกแบบโครงข่ายประสาทเทียมคอนโวลูชัน (CNN) เครื่องมือในการฝึกสอนสร้างโมเดล ได้แก่ ภาษา Python และไลบรารีมาตรฐานสำหรับการเรียนรู้เชิงลึกที่มีชื่อว่า TensorFlow ซึ่งไลบรารีนี้จะทำการแบ่งชุดข้อมูลรูปภาพจำนวน 400 รูปออกเป็นสองส่วนคือชุดข้อมูลฝึกสอน (Training Data) และชุดข้อมูลทดสอบ (Testing Data) โดยในแต่ละครั้งที่ทดสอบการเรียนรู้ของโมเดล โปรแกรมจะทำการแบ่งข้อมูลออกมาเป็นข้อมูลฝึกสอนจำนวน 80% และชุดข้อมูลทดสอบจำนวน 20% ของข้อมูลทั้งหมดโดยประมาณ นอกจากนี้โครงข่ายประสาทเทียมคอนโวลูชันยังสามารถเรียนรู้และจำแนก (Learn Feature & Classification) เองได้โดยที่จำเป็นต้องสร้าง Feature Extraction ขั้นตอนการฝึกสอนและสร้างโมเดลดังรูปที่ 9 การฝึกสอนและสร้างโมเดลจะดำเนินการทั้งหมดจำนวน 500 รอบ ดังแสดงในรูปที่ 10 เมื่อทำการฝึกสอนจนครบแล้วจะได้ไฟล์โมเดลจำนวน 2 ไฟล์ ได้แก่ retrained_graph.pb และ retrained_labels.txt ซึ่งเป็นไฟล์ที่จะนำไปใช้สร้างระบบต่อไป



รูปที่ 9. แสดงขั้นตอนการจำแนกประเภทของรูปภาพ

```

INFO:tensorflow:2020-09-29 09:30:31.720673: Step 499
INFO:tensorflow:2020-09-29 09:30:31.720673: Step 499
INFO:tensorflow:2020-09-29 09:30:31.804956: Step 499
INFO:tensorflow:Final test accuracy = 90.2% (N=82)
  
```

รูปที่ 10. แสดงตัวอย่างการเทรนโมเดล

3.3 การประเมินประสิทธิภาพโมเดล

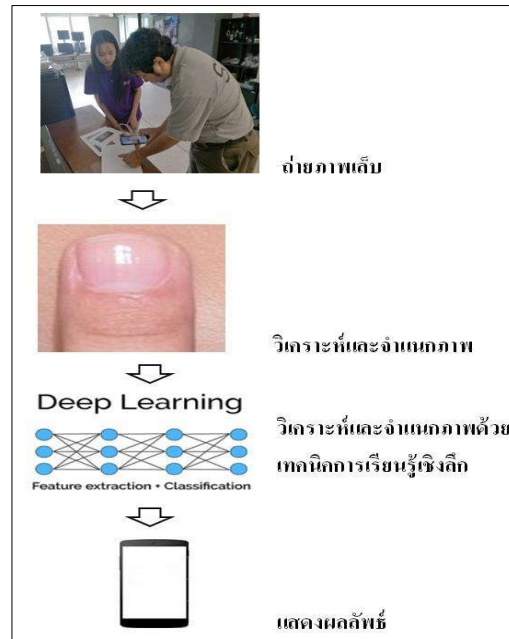
การประเมินประสิทธิภาพโมเดลในงานวิจัยนี้ใช้การวัดค่าความถูกต้อง (Accuracy) ในการจำแนกประเภทภาพ ซึ่งเป็นค่าที่ได้จากวิธีการทดสอบเพื่อหาค่าพยากรณ์ความถูกต้องของข้อมูลโดยคิดเป็นค่าร้อยละ ใช้สูตรคำนวณดังสมการที่ 1 [14]

$$\text{Accuracy} = \left(\frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}} \right) \times 100 \quad (1)$$

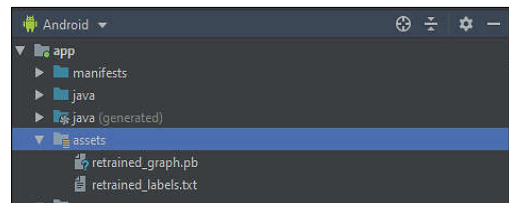
โดย TP คือ คลาสบอกว่าจริงและทำนายว่าจริง TN คือ คลาสบอกว่าไม่จริงและทำนายว่าไม่จริง FP คือ คลาสบอกว่าจริงและทำนายว่าไม่จริง FN คือ คลาสบอกว่าไม่จริงและทำนายว่าจริง

3.4 การพัฒนาระบบ

การออกแบบและพัฒนาระบบเป็นขั้นตอนของการนำโมเดลจำแนกภาพที่ได้ไปพัฒนาส่วนติดต่อผู้ใช้ในรูปแบบแอปพลิเคชันบนสมาร์ตโฟนระบบปฏิบัติการแอนดรอยด์ เพื่อให้ง่ายและสะดวกต่อการใช้งาน โดยดำเนินการออกแบบขั้นตอนการทำงานของแอปพลิเคชัน ดังรายละเอียดในรูปที่ 11



รูปที่ 11. แสดงขั้นตอนการทำงานของแอปพลิเคชัน



รูปที่ 12. ตัวอย่างการนำโมเดลเข้าใช้งานบน Android Studio

เครื่องมือในการพัฒนาระบบใช้โปรแกรม Android Studio และภาษา JAVA โดยการนำไฟล์โมเดลไปไว้ในโปรเจกต์ที่สร้างด้วยโปรแกรม Android Studio (รูปที่ 12) การทำงานของระบบจะสามารถถ่ายภาพเล็บด้วยกล้องสมาร์ตโฟน จากนั้นทำการประมวลผลและจำแนกภาพด้วยโมเดลการเรียนรู้เชิงลึกที่ได้ฝึกสอนไว้เพื่อตรวจสอบรูปภาพ หลังจากนั้นระบบจะแสดงผลลัพธ์ด้วยการวิเคราะห์จำแนกกลุ่มโรคที่ตรวจพบพร้อมแสดงเปอร์เซ็นต์ความถูกต้องที่หน้าจอแอปพลิเคชัน

3.5 การประเมินประสิทธิภาพระบบ

การประเมินประสิทธิภาพระบบ ดำเนินการประเมินประสิทธิภาพการจำแนกกลุ่มโรคจากภาพเล็บของแอปพลิเคชันและประเมินผลความพึงพอใจจากผู้ใช้ โดยได้ทำการติดตั้งแอปพลิเคชันบนสมาร์ตโฟน จากนั้นทำการทดสอบประสิทธิภาพ

แอปพลิเคชัน โดยให้กลุ่มตัวอย่างทดลองใช้งานและประเมินผลระบบ ดำเนินการสรุปผลการประเมินประสิทธิภาพของระบบ ด้วยการหาค่าความถูกต้อง ค่าเฉลี่ยและค่าส่วนเบี่ยงเบนมาตรฐาน

4. ผลการทดลอง

เพื่อหาประสิทธิภาพของระบบ งานวิจัยนี้ได้แบ่งการทดลองไว้ 3 ส่วน ส่วนแรกเป็นการทดสอบหาประสิทธิภาพของโมเดล ด้วยการหาค่าความถูกต้องของการเรียนรู้ภาพ ส่วนที่ 2 ทดสอบประสิทธิภาพการจำแนกภาพตามกลุ่มของโรคต่างๆ ของแอปพลิเคชัน และส่วนที่ 3 ประเมินการใช้ระบบกับผู้ใช้งาน มีรายละเอียดดังนี้

4.1 ผลการทดสอบประสิทธิภาพโมเดล

ผลการทดสอบประสิทธิภาพโมเดลที่ได้จากการนำรูปภาพ 400 ภาพมาทำการเรียนรู้กับโครงข่ายประสาทเทียมคอนโวลูชัน ซึ่งงานวิจัยนี้ทดลองเปรียบเทียบโมเดลการเรียนรู้แบบ MobileNet จำนวน 3 แบบ ได้แก่ MobileNet 0.50, MobileNet 0.75 และ MobileNet 1.0 ที่ให้ค่าความแม่นยำสูง [8] มีประสิทธิภาพในการจำแนกข้อมูลรูปภาพได้เร็ว ไฟล์โมเดลขนาดเล็ก สามารถนำไปใช้งานได้บนสมาร์ตโฟน โดยกำหนดค่าพารามิเตอร์ที่ใช้ในการทดลองดังแสดงในตารางที่ 2

ตารางที่ 2. การกำหนดค่าพารามิเตอร์ที่ใช้ในการทดลอง

ชื่อพารามิเตอร์	ค่าที่กำหนด
ขนาดรูปภาพนำเข้า (Input Size)	224x224 พิกเซล
จำนวนรอบในการฝึกสอน (Epoch)	500 รอบ
อัตราการเรียนรู้ (Learning Rate)	0.001

ตารางที่ 3. แสดงผลการเปรียบเทียบประสิทธิภาพ โมเดล

Model	Loss (%)	Accuracy (%)
MobileNet 1.0	9.80	90.20
MobileNet 0.75	12.20	87.80
MobileNet 0.50	10.80	89.20

ในการฝึกสอนโมเดลได้ทำการทดลองเปรียบเทียบประสิทธิภาพโมเดล 3 แบบ เพื่อศึกษาว่าโมเดลใดมีประสิทธิภาพ

ภาพดีที่สุด รายละเอียดผลการทดลองดังตารางที่ 3 ซึ่งจากการเปรียบเทียบประสิทธิภาพโมเดลบนชุดข้อมูลฝึกสอน พบว่าโมเดล MobileNet 1.0 ได้ค่าความถูกต้องมากที่สุด คือ 90.20% รองลงมา คือ MobileNet 0.50 และ MobileNet 0.75 ได้ค่าความถูกต้องเท่ากับ 89.20%, 87.80% ตามลำดับ ดังนั้นงานวิจัยจึงเลือกใช้โมเดล MobileNet 1.0 ซึ่งได้ค่าความถูกต้องมากที่สุด

ตารางที่ 4. ผลการทดสอบ โดยใช้โมเดล MobileNet 1.0

Category	Loss (%)	Accuracy (%)
ไขข้ออักเสบ	1.00	99.00
สะเก็ดเงิน	9.00	91.00
มะเร็งผิวหนัง	0.00	100.00
โลหิตจาง	10.00	90.00
ปกติ	29.00	71.00
ค่าความถูกต้องเฉลี่ย		90.20

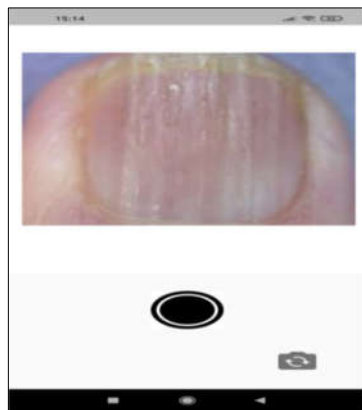
จากตารางที่ 4 แสดงผลการสร้างโมเดลเพื่อจำแนกกลุ่มโรค โดยการทดลองแสดงค่าความถูกต้อง (Accuracy) และค่าความผิดพลาด (Loss) ในการจำแนก เมื่อพิจารณาในแต่ละประเภท พบว่าการจำแนกภาพเล็บโรคมะเร็งผิวหนังมีค่าความถูกต้องของการจำแนกสูงกว่าประเภทอื่น คือ 100.00% การจำแนกภาพเล็บปกติมีความถูกต้องน้อยที่สุด คือ 71.00% ผลลัพธ์ค่าความถูกต้องเฉลี่ยการจำแนกภาพทั้ง 5 ประเภทของโมเดลเท่ากับ 90.20% ซึ่งมีประสิทธิภาพความถูกต้องอยู่ในระดับดี ดังนั้นจึงนำโมเดลนี้ไปประยุกต์ใช้ในการพัฒนาแอปพลิเคชันวิเคราะห์จำแนกกลุ่มโรคจากภาพถ่ายเล็บต่อไป

4.2 ผลการพัฒนาแอปพลิเคชัน

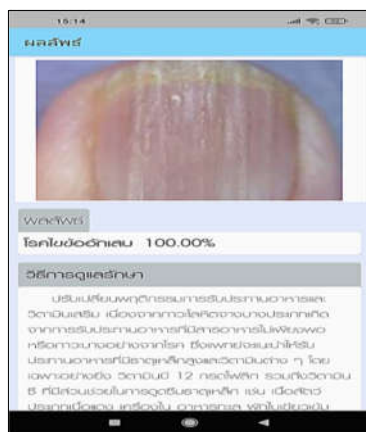
ผลการพัฒนาแอปพลิเคชันวิเคราะห์จำแนกกลุ่มโรคจากภาพถ่ายเล็บ แสดงรายละเอียดแอปพลิเคชันดังนี้ หน้าจอหลักแอปพลิเคชันแสดงดังรูปที่ 13 เมื่อผู้ใช้เลือกเมนูถ่ายภาพ ผู้ใช้งานต้องใช้กล้องสมาร์ตโฟนถ่ายภาพเล็บ ดังรูปที่ 14 จากนั้นแอปพลิเคชันจะเรียกใช้การประมวลผลจากโมเดลการเรียนรู้เชิงลึกที่ได้ฝึกสอนไว้เพื่อวิเคราะห์ภาพเล็บและแสดงผลโรคที่สามารถวิเคราะห์ได้จากลักษณะเล็บทางหน้าจอพร้อมเปอร์เซ็นต์ค่าความถูกต้อง (รูปที่ 15)



รูปที่ 13. หน้าจอหลักแอปพลิเคชัน



รูปที่ 14. หน้าจอการถ่ายรูป



รูปที่ 15. หน้าจอผลลัพธ์

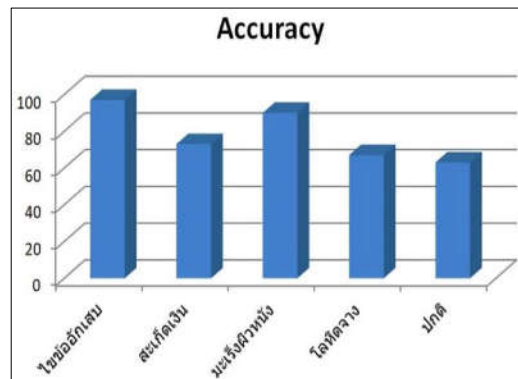
การตรวจสอบประสิทธิภาพของแอปพลิเคชัน โดยการทดสอบการจำแนกภาพกลุ่มโรคจากภาพเล็บ 5 ประเภท ประเภทละ 30 ภาพ ชุดข้อมูลภาพทดสอบในขั้นตอนี้รวบรวมจากเว็บไซต์ที่เป็นแหล่งแจกจ่ายรูปที่อนุญาตให้ใช้ได้ การ

ทดสอบแอปพลิเคชันใช้ตัวชี้วัด คือ ค่าความถูกต้อง (Accuracy) ใช้สมการดังนี้

$$\text{Accuracy} = \left(\frac{\text{No. of Correct Classify}}{\text{No. of Total Classify}} \right) \times 100 \quad (2)$$

ตารางที่ 5. แสดงผลการทดสอบการใช้งานแอปพลิเคชัน

Category	จำแนกถูกต้อง	จำแนกผิดพลาด	Accuracy (%)
ไขข้ออักเสบ	29	1	97.00
สะเก็ดเงิน	22	8	73.00
มะเร็งผิวหนัง	27	3	90.00
โลหิตจาง	20	10	67.00
ปกติ	19	11	63.00
ค่าความถูกต้องเฉลี่ย			78.00



รูปที่ 16. กราฟแสดงค่าความถูกต้องในการจำแนกภาพของแอปพลิเคชัน

จากตารางที่ 5 และรูปที่ 16 แสดงผลการทดสอบประสิทธิภาพการจำแนกภาพของแอปพลิเคชัน โดยสรุปค่าเฉลี่ยความถูกต้องในการจำแนกภาพโรคจากลักษณะเล็บเท่ากับ 78.00% ซึ่งมีประสิทธิภาพความถูกต้องอยู่ในระดับดี โรคที่แอปพลิเคชันสามารถทำนายได้ถูกต้องมากที่สุดได้แก่โรคไขข้ออักเสบ ส่วนประเภทภาพที่ทำนายได้ถูกต้องต่ำที่สุดได้แก่ ภาพเล็บปกติ ซึ่งการจำแนกภาพเล็บที่ผิดพลาดเนื่องจากภาพเล็บปกติ มีลักษณะสีและรูปทรงของเล็บคล้ายคลึงกับภาพเล็บโรคโลหิตจาง จึงทำให้มีบางรูปภาพของภาพเล็บปกติถูกจำแนกโรคผิดเป็นภาพเล็บโรคโลหิตจาง จึงทำให้ภาพเล็บปกติมีความถูกต้องต่ำกว่าภาพโรคอื่นๆ ในทางตรงกันข้ามภาพเล็บโรคไขข้ออักเสบ มีลักษณะที่แตกต่างจากโรคอื่นๆ อย่างชัดเจน

จึงทำให้โมเดลสามารถรู้จำลักษณะของโรคได้อย่างถูกต้อง แม่นยำมากกว่าโรคชนิดอื่นๆ

4.3 ผลการประเมินการใช้งานระบบ

ผลการประเมินการใช้งานระบบวิเคราะห์โรคจากภาพถ่ายเล็บ ได้นำไปทดสอบการใช้งาน ซึ่งกลุ่มตัวอย่างที่ใช้ในการศึกษา ความพึงพอใจ ได้แก่ ผู้เชี่ยวชาญด้านการแพทย์จำนวน 1 คน ผู้เชี่ยวชาญด้านเทคโนโลยีจำนวน 3 คน และกลุ่มผู้ใช้ทั่วไป จำนวน 26 คน รวมจำนวน 30 คน โดยวิเคราะห์จาก แบบสอบถามเพื่อศึกษาความพึงพอใจของกลุ่มตัวอย่าง สถิติที่ใช้ในการวิจัยได้แก่ ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐาน แบ่งระดับการให้คะแนนจากน้อยที่สุดไปหามากที่สุด (ระดับ 1 ถึง 5) ระดับคุณภาพวัดจากค่าเฉลี่ยของข้อมูลโดยมีความหมายดังนี้

4.51-5.00 ดีมาก	3.51-4.50 ดี
2.51-3.50 ปานกลาง	1.51-2.50 พอใช้
1.00-1.50 ควรปรับปรุง	

และจากการนำระบบไปให้กลุ่มตัวอย่างประเมินความพึงพอใจ ผลการประเมินมีรายละเอียดดังตารางที่ 6

ตารางที่ 6. แสดงผลการประเมินการใช้ระบบ

รายการประเมิน	Mean	S.D.	การแปลผล
1.ด้านตรงตามความต้องการของผู้ใช้ (Functional Requirement Test)	4.07	0.51	ดี
2.ด้านการทำงานได้ตามฟังก์ชันงาน (Functional Test)	4.13	0.67	ดี
3.ด้านความสะดวกและง่ายต่อการใช้งาน (Usability Test)	4.27	0.44	ดี
4.ด้านความน่าเชื่อถือของระบบ (Reliability Test)	3.90	0.30	ดี
ค่าเฉลี่ย	4.09	0.48	ดี

จากตารางที่ 6 ผลการประเมินความพึงพอใจการใช้งานระบบของกลุ่มตัวอย่างที่มีต่อระบบที่พัฒนาขึ้น พบว่าภาพรวมความคิดเห็นของกลุ่มตัวอย่างต่อระบบวิเคราะห์โรคจาก

ภาพถ่ายเล็บอยู่ในเกณฑ์ระดับดี มีค่าเฉลี่ยเท่ากับ 4.09 และค่าส่วนเบี่ยงเบนมาตรฐาน (S.D.) เท่ากับ 0.48

กลุ่มตัวอย่างมีความพึงพอใจมากที่สุดในด้านความสะดวกและง่ายต่อการใช้งาน ค่าเฉลี่ย 4.27 รองลงมาได้แก่ ด้านการทำงานได้ตามฟังก์ชันงาน และด้านตรงตามความต้องการของผู้ใช้ ค่าเฉลี่ย 4.13, 4.07 ตามลำดับ

5. สรุปและอภิปรายผล

การศึกษานี้เป็นส่วนหนึ่งของความพยายามที่จะพัฒนาระบบวิเคราะห์โรคเบื้องต้นจากภาพถ่ายเล็บด้วยการประยุกต์ใช้เทคนิคการเรียนรู้เชิงลึกแบบโครงข่ายประสาทเทียมคอนโวลูชันสำหรับการจำแนกภาพ ซึ่งสามารถนำไปใช้ประโยชน์ในการช่วยคัดกรองความเสี่ยงการเกิดโรคจากลักษณะเล็บได้ด้วยตัวเองอย่างมีประสิทธิภาพและสะดวกรวดเร็ว การศึกษานี้ใช้ชุดข้อมูลภาพถ่ายของเล็บ 5 ประเภท ได้แก่ ภาพเล็บโรคไขข้ออักเสบ (Rheumatoid) ภาพเล็บโรคสะเก็ดเงิน (Psoriasis) ภาพเล็บโรคมะเร็งผิวหนัง (Melanoma) ภาพเล็บโรคโลหิตจาง (Anemia) และภาพเล็บปกติ (Normal) เครื่องมือในการฝึกสอนและสร้างโมเดลใช้ภาษา Python และไลบรารี TensorFlow เพื่อทำการจำแนกภาพตามกลุ่มของโรคต่างๆ ผลการทดลองจากตารางที่ 3 แสดงผลการเปรียบเทียบประสิทธิภาพของโมเดลแต่ละวิธี พบว่าโมเดลการเรียนรู้แบบ MobileNet 1.0 มีค่าความถูกต้องสูงสุดเท่ากับ 90.20% สอดคล้องกับงานวิจัยของ [8] ซึ่งพบว่าโมเดลที่สร้างขึ้นโดยใช้โครงข่ายแบบ MobileNet 1.0 มีความสามารถในการจำแนกภาพได้อย่างมีประสิทธิภาพสูงที่สุด ดังนั้นงานวิจัยนี้จึงนำโมเดลที่ได้จากโครงข่ายการเรียนรู้แบบ MobileNet 1.0 ไปใช้พัฒนาส่วนติดต่อผู้ใช้งานในรูปแบบแอปพลิเคชันบนสมาร์ตโฟนระบบปฏิบัติการแอนดรอยด์

ในขั้นตอนการตรวจสอบประสิทธิภาพของแอปพลิเคชันซึ่งวัดจากร้อยละความถูกต้องของการจำแนกภาพจำนวนประเภทละ 30 ตัวอย่าง พบว่าแอปพลิเคชันสามารถจำแนกภาพตามกลุ่มของโรคต่างๆ คิดเป็นร้อยละความถูกต้องเฉลี่ย 78.00% ผลการทดลองชี้ให้เห็นว่าแอปพลิเคชันสามารถจำแนกภาพโรคได้อย่างมีประสิทธิภาพ

ส่วนผลการประเมินค่าความถูกต้องในการจำแนกภาพแต่ละประเภท (ตารางที่ 5) พบว่าการจำแนกภาพเล็บโรคไขข้ออักเสบ มีค่าความถูกต้องสูงสุด 97.00% ส่วนผลการจำแนกภาพเล็บปกติ มีค่าความถูกต้องต่ำที่สุด ค่าความถูกต้อง 63.00%

สาเหตุที่การจำแนกภาพเล็บปกติ มีผลลัพธ์ค่าความถูกต้องค่อนข้างน้อย เนื่องจากชุดข้อมูลภาพที่นำมาฝึกสอนของภาพเล็บปกติและภาพเล็บโรคโลหิตจาง มีลักษณะสีและรูปทรงของเล็บค่อนข้างคล้ายคลึงกัน ทำให้มีบางรูปภาพของภาพเล็บปกติถูกจำแนกโรคผิดเป็นภาพเล็บโรคโลหิตจาง จึงทำให้ภาพเล็บปกติมีความถูกต้องต่ำกว่าภาพโรคอื่นๆ นอกจากนี้จากการทดสอบประสิทธิภาพแอปพลิเคชันบนสมาร์ตโฟนยังพบว่า ความถูกต้องของการจำแนกภาพลดลงจากขั้นตอนการสร้างโมเดล (ตารางที่ 4 และตารางที่ 5) สาเหตุเนื่องมาจากภาพที่นำไปใช้ในการฝึกสอนมีการควบคุมสภาวะแวดล้อม เช่น แสงพื้นหลัง แต่เมื่อถ่ายภาพจากกล้องสมาร์ตโฟน ซึ่งไม่สามารถควบคุมสภาพแวดล้อมเหมือนขั้นตอนการฝึกสอนได้ จึงทำให้การถ่ายภาพในช่วงเวลาที่มีแสงสว่างมากหรือเกิดการสะท้อนแสง จะส่งผลให้โมเดลประมวลผลภาพผิดพลาดและค่าความถูกต้องการจำแนกภาพจะลดลง ซึ่งสอดคล้องกับงานวิจัยของ [15] ที่พบว่าแสงสะท้อนจากภาพถ่ายมีผลต่อการประมวลผลภาพของโมเดลการเรียนรู้เชิงลึก

ในส่วนของการประเมินความพึงพอใจจากผู้ใช้งาน ซึ่งประกอบด้วยกลุ่มตัวอย่างจำนวน 30 คน ผลการวิเคราะห์ระดับความพึงพอใจโดยรวมสรุปได้ว่า แอปพลิเคชันที่พัฒนาขึ้นสามารถใช้งานได้มีประสิทธิภาพและมีความพึงพอใจอยู่ในระดับดี มีค่าเฉลี่ยเท่ากับ 4.09 แสดงว่าแอปพลิเคชันที่พัฒนาขึ้นนี้สามารถนำไปใช้เป็นเครื่องมือตรวจสอบปัญหาสุขภาพโดยใช้ภาพถ่ายเล็บ ซึ่งจะช่วยให้ผู้ใช้งานสามารถตรวจสอบคัดกรองความเสี่ยงโรคด้วยตนเองเบื้องต้นได้ทันทีก่อนที่จะไปพบแพทย์เพื่อวินิจฉัยโดยละเอียดต่อไป

ข้อเสนอแนะในการดำเนินงานวิจัยต่อไปในอนาคตมีดังนี้ เพื่อให้ระบบสามารถใช้งานได้อย่างหลากหลายจำเป็นต้องมีข้อมูลที่ใช้ในการฝึกสอนที่หลากหลายและมีจำนวนเพียงพอที่จะใช้แทนกลุ่มข้อมูลได้ทุกกลุ่ม ประสิทธิภาพของระบบการจำแนกภาพสามารถปรับปรุงได้ โดยการปรับปรุงกระบวนการเตรียมข้อมูล เพื่อให้โมเดลการเรียนรู้สามารถดึงคุณลักษณะเด่นของโรคให้เกิดความแตกต่างระหว่างข้อมูลโรคแต่ละประเภทให้เพิ่มมากขึ้น จะทำให้ระบบสามารถจำแนกภาพอย่างมีประสิทธิภาพเพิ่มมากยิ่งขึ้น

เอกสารอ้างอิง

- [1] M. Ngarun, "Fingernails tell disease," Bangkok: Samit Press, 2006.
- [2] P. Asavanon, "Look at the nails to tell the disease," *Chulalongkorn Hospital*, 2020. [Online]. Available: [https://chulalongkornhospital.go.th/kcmh/line/Look at the nails to tell the disease](https://chulalongkornhospital.go.th/kcmh/line/Look%20at%20the%20nails%20to%20tell%20the%20disease). [Accessed Oct. 5, 2020].
- [3] Institute of Dermatology, "Nails Disease," *Institute of Dermatology*, 2020. [Online]. Available: <https://www.hfocus.org/content/2018/01/15247> [Accessed Sep. 20, 2020].
- [4] A. Rasminsky, "What These 8 Fingernail Signs Say About Your Health," *Healthline*, 2020. [Online]. Available: <https://www.healthline.com/health/beauty-skin-care/healthy-nails#improve-nail-health> [Accessed Sep. 23, 2020].
- [5] K. Gopalakrishnan, S. Khaitan, A. Choudhary, and A. Agrawal, "Deep Convolutional Neural Networks with Transfer Learning for Computer Vision-based Data-driven Pavement Distress Detection," *Construction and Building Materials*, vol. 157, pp. 322-330, 2017.
- [6] Z. Q. Zhao, P. Zheng, S. T. Xu, and X. Wu, "Object detection with deep learning: A review," *IEEE Transactions on Neural Networks and Learning Systems*, pp. 3212-3232, 2019.
- [7] G. Lin, and W. Shen, "Research on Convolutional Neural Network based on Improved Relu - piecewise Activation Function," *Procedia Computer Science*, vol. 131, pp. 977-984, 2018.
- [8] A. Howard *et al.*, "Mobilenets: Efficient Convolutional Neural Networks for Mobile Vision Applications," *arXiv preprint arXiv:1704.04861*, 2017.
- [9] J. Orlando, E. Prokofyeva, M. Fresno and B. Blaschko, "An Ensemble Deep Learning based Approach for Red Lesion Detection in Fundus Images," *Computer Methods and Programs in Biomedicine*, vol.153, pp. 115-127, 2018.

- [10] T. Metkarunchit and K. Charoenpojvajan, "Detection of COVID-19 using Deep Learning with Images," *TNI Journal of Engineering and Technology*, vol. 8, no. 2, pp. 8-17, 2020.
- [11] C. Chin, B. Lin, and G. Wu, "An Automated Early Ischemic Stroke Detection System using CNN Deep Learning Algorithm," IEEE 8th International Conference on Awareness Science and Technology, pp. 368-372, 2017.
- [12] Q. Que, Z. Tang, et al., "CardioXNet: Automated Detection for Cardiomegaly Based on Deep Learning," 40th Annual International Conference of the IEEE Engineering in Medicine and Biology Society, pp. 612-615, 2018.
- [13] N. Lamsamut and S. Valuvanathorn, "Stroke Disease Classification on Computerized Tomography Scan Images Using Convolutional Neural Network," The 17th National Conference on Computing and Information Technology, pp. 43-48, 2021.
- [14] S. Phimpisan and N. Sriwiboon, "Image Processing for Fundus Image Classification using Deep Learning," *Journal of Information Science and Technology*, vol. 10, no. 2, pp. 19-25, 2020.
- [15] R. Petagon and O. Pantho, "Drone for Detecting Forest Fires using Deep Learning Technique," *Sripatum Review of Science and Technology*, vol. 12, no. 1, pp. 65-78, 2020.



ระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยการประมวลผล
ภาพร่วมกับไลบรารีการรู้จำใบหน้า

**Face Registration and Student's Class Attending Monitoring
System using Image Processing with Library Face Recognition**

เอกรัตน์ สุขสุคนธ์*

*Aekkarat Suksukont**

สาขาเทคโนโลยีคอมพิวเตอร์และนวัตกรรม คณะวิทยาศาสตร์และเทคโนโลยี วิทยาลัยเซาท์อีสต์บางกอก

Program in Computer Technology and Innovation, Faculty of Science and Technology,

Southeast Bangkok College

Received: October 01, 2021; Revised: November 02, 2021; Accepted: November 24, 2021; Published: December 28, 2021

ABSTRACT – This research present face registration and student's class attending monitoring system using image processing with library face recognition. It will be useful to check-in all the student's class attending and also to use the modern technology to register for classroom attendance. The effectiveness of the student's face recognition system was divided into two parts. In the term of comparison between the students' face images with the database image. Then these image were pass through to Pre-Processing to create a database of student's faces and was compared with those database. In the term using the real-time face images (registration images), the registration images were created using the face recognition image system. The experiment results shown that, in the case of random 100 face image per a student gave the best recognition performance accuracy as 92%. Also, this designed system can view historical record data and transfer it to be document file.

KEYWORDS: Image processing, Face detection, Face recognition, Computer vision, Library

บทคัดย่อ -- งานวิจัยนี้นำเสนอระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยหลักการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้า ในงานวิจัยนี้ได้ใช้กระบวนการทดสอบประสิทธิภาพของระบบรู้จำใบหน้าของนักศึกษาโดยการเปรียบเทียบภาพใบหน้าของนักศึกษากับภาพต้นฉบับด้วยการลงทะเบียนในระบบ จากนั้นเข้าสู่กระบวนการเตรียมภาพสำหรับการทดลอง เพื่อสร้างฐานข้อมูลภาพใบหน้าของนักศึกษา และนำมาเปรียบเทียบกับภาพใบหน้าจากฐานข้อมูล โดยทดลองใช้อัตราการสุ่มของจำนวนภาพใบหน้าที่แตกต่างกัน จากผลการทดลอง พบว่า การสุ่มบันทึกภาพใบหน้าของนักศึกษา จำนวน 100 ภาพต่อนักศึกษา 1 คน ได้ผลการทดลองที่แม่นยำที่สุด โดยระบบรู้จำใบหน้าของนักศึกษาสามารถระบุตัวตนของนักศึกษาได้ถูกต้องแม่นยำอยู่ที่ 92% และระบบนี้สามารถดูข้อมูลย้อนหลังของนักศึกษาในการเข้าห้องเรียน และสามารถนำเอาข้อมูลออกมาใช้งานในรูปแบบของไฟล์เอกสาร

คำสำคัญ: การประมวลผลภาพ, การตรวจจับใบหน้า, การรู้จำใบหน้า, คอมพิวเตอร์วิชั่น, ไลบรารี

*Corresponding Author: ekk_ele@hotmail.com

1. บทนำ

ในปัจจุบันการลงเวลาเข้าชั้นเรียนของนักศึกษามีหลากหลายรูปแบบที่แตกต่างกันออกไป ตัวอย่างเช่น การขานชื่อ การส่งใบลงเวลาให้กัน การใช้แถบแม่เหล็ก รวมถึงการใช้อุปกรณ์ลู่วิ่ง ซึ่งแต่ละรูปแบบได้มีปัญหาที่เกิดขึ้นแตกต่างกัน ได้แก่ การลงเวลาเข้าชั้นเรียนแทนกัน การฝากบัตรประจำตัวสแกน การลงเวลาเข้าชั้นเรียนออนไลน์แต่ไม่เข้าชั้นเรียน จากปัญหาข้างต้นผู้วิจัยได้พัฒนาระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้า ซึ่งเทคโนโลยีเกี่ยวกับการประมวลผลสัญญาณภาพในปัจจุบันนั้นเป็นที่นิยมใช้งานและพัฒนาขึ้นอย่างมาก รองรับได้หลายภาษา และรองรับอุปกรณ์ที่เป็นฮาร์ดแวร์มากยิ่งขึ้น จากงานวิจัยของ [1] การค้นหาพื้นที่ใบหน้าบุคคลและวัตถุแปลกปลอมบริเวณดวงตา โดยใช้เทคนิคโมเดลสี RGB ร่วมกับ HSV และใช้เทคนิคการแบ่งส่วนใบหน้าในการค้นหาใบหน้าบุคคล จากผลการทดลอง พบว่า สามารถค้นหาพื้นที่ใบหน้าได้ 86 เปอร์เซ็นต์ จากงานวิจัยของ [2] เป็นการตรวจสอบนักศึกษาเข้าเรียนโดยใช้เทคนิคการรู้จำใบหน้า โดยใช้การเปรียบเทียบ 3 วิธี คือ เทคนิค Eigen Face Recognition เทคนิค Fisher Face Recognition และเทคนิค Local Binary Pattern Histograms Recognition: LBPH เพื่อนำไปใช้ในระบบที่พัฒนาซึ่งจากการทดลองพบว่า วิธี LBPH มีความถูกต้องในการระบุตัวตนได้ 94.21 เปอร์เซ็นต์ จากงานวิจัยของ [3] เป็นการใช้เทคนิค LBPH ในการสกัดคุณลักษณะเด่นของใบหน้า การทดลองข้อมูลการทดสอบจำนวน 268 จาก 342 ได้รับการยอมรับว่าถูกต้องและแม่นยำ ส่งผลให้มีความแม่นยำ 78.4 เปอร์เซ็นต์ จากงานวิจัยของ [4] ใช้กระบวนการสร้างฐานข้อมูลรูปภาพและการพัฒนาระบบในลักษณะเว็บแอปพลิเคชันเพื่อนำไปใช้งาน โดยใช้ Haar-Like Feature ในการตรวจจับใบหน้า และใช้ LBPH ในการรู้จำใบหน้า โดยเบื้องต้นมีความแม่นยำของการรู้จำใบหน้าที่ 48 เปอร์เซ็นต์ จากงานวิจัยของ [5] ได้นำเสนอการสกัดคุณลักษณะของใบหน้าโดยใช้เทคนิค Local Binary Pattern: LBP เพื่อดึงคุณลักษณะเด่นของใบหน้าพร้อมกับประเมินผลของการตรวจจับใบหน้า จากการวิจัยพบว่า เมื่อใช้เทคนิคนี้ในการตรวจจับใบหน้าสามารถตรวจจับได้เป็นอย่างดี สามารถนำค่าที่ได้ไปใช้ในการรู้จำบุคคลและสามารถประยุกต์ใช้งานได้กับวัตถุชนิดอื่นได้ งานวิจัยของ [6] ได้ทำการรู้จำใบหน้าโดยใช้ค่า

ของ Eigen Faces เพื่อจำแนกลักษณะจำเพาะบุคคล โดยที่อัตราความผิดพลาดอยู่ที่ค่าสัมประสิทธิ์ของการกำหนดพารามิเตอร์ที่มีความสัมพันธ์กับคุณลักษณะของใบหน้า งานวิจัยของ [7] และ [8] ประยุกต์ใช้งานกล้องวิดีโอตรวจจับวัตถุในลักษณะที่แตกต่างกัน งานวิจัยของ [9] และ [10] เป็นการเลือกใช้คุณสมบัติของไลบรารี OpenCV ในการสกัดคุณลักษณะเด่นของภาพและบัตรประจำตัวพนักงาน โดยกำหนดอัลกอริทึมเพื่อเปรียบเทียบค่าต่าง ๆ จากการวิจัย พบว่า ไลบรารีนี้มีคุณสมบัติเด่นที่ทำให้สามารถนำไปใช้เปรียบเทียบและจับคู่ขึ้นงานได้ดี

จากงานวิจัยที่กล่าวมาข้างต้นนั้น พบว่า ระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยหลักการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้า มีการประยุกต์ใช้ไลบรารี OpenCV เข้ามาช่วยในการเพิ่มประสิทธิภาพของระบบ ซึ่งส่วนสำคัญของงานวิจัยนี้อยู่ที่การพัฒนาเทคนิคเพื่อเพิ่มประสิทธิภาพของการประมวลผลภาพเพื่อการรู้จำใบหน้า เพื่อเพิ่มความแม่นยำในการระบุใบหน้าของแต่ละบุคคลเพื่อให้ระบบมีความเสถียรมากยิ่งขึ้น

2. ระเบียบวิธีวิจัย

ระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยหลักการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้าของนักศึกษาวิทยาลัยเซาธ์อีสท์บางกอก จำนวน 100 คน โดยเริ่มจากการออกแบบระบบที่ใช้ในการทดลอง จากนั้นทำการเก็บภาพใบหน้าของนักศึกษาเพื่อสร้างฐานข้อมูลภาพใบหน้า โดยใช้ไลบรารีการรู้จำใบหน้าในการสกัดคุณลักษณะเด่นของใบหน้า เพื่อนำมาใช้เปรียบเทียบภาพใบหน้าที่ฐานข้อมูลใบหน้าสำหรับการระบุตัวตนของนักศึกษาที่ถูกต้องและแม่นยำ

2.1 เครื่องมือที่ใช้ในการวิจัย

2.1.1 การประเมินประสิทธิภาพของระบบ

ในการเปรียบเทียบประสิทธิภาพของระบบ แบ่งออกเป็น 2 ส่วน คือ

- 1) การเปรียบเทียบฐานข้อมูลใบหน้าของนักศึกษากับภาพต้นฉบับ
- 2) การเปรียบเทียบฐานข้อมูลใบหน้าด้วยการลงทะเบียนใบหน้า

2.1.2 การประเมินความพึงพอใจของผู้ใช้ระบบ

แบบประเมินผลความพึงพอใจของระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้าของนักศึกษาวិทยาลัยเซาริฮิสต์บางกอก โดยใช้แบบสอบถาม ซึ่งแบ่งเป็น 3 ส่วน ดังนี้

1) ข้อมูลพื้นฐานเกี่ยวกับผู้ตอบแบบสอบถาม เพื่อประเมินความพึงพอใจของระบบเป็นแบบตรวจสอบรายการ

2) การกำหนดรูปแบบของคำถามเป็นคำถามวัดประสิทธิภาพของเครื่องมือวิจัย โดยมีผู้ทดลองใช้ระบบฯ ประเมินความพึงพอใจ จำนวน 10 ท่าน โดยเป็นคณาจารย์ในคณะวิทยาศาสตร์และเทคโนโลยี วิทยาลัยเซาริฮิสต์บางกอก

ระดับการให้คะแนนเฉลี่ยในแต่ละระดับชั้นใช้เทคนิคในการคำนวณช่องกว้างของชั้นของข้อมูลที่มีต่อเนื่อง โดยสามารถแปลความหมายของระดับความสำคัญของคะแนนได้ดังนี้

$$\text{พิสัย} = \frac{\text{คะแนนสูงสุด} - \text{คะแนนต่ำสุด}}{\text{จำนวนชั้น}} \quad (1)$$

$$= \frac{5 - 1}{5} = 0.8$$

คะแนนเฉลี่ย 4.21 – 5.00 หมายถึง ผู้ใช้มีความพึงพอใจต่อระบบมากที่สุด

คะแนนเฉลี่ย 3.41 – 4.20 หมายถึง ผู้ใช้มีความพึงพอใจต่อระบบมาก

คะแนนเฉลี่ย 2.61 – 3.40 หมายถึง ผู้ใช้มีความพึงพอใจต่อระบบปานกลาง

คะแนนเฉลี่ย 1.81 – 2.60 หมายถึง ผู้ใช้มีความพึงพอใจต่อระบบน้อย

คะแนนเฉลี่ย 1.00 – 1.80 หมายถึง ผู้ใช้มีความพึงพอใจต่อระบบน้อยที่สุด

3) ข้อคิดเห็นและข้อเสนอแนะเพิ่มเติมเป็นแบบคำถามปลายเปิด

2.2 การหาขอบภาพ (Edge Detection)

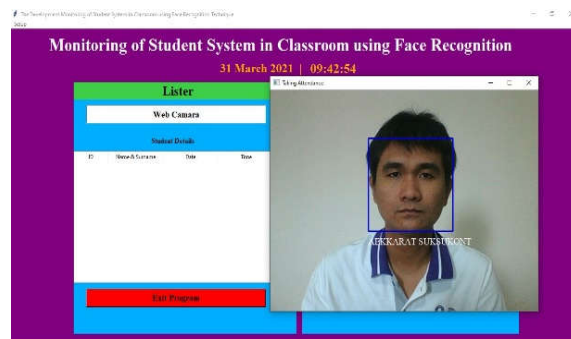
การหาขอบภาพโดยวิธีโซเบล เป็นการหาขอบภาพโดยใช้เทมเพลตขนาด 3x3 สองเทมเพลต โดยเทมเพลตแรกจะใช้หาค่าความแตกต่างในแนวนอน (X_{diff}) และค่าความแตกต่างในแนวตั้ง (Y_{diff}) ดังรูปที่ 1

$$X_{diff} = \begin{bmatrix} -1 & 0 & 1 \\ -2 & 0 & 2 \\ -1 & 0 & 1 \end{bmatrix} \quad Y_{diff} = \begin{bmatrix} 1 & 2 & 3 \\ 0 & 0 & 0 \\ -1 & -2 & -1 \end{bmatrix}$$

รูปที่ 1. การหาขอบภาพด้วยวิธีโซเบล

2.3 วิธีการค้นหาใบหน้า (Face Detection)

หลักการค้นหาใบหน้า คือ การตรวจหาบริเวณที่เป็นใบหน้าของบุคคลในภาพ หนึ่ง ซึ่งถูกออกแบบมาให้ทำการเปรียบเทียบใบหน้าบุคคลที่ถูกเก็บไว้เป็นฐานข้อมูลซึ่งข้อมูลใบหน้าที่ใช้ในขั้นตอนการสร้างแม่แบบหรือข้อมูลค้นฉบับและขั้นตอนการเปรียบเทียบอาจแตกต่างกันไป โดยทั่วไปภาพจะประกอบไปด้วยส่วนที่เป็นใบหน้าและส่วนที่เป็นพื้นหลัง ฟังก์ชันสำหรับการตรวจจับใบหน้าที่มีอยู่ในไลบรารี OpenCV โดยที่ผลลัพธ์ของการค้นหาใบหน้าจะได้เป็นตำแหน่ง x, y ในภาพของจุดซ้ายบน (x1, y1) และจุดขวาล่าง (x2, y2) ของสี่เหลี่ยมที่ล้อมบริเวณใบหน้า ดังรูปที่ 2



รูปที่ 2. การค้นหาภาพใบหน้าบุคคล

2.4 Open Source Computer Vision Library

ไลบรารี OpenCV นำมาใช้สำหรับการประมวลผลสัญญาณภาพและงานทางด้านการมองเห็นของคอมพิวเตอร์ (Computer Vision) ไลบรารีนี้ถูกพัฒนาขึ้นด้วยภาษา C และ C++ และยังมี Interface ที่ไว้เชื่อมต่อกับ Tool อื่น ๆ และรองรับได้หลายภาษา เช่น Python, Ruby, Matlab เป็นต้น นอกจากนี้ OpenCV ยังเป็นไลบรารีที่ถูกสร้างขึ้นเพื่อให้ผู้ใช้หรือนักพัฒนาสามารถใช้ฟังก์ชันในไลบรารีมาพัฒนาชิ้นงานที่มีความซับซ้อนโดยใช้เวลาเพียงไม่นาน OpenCV ประกอบด้วย Data Structure และ Algorithm

3. ผลการศึกษา

3.1 การเก็บข้อมูลประชากร

ในการวิจัยครั้งนี้ ผู้วิจัยดำเนินการเก็บรวบรวมข้อมูลใบหน้าของนักศึกษาวิทยาลัยเซาธ์อีสท์บางกอก จำนวน 100 คน โดยทำการถ่ายภาพระยะห่างจากนักศึกษา 1.50 เมตร ภาพใบหน้าที่รับเข้ามาเป็นภาพสีแบบ RGB โดยทำการจัดเก็บภาพใบหน้าเพื่อสร้างเป็นฐานข้อมูลภาพใบหน้าของนักศึกษา ซึ่งลักษณะของภาพจะมีภาพพื้นหลังที่เหมือนกัน รวมทั้งเห็นบริเวณดวงตาทั้งสองข้าง จมูก และปากที่ชัดเจน ภาพที่ได้เป็นการถ่ายภาพในช่วงเวลาที่ใกล้เคียงกัน โดยใช้แสงธรรมชาติภายในห้องเรียนมีขนาดและความละเอียดของภาพที่เท่ากัน

เมื่อทำการจัดเก็บภาพใบหน้าในลักษณะที่แตกต่างกัน 5 อริยาบทแล้ว จากนั้นเข้าสู่กระบวนการเพื่อหาตำแหน่งของใบหน้าจากภาพต้นฉบับ แล้วจึงทำการรู้จำใบหน้า จากนั้นสกัดคุณลักษณะเด่นของใบหน้านั้น ๆ ออกมา และนำไปสร้างเป็นฐานข้อมูลภาพ ใบหน้า เมื่อระบบนำเข้าภาพใบหน้าของนักศึกษาจากภาพถ่ายหรือวิดีโอก็จะนำมาเปรียบเทียบกับฐานข้อมูลภาพใบหน้าของนักศึกษา โดยในการทดลองนี้จะสามารถหาประสิทธิภาพความแม่นยำของระบบและค่าความผิดพลาดของระบบได้

3.2 กระบวนการเตรียมภาพใบหน้า

เมื่อได้ภาพต้นฉบับ ใบหน้าของนักศึกษาแล้ว จากนั้นทำการครอบ (Crop) เฉพาะบริเวณใบหน้าของนักศึกษา เพื่อสกัดลักษณะเฉพาะของใบหน้าของแต่ละบุคคล หลังจากนั้นทำการแปลงภาพสีเป็นภาพระดับสีเทา ดังรูปที่ 3

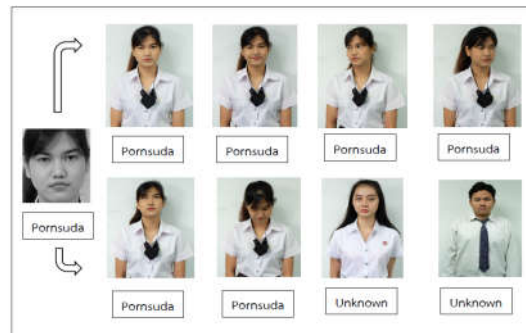


รูปที่ 3. กระบวนการครอบภาพใบหน้าและแปลงภาพสีเป็นภาพสีเทา

3.3 การสร้างฐานข้อมูลภาพใบหน้า

3.3.1 การสร้างฐานข้อมูลภาพใบหน้าของนักศึกษาจากภาพต้นฉบับ

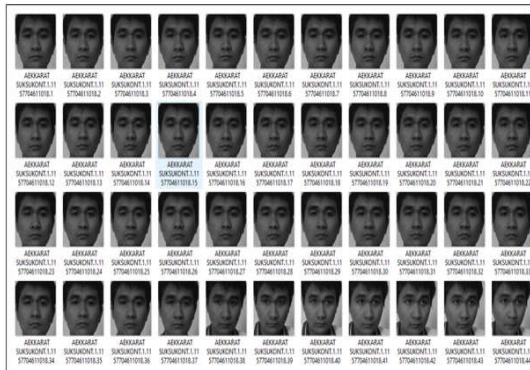
เมื่อนำภาพใบหน้าของนักศึกษามาผ่านกระบวนการเตรียมภาพใบหน้า จากนั้นนำภาพที่ได้มาสร้างเป็นฐานข้อมูลภาพใบหน้าของนักศึกษา ในที่นี้จะนำภาพใบหน้าของนักศึกษาในอริยาบทที่แตกต่างกัน ได้แก่ ภาพหน้าตรง ภาพอมยิ้ม ภาพมองด้านซ้าย ภาพมองด้านขวา ภาพก้มหน้า และภาพเงยหน้า มาเปรียบเทียบกับภาพต้นฉบับที่สร้างเป็นฐานข้อมูลไว้แล้ว โดยใช้ภาษาไพธอนร่วมกับ OpenCV ในการวิเคราะห์และระบุตัวตนของนักศึกษา ซึ่งในขั้นตอนนี้จะใช้ไลบรารี Face Recognition เข้ามาช่วยในการรู้จำใบหน้าของนักศึกษา ดังรูปที่ 4 ระบบสามารถแสดงรายชื่อและระบุตัวตนของนักศึกษาค้นนั้น ๆ ได้แต่หากนำภาพของบุคคลอื่นเข้ามาเปรียบเทียบ ระบบจะแสดงคำว่า “Unknown” คือ ไม่รู้จักบุคคลคนนั้น



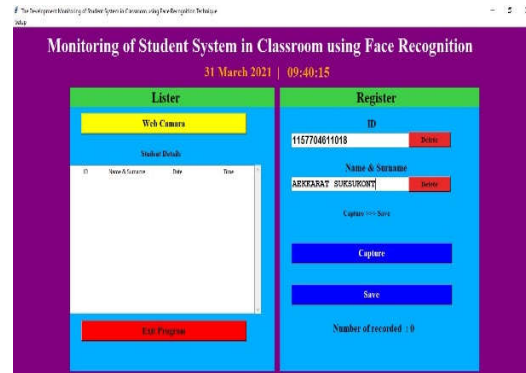
รูปที่ 4. ผลการเปรียบเทียบ ใบหน้าและระบุตัวตนของนักศึกษา

3.3.2 การสร้างฐานข้อมูลใบหน้าด้วยการลงทะเบียนใบหน้า

เริ่มจากให้นักศึกษาทำการลงทะเบียนโดยการใส่รายละเอียดประจำตัวของตนเอง ได้แก่ รหัสประจำตัวนักศึกษา ชื่อและนามสกุล จากนั้นจะทำการเปิดกล้องวิดีโอเพื่อทำการจับภาพใบหน้าและบันทึกภาพใบหน้าของนักศึกษาเพื่อสร้างเป็นฐานข้อมูลภาพใบหน้าของนักศึกษา ในที่นี้จะการจับภาพใบหน้าและบันทึกภาพของนักศึกษา จำนวน 100 ภาพต่อนักศึกษา 1 คน เพื่อให้ได้ฐานข้อมูลภาพที่มากพอสำหรับการรู้จำและการนำมาเปรียบเทียบเพื่อระบุตัวตนแสดงรายชื่อของนักศึกษาที่มีความแม่นยำมากที่สุด ดังรูปที่ 5



รูปที่ 5. ข้อมูลภาพใบหน้าของนักศึกษาที่ได้จากการลงทะเบียน



รูปที่ 7. หน้าต่างของโปรแกรมลงทะเบียน

3.4 การสร้างโปรแกรม

ระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียน ด้วยหลักการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้า ในที่นี้ใช้โปรแกรม Visual Studio Code สำหรับเขียนภาษาไพธอน ร่วมกับไลบรารีอื่น ๆ สำหรับไลบรารีในการรู้จำที่สำคัญสำหรับ ภาษาไพธอนติดตั้ง ดังรูปที่ 6 และเมื่อทำการติดตั้งไลบรารี เรียบร้อยแล้ว จึงทำการสร้างโปรแกรมลงทะเบียนใบหน้าได้ หน้าต่างของโปรแกรม ดังรูปที่ 7

Package	Version
altgraph	0.17.2
click	7.1.2
cmake	3.18.4.post1
cycler	0.10.0
dlib	19.8.1
face-recognition	1.3.0
face-recognition-models	0.3.0
future	0.18.2
importlib-metadata	4.8.1
kiwisolver	1.3.1
matplotlib	2.0.0
numpy	1.19.4
opencv-contrib-python	4.4.0.46
opencv-python	4.4.0.46
pandas	1.1.4
pefile	2021.9.3
Pillow	8.0.1
pip	21.0.1
pyinstaller	4.5.1
pyinstaller-hooks-contrib	2021.3
pyarsing	2.4.7
python-dateutil	2.8.1
pytz	2020.4
pywin32-ctypes	0.2.0
setuptools	40.6.2
six	1.15.0
typing-extensions	3.10.0.2
zipp	3.6.0

รูปที่ 6. ไลบรารีที่ใช้สำหรับภาษาไพธอน

ในการตรวจสอบรายชื่อของนักศึกษาเข้าชั้นเรียน สามารถ เข้าไปในไฟล์ของโปรแกรมระบบลงทะเบียนใบหน้าและ ตรวจสอบนักศึกษาเข้าห้องเรียนด้วยการประมวลผลภาพ ร่วมกับไลบรารีการรู้จำใบหน้า จากไฟล์ชื่อ “Attendance” จะ พบไฟล์นามสกุล .CSV สามารถนำมาเปิดใช้งานด้วยโปรแกรม Microsoft Excel โดยที่ระบบจะทำการบันทึกเป็น วันต่อวัน สามารถตรวจสอบข้อมูลย้อนหลังได้ ดังรูปที่ 8

	A	B	C	D	E	F
1	Id	Name		Date	Time	
2						
3	115770461101	AEKKARAT SUKSUKONT		5/10/2021	17:41:30	
4						
5						

รูปที่ 8. รายชื่อนักศึกษาที่เข้าห้องเรียน

3.5 ผลการทดลองการเปรียบเทียบฐานข้อมูลใบหน้าของ นักศึกษากับภาพต้นฉบับ

ในที่นี้จะนำภาพใบหน้าของนักศึกษาในอริยาบทที่แตกต่างกัน ได้แก่ ภาพหน้าตรง ภาพอมยิ้ม ภาพมองด้านซ้าย ภาพมอง ด้านขวา ภาพก้มหน้า และภาพเงยหน้า มาเปรียบเทียบกับภาพ ต้นฉบับที่สร้างเป็นฐานข้อมูลไว้โดยใช้ภาษาไพธอน ร่วมกับ OpenCV ในการวิเคราะห์และระบุตัวตนของนักศึกษา ผล ปรากฏดังตารางที่ 1 และตารางที่ 2

ตารางที่ 1. ผลการเปรียบเทียบภาพใบหน้าของนักศึกษาในสภาพแวดล้อมที่ภาพพื้นหลังไม่มีสวคล้ายกับภาพต้นฉบับ

นักศึกษา	จำนวนนักศึกษา (คน)	ความแม่นยำ (คน)	ความผิดพลาด (คน)
เพศชาย	56	55	1
เพศหญิง	44	41	3
รวม	100	96	4

จากตารางที่ 1 เมื่อนำภาพของนักศึกษา จำนวน 100 คน มาเปรียบเทียบกับฐานข้อมูลภาพใบหน้าของนักศึกษาที่ภาพพื้นหลังไม่มีสวคล้ายกับภาพต้นฉบับ พบว่า ระบบรู้จำใบหน้าของนักศึกษาสามารถระบุตัวตนของนักศึกษาได้ถูกต้องแม่นยำถึง 96 คน และเกิดความผิดพลาดที่ไม่สามารถระบุตัวตนได้ จำนวน 4 คน

ตารางที่ 2. ผลการเปรียบเทียบภาพใบหน้าของนักศึกษาในสภาพแวดล้อมที่ภาพพื้นหลังมีสวคล้ายกับภาพต้นฉบับ

นักศึกษา	จำนวนนักศึกษา (คน)	ความแม่นยำ (คน)	ความผิดพลาด (คน)
เพศชาย	56	52	4
เพศหญิง	44	38	6
รวม	100	90	10

จากตารางที่ 2 เมื่อนำภาพของนักศึกษา จำนวน 100 คน มาเปรียบเทียบกับฐานข้อมูลภาพใบหน้าของนักศึกษาที่ภาพพื้นหลังมีสวคล้ายกับภาพต้นฉบับ พบว่า ระบบรู้จำใบหน้าของนักศึกษาสามารถระบุตัวตนของนักศึกษาได้ถูกต้องแม่นยำถึง 90 คน และเกิดความผิดพลาดที่ไม่สามารถระบุตัวตนได้ จำนวน 10 คน ซึ่งความผิดพลาดเกิดจากภาพที่นำมาเปรียบเทียบไม่ใช่คนเดียวกันกับบุคคลต้นฉบับหรือไม่ได้มีการสร้างฐานข้อมูลภาพใบหน้าของนักศึกษา และภาพนั้นอาจมีตำแหน่งของดวงตาและระยะของตำแหน่งจมูกและตำแหน่งปากที่ไม่ชัดเจน จึงทำให้ระบบไม่สามารถประมวลผลและระบุตัวตนของนักศึกษาได้

3.6 ผลการเปรียบเทียบฐานข้อมูลใบหน้าด้วยการลงทะเบียนใบหน้า

กระบวนการสร้างฐานข้อมูลใบหน้าด้วยการลงทะเบียนใบหน้า โดยให้นักศึกษาทำการลงทะเบียน โดยการใส่รายละเอียดประจำตัวของตนเอง ได้แก่ รหัสประจำตัวนักศึกษา ชื่อและนามสกุล จากนั้นจะทำการเปิดกล้องเพื่อทำการจับภาพใบหน้า และบันทึกภาพใบหน้าของนักศึกษาเพื่อสร้างเป็นฐานข้อมูลภาพใบหน้าของนักศึกษา

ในกระบวนการสร้างฐานข้อมูลภาพใบหน้าของนักศึกษา จะใช้อัตราการสุ่มบันทึกภาพใบหน้า จำนวน 100 ภาพจำนวนนักศึกษา 1 คน ดังนั้นจะได้ภาพของฐานข้อมูลใบหน้า ทั้งหมด 10,000 ภาพ ดังตารางที่ 3

ตารางที่ 3. ผลการเปรียบเทียบฐานข้อมูล ใบหน้าด้วยการลงทะเบียนใบหน้า

นักศึกษา	จำนวนนักศึกษา (คน)	ความแม่นยำ (คน)	ความผิดพลาด (คน)
เพศชาย	56	53	3
เพศหญิง	44	39	5
รวม	100	92	8

จากตารางที่ 3 ความแม่นยำในการระบุตัวตนของนักศึกษา อยู่ที่ 92 คน และความผิดพลาดอยู่ที่ 8 คน โดยใช้เวลาในการจับภาพสร้างฐานข้อมูลภาพใบหน้าเป็นระยะเวลา 50 วินาที เหตุที่ระบบไม่สามารถระบุตัวตนของนักศึกษาได้ในกรณีที่การจับภาพของกล้องวิดีโอมีการเคลื่อนไหวของตัวนักศึกษาอยู่ตลอดเวลาทำให้การจับภาพที่ไม่นิ่งพอ ส่งผลต่อการประมวลผลของข้อมูล การเก็บข้อมูลของนักศึกษาที่น้อยเกินไปและในกรณีที่เก็บภาพใบหน้าของนักศึกษามากเกินไปจนทำให้ข้อมูลภาพใบหน้าของศึกษาไปผสมกับใบหน้าของนักศึกษาคนอื่นที่มีลักษณะของโครงหน้าที่คล้ายคลึงกัน รวมไปถึงกรณีที่นักศึกษาบางคนในขณะที่บันทึกภาพไม่ใส่แว่นตาแต่เวลาที่จะทำการเข้าห้องเรียนนั้นสวมแว่นตาทำให้ระบบประมวลผลว่านักศึกษาเป็นนักศึกษาคนละคน

3.7 ผลประเมินความพึงพอใจของผู้ใช้ระบบ

จากการทดสอบทดสอบประสิทธิภาพความแม่นยำในการรู้จำใบหน้าของนักศึกษาแล้ว ผู้วิจัยได้ทำการสร้างแบบประเมินความพึงพอใจของผู้ใช้งาน โดยมีคณาจารย์ของคณะวิทยาศาสตร์และเทคโนโลยี วิทยาลัยเซาธ์อีสท์บางกอก เป็นผู้ตอบแบบสอบถาม ผลจากการประเมินความพึงพอใจของผู้ใช้ระบบฯ ดังตารางที่ 4

ตารางที่ 4. ผลการประเมินความพึงพอใจของผู้ทดลองใช้ระบบ

รายการ	ค่าเฉลี่ย	แปลผล
1. ความสอดคล้องของระบบกับจุดประสงค์ในการนำไปใช้งานจริง	4.20	พึงพอใจมาก
2. ประโยชน์ของระบบที่มีต่อองค์กรหรือสังคม	4.60	พึงพอใจมากที่สุด
3. ประสิทธิภาพและความแม่นยำของระบบ	3.80	พึงพอใจมาก
4. ความสะดวกสบายของผู้ใช้งาน	4.30	พึงพอใจมากที่สุด
5. ความปลอดภัยในการใช้ระบบ	4.60	พึงพอใจมากที่สุด
6. ความเสถียรของซอฟต์แวร์	4.00	พึงพอใจมาก

4. สรุปผลการวิจัย

4.1 สรุปผลการทดลอง

ระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้าเป็นการศึกษาแนวคิดหลักการ ทฤษฎีที่เกี่ยวข้องกับการรู้จำใบหน้า ซึ่งผลที่ได้จากการทดลองในการเปรียบเทียบภาพใบหน้าของนักศึกษา กับภาพต้นฉบับ ความแม่นยำในการรู้จำอยู่ที่ 96% และการเปรียบเทียบฐานข้อมูลใบหน้าด้วยการลงทะเบียนใบหน้านั้นมีความแม่นยำอยู่ที่ 92% เหตุที่การเปรียบเทียบภาพใบหน้าที่ของนักศึกษา กับภาพต้นฉบับมีความแม่นยำที่สูงกว่าเพราะภาพที่นำมาเปรียบเทียบนั้นเป็นภาพนิ่ง ส่วนการเปรียบเทียบฐานข้อมูลใบหน้าด้วยการลงทะเบียนใบหน้านั้นเป็นภาพจากกล้องวิดีโอแบบเวลาจริง โดยระบบสามารถดูข้อมูลของ

นักศึกษาเข้าห้องเรียนแบบย้อนหลังและสามารถดึงข้อมูลออกมาใช้งานได้อีกด้วย

4.2 ข้อเสนอแนะ

ระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยการประมวลผลภาพร่วมกับไลบรารีการรู้จำใบหน้าสามารถประยุกต์ใช้งานกับฮาร์ดแวร์หรือคอมพิวเตอร์ขนาดเล็กแบบพกพาหรือสามารถนำไปประยุกต์ใช้กับงานทางด้านอุตสาหกรรมหรืองานเกี่ยวกับด้านธุรกิจ การตรวจสอบรายชื่อเข้า-ออกของพนักงานบริษัท เป็นต้น ทั้งนี้การเพิ่มประสิทธิภาพความแม่นยำของระบบอาจนำเทคนิคแบบจำลองสี หรือการตรวจจับตำแหน่งของใบหน้าแบบอื่น ๆ มาใช้งานร่วมด้วย เพื่อให้ระบบมีการแยกแยะและระบุตัวตนของบุคคลคนนั้น ซึ่งอาจทำให้สามารถเพิ่มประสิทธิภาพการรู้จำใบหน้าที่สูงขึ้นตามไปด้วย

5. กิตติกรรมประกาศ

งานวิจัยเรื่องนี้ ผู้วิจัยได้รับทุนภายในจากสถาบันวิจัยและพัฒนา วิทยาลัยเซาธ์อีสท์บางกอก ผู้วิจัยจึงขอขอบพระคุณวิทยาลัยเซาธ์อีสท์บางกอกที่ให้ทุนทรัพย์สนับสนุนในการวิจัยครั้งนี้ รวมทั้งคณาจารย์ทุกท่านที่ให้คำปรึกษาในด้านวิชาการ และผู้วิจัยขอขอบคุณผู้เข้าร่วมโครงการวิจัยทุกท่านที่ให้ข้อมูลและความร่วมมือเกี่ยวกับการวิจัยการวิจัยด้วยดีมาตลอด

เอกสารอ้างอิง

- [1] Suksukont A., "Face Detection and Objects on Eyes Boundary using Color Model with Image Processing," Rajamangala University of Technology Tawan-ok Research Journal, Vol.14, No.1, pp. 42-53, 2021.
- [2] Triprapin K., Naudom P. and Kongchai P., "Attendance Monitoring System with Face Recognition Technologies," Science and Technology Journal, Vol.20, No.2, pp. 92 – 105, 2018.
- [3] Suherwin, Zainuddin Z., Ilham A.A., "The Performance of Face Recognition Using the Combination of Viola-Jones," Local Binary Pattern Histogram and Euclidean

- Distance, International Conference on Informatics and Computational Sciences, pp. 1-4, 2020.
- [4] Jaturawatthana P., Phongmanawut P. and Phankokkruad M., Development of A Learning Record System with Face Detection and Recognition, Lat Krabang Information Technology Journal, Vol.15, No.1, pp. 1-11, 2017.
- [5] Deng W., Hu J. and Guo J., “Compressive Binary Patterns: Designing a Robust Binary Face Descriptor with Random-Field Eigenfilters,” IEEE Transactions on Pattern Analysis and Machine Intelligence, 41(3) pp. 758-767, 2019.
- [6] Shatnawi Y., Alsmirat M. and Al-Ayyoub M., “Face Recognition using Eigen-Faces and Extension Neural Network,” International Conference on Computer Systems and Applications, pp. 1-7, 2019.
- [7] Xiao J., Li S. and Xu Q., “Video-Based Evidence Analysis and Extraction in Digital Forensic Investigation,” Deep Learning: Security and Forensics Research Advances and Challenges Vol.7, pp. 5432-5442, 2019.
- [8] Shahbaz A. and Jo K.H., “Moving Object Detection based on Deep Atrous Spatial Features for Moving Camera,” International Symposium on Industrial Electronics, pp. 67-70, 2021.
- [9] Kushal M., Kushal Kumar B.V., Charan Kumar M.J. and Pappa M., “ID Card Detection with Facial Recognition using Tensorflow and OpenCV,” International Conference on Inventive Research in Computing Applications, pp. 742-746, 2020.
- [10] Noble F.K., “Comparison of OpenCV’s Feature Detectors and Feature Matchers,” Proceeding of International Conference on Mechatronics and Machine Vision in Practice, pp. 1-6, 2016.

Secure Mutual Authentication Protocol Based on Wireless Body Area Networks

Chalee Thammarat and Chian Techapanupreeda*

Faculty of Engineering and Technology, Mahanakorn University of Technology

Received: October 06, 2021; Revised: December 19, 2021; Accepted: December 23, 2021; Published: December 28, 2021

ABSTRACT – Data sent from wireless body area networks to healthcare professionals or doctors include sensitive information which needs to be protected from unauthorized access. A mutual authentication protocol is a security feature that can prevent man-in-the-middle and spoofing attacks. A number of mutual authentication protocols based on wireless body area networks have been proposed; however, these impose high cryptographic operation costs, energy costs, and time costs, and also lack some security properties. In this research, we propose an efficient mutual authentication protocol for secure data exchange to send personal health records from a smartphone device to a doctor. The proposed protocol leads to a reduction in the cryptographic operation, energy, and time costs, and uses fewer resources than previous protocols. Although our approach utilizes a one-way hash function rather than encryption, it still provides the necessary security properties, unlike existing protocols. We also formally verify our approach using the Scyther tool and AVISPA. The results show that the proposed protocol has been verified as being resistant to attack as designed.

KEYWORDS: Security, eHealth, Information Security, Scyther, AVISPA, PHRs, PHIs

1. Introduction

A personal health record (PHR) contains patient health information needed by health care workers [1]. A PHR may also be used for otherwise healthy people who want to record their health status. The security requirements for PHRs are confidentiality, integrity, and authentication [2, 3], as these protect against threats from attackers seeking to access personal health information, e.g., by altering, eavesdropping on, or denying health information. Many researchers have devised authentication protocols to support all essential aspects of security for wireless body area network (WBAN) data [17-20]. For example, the authors of [17] proposed a key exchange protocol for WBANs that achieved authentication between a control node and a secondary node, between a control node and a primary node, and between a secondary node and a primary node. A timestamp was utilized to guarantee the freshness of the message, and this formed the main

security protocol. Their proof of this security protocol used BAN logic. The study in [18]

Introduced a protocol to support selective authentication between nodes and key exchange in a WBAN. This protocol provided the desired security properties, and imposed light computation and communication overheads. In this approach, a random number was adopted instead of a timestamp to reduce the complexity and the cost. The BAN logic model was used to prove this security protocol. The authors of [19] devised a robust anonymous authentication protocol for healthcare applications using a wireless medical sensor network (WMSN). This was suitable for healthcare applications based on a WMSN, and offered strong security and computational efficiency. Both a one-way hash function and symmetric key encryption were applied to ensure the security of the protocol. A formal security analysis was given that used the BAN logic model. In [20], an efficient three-party authentication protocol was suggested for WBANs which used a two-hop star network topology. It made

*Corresponding Author: chalee23@gmail.com

use of three entities: a central device, a primary node and a secondary node. BAN logic and the AVISPA tool were used to provide a security proof for this protocol.

However, the approaches put forward in [17-20] do not cover all of the necessary security properties, such as mutual authentication, integrity and computational cost. In this paper, we propose a lightweight mutual authentication protocol for WBANs that can maintain the security of sensitive information. It can also be used to solve problems with existing protocols and to overcome their limitations, as a mutual authentication method for a WBAN protocol is still lacking.

This paper is organized as follows. Section 2 discusses some related works. Section 3 presents our proposed protocol, and in Section 4, we analyze the security of our approach. Section 5 compares the performance of our method with existing protocols, and Section 6 concludes the paper and suggests future work.

2. Related Works

2.1. Wireless Body Area Network (WBAN)

The application and usage of WBANs differ depending on which devices are used [4, 6-8]. For medical applications, they can provide valuable information monitoring via wireless communication [5], such as data from electrocardiogram (ECG), heart rate, or blood pressure sensors. Many researchers have presented protocols for medical body area networks. For instance, the authors of [9] designed a protocol based on lightweight identity-based encryption, to provide security and privacy for a body sensor network. In [10], a secure healthcare system for IoT was proposed, based on elliptic curve cryptography (ECC). Both of these approaches used strong encryption that consumed considerable computational resources but was suitable for body sensor networks and healthcare applications. As mentioned above, the security requirements for sending PHR data via a network are confidentiality, integrity, and authentication. Consequently, when sending PHRs via WBAN devices, these security requirements require consideration.

2.2. IEEE 802.15 Security

The security of body area networks relies on IEEE 802.15, and in this research, we will focus only on IEEE 802.15.1 (Bluetooth), IEEE 802.15.4 (ZigBee), and

IEEE 802.15.6 Task Group 6 (TG6), which are described below.

2.2.1. The IEEE 802.15 Task Group 6 is developing a communication standard that is optimized for low-power devices for operation on, in or around the human body (not limited to humans), to serve a variety of applications including medical, consumer electronics, and personal entertainment. The security of TG6 supports confidentiality, integrity and authentication but does not consider authorization and non-repudiation of the message. It is optimized for low-power devices and operation with the human body, i.e., for smartwatches or blood pressure tags [11].

2.2.2. Bluetooth (IEEE 802.15.1) is a wireless technology band at 2.4 GHz. The security of Bluetooth supports only two security properties, confidentiality and authentication, and is not applicable to body area networks. The reader can find more information on Bluetooth in [12].

2.2.3. ZigBee (IEEE 802.15.4) can be applied to create a personal area network, in which symmetric key encryption is used to secure the communication between devices. However, Zigbee supports only two security properties, which are confidentiality and authentication [13].

3. Proposed Protocol

In this section, we propose a mutual authentication protocol for a WBAN. Our protocol provides mutual authentication between P and AGW , and between D and AGW . The details of the mutual authentication process are explained below.

3.1. Notation

Our proposed protocol uses the symbols and notation given in Table 1.

Table 1. Notation used in the proposed protocol.

Symbol	Definition
P	A smartphone of a patient who owns personal health information
D	Doctors, hospital, professional care or clinical
AGW	Authentication gateway
P_{ID}	The identity of a patient
D_{ID}	The identity of the infirmity

Symbol	Definition
SK_{A-B}	Symmetric key between party A and party B
$\{M\}_{SK_{A-B}}$	A message encrypted with a symmetric key between party A and party B
$h(M)$	Hash function of message M
N_A	A nonce is issued by party A
T_A	Current timestamp created by party A

3.2. Network Model

Our network model includes body sensors and the P , D and AGW entities, as shown in Figure 1. The body sensors need to be verified, and are resource-limited devices. P acts as an intermediate node between the body sensors and AGW , and has more resources than the body sensors. It is usually a portable device such as a smartphone, tablet or notebook. AGW is rich in resources (i.e., has more resources than P), and is usually a server. We assume that the communication between all entities is flexible.

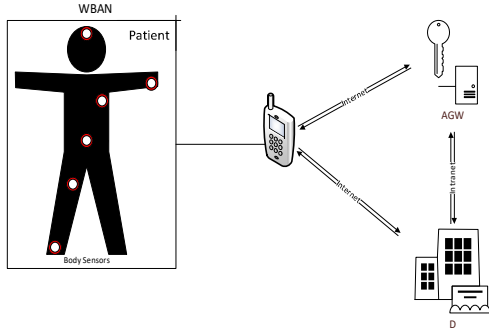


Figure 1. Network model of the proposed protocol.

3.3. Registration

In the registration phase, P and D register with AGW to share the key between P and AGW , and between D and AGW . Registration is performed via a secure channel.

3.3.1. Registration of P with AGW

P connects to AGW via a secure communication channel, and sends a request to AGW . When AGW receives this request, it generates P_{ID} and the key SK_{P-AGW} and sends these back to P . Note that SK_{P-AGW} is shared between P and AGW .

3.3.2. Registration of D with AGW

D connects with AGW via a secure communication channel, and sends a request to AGW . When AGW receives this request, it generates D_{ID} and the key SK_{D-AGW} and sends these back to D . Note that SK_{D-AGW} is shared between D and AGW .

3.3. Mutual Authentication Protocol

We propose a protocol to prevent misuse of patient health information and to protect against threats. A smartphone is used to read and check the sensors on the patient's body, such as blood pressure monitors. If the blood pressure (BPI) is equal to or more than 140SYS/90DIA mm HG, the smartphone sends the patient's identity to the emergency medical services (as an emergency case). In this case, the system sends P_{ID} , D_{ID} , and BPI directly to the emergency medical service (EMS) and calls for help. This because hypertension may threaten a patient's life, by causing a heart attack, stroke, or aneurysm. The patient therefore needs immediate treatment in order to save their life.

In contrast, if a patient's BPI is in the normal range, the system sends information to the devices when it needs a doctor or hospital to retrieve it directly when needed. The details of the proposed protocol are set out below.

M1: $D \rightarrow AGW$: D_{ID} , P_{ID} , N_D , T_D , $h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW})$

M2: $AGW \rightarrow P$: D_{ID} , P_{ID} , N_D , T_D , $h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW})$, $h(D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), SK_{P-AGW}))$

M3: $P \rightarrow AGW$: N_P , $h(N_P, h(D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), SK_{P-AGW}), SK_{P-AGW}))$

M4: $AGW \rightarrow P$: N_{AGW} , T_{AGW} , $\{SK_{P-D}, h(D_{ID}, P_{ID}, N_D, N_P, N_{AGW}, T_D, T_{AGW}, SK_{P-D}, SK_{P-AGW})\}_{SK_{P-AGW}}$

M5: $AGW \rightarrow D$: N_{AGW} , T_{AGW} , $\{SK_{P-D}, h(D_{ID}, P_{ID}, N_D, N_P, N_{AGW}, T_D, T_{AGW}, SK_{P-D}, SK_{D-AGW})\}_{SK_{D-AGW}}$

In message M1, D sends D_{ID} , P_{ID} , N_D , T_D , $h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW})$ to AGW to request a connection with P . After receiving message M1 from D , AGW uses D_{ID} , P_{ID} , N_D , T_D and SK_{D-AGW}

to compute the hash value and compares it with $h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW})$. If the two hash values are equal, AGW will continue with M2; otherwise, it terminates the connection. Note that the message contains $h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW})$ and is considered a message authentication code (MAC) that can ensure the integrity of the message.

After verifying that D is the originator of the message, AGW will send $D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), h(D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), SK_{P-AGW}))$ to P in message M2. This message can be used to ensure that AGW is the sender, since AGW possesses both the symmetric key SK_{D-AGW} and SK_{P-AGW} . Once P has received message M2 from AGW , it will check the correctness of the message by checking the hash value of $D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), h(D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), SK_{P-AGW}))$. If this is correct, P sends $N_P, h(N_P, h(D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), SK_{P-AGW}), SK_{P-AGW})$ to AGW in message M3. Otherwise, P terminates the connection.

When AGW has received message M3 and has checked that the message was sent from P , it will send a nonce, the timestamp of AGW , the shared SK_{P-D} key and $h(D_{ID}, P_{ID}, N_D, N_P, N_{AGW}, T_D, T_{AGW}, SK_{P-D}, SK_{P-AGW})$, encrypted with SK_{P-AGW} , to P in message M4.

AGW then sends a nonce, the timestamp of AGW , the shared SK_{P-D} key and $h(D_{ID}, P_{ID}, N_D, N_P, N_{AGW}, T_D, T_{AGW}, SK_{P-D}, SK_{P-AGW})$, encrypted with SK_{D-AGW} , to D in message M5. The goal of this step is to send the shared symmetric key SK_{P-D} to P and D to allow them to exchange personal health information via a secure channel.

It can be seen that the proposed protocol ensures mutual authentication between P and AGW , and between D and AGW . Each message in the proposed protocol can be used to identify the sender of the original message. We use only symmetric cryptographic operations, a MAC and a hash function to provide mutual authentication, and this results in lightweight protocol that is suitable for a WBAN.

4. Security Analysis

4.1. Informal Security Analysis

In this section, we present a security analysis of the proposed protocol to prove that it provides the necessary security, as follows:

4.1.1. Mutual authentication: This is an important security property that is used to identify the sender and the receiver of the messages. The proposed protocol deploys a MAC to provide mutual authentication between entities, which can expressed as in the message below:

M2: $AGW \rightarrow P$: $D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), h(D_{ID}, P_{ID}, N_D, T_D, h(D_{ID}, P_{ID}, N_D, T_D, SK_{D-AGW}), SK_{P-AGW}))$

AGW cannot deny sending the original message to P , as only AGW possesses both the symmetric keys SK_{D-AGW} and SK_{P-AGW} . Hence, only AGW can construct this message or the original message.

M5: $AGW \rightarrow D$: $N_{AGW}, T_{AGW}, \{SK_{P-D}, h(D_{ID}, P_{ID}, N_D, N_P, N_{AGW}, T_D, T_{AGW}, SK_{P-D}, SK_{P-AGW})\}_{SK_{D-AGW}}$

AGW cannot deny sending the original message to D , as only AGW possesses both the symmetric keys SK_{D-AGW} and SK_{P-D} . Hence, only AGW can construct this message or the original message.

4.1.2. Integrity: This property ensures at the recipient's end that the information in the received message has not been altered by an attacker during the exchange of messages. The proposed protocol utilizes a cryptographic hash function and a MAC to guarantee message integrity.

4.1.3. Confidentiality: Personal health information should not made be available or disclosed to unauthorized persons, and should be protected from disclosure to an attacker. Health information should be confidential, and made available only to authorized doctors. The proposed protocol applies a symmetric key to encrypt the messages that are exchanged between parties. This can ensure that the protocol provides message confidentiality.

4.1.4. Replay attack: In this scenario, an attacker records old messages and then resends them, as these are valid message transmissions. The attacker therefore gets the same messages as the legitimate parties. The proposed protocol uses a nonce and a timestamp at each step of the protocol, which can prevent replay attacks.

4.1.5. Eavesdropping attack: In this case, an attacker secretly listens to a conversation transmitted over the air between parties, to obtain medical information about the victim. The goal of the attacker is to learn the content of the exchanged message. An attacker that eavesdrops on medical information can collect a large amount of information. To prevent this, the proposed protocol uses symmetric key encryption for the message exchange.

4.1.6. Data modification: An attacker could edit a message and send it on to the receiver during the communication process, which could result in a false diagnosis. Data modification cannot occur in our scheme, since we use symmetric cryptography, including a hash function, in each step.

4.1.7. MITM attack: An attacker cannot analyze a transmitted message or fraudulently pose as one of the parties (i.e., the sender or receiver), since the proposed protocol uses a cryptographic hash function and symmetric key cryptography to maintain the confidentiality of messages and the message integrity. Furthermore, our protocol applies a MAC to identify the sender and the receiver, who share the same symmetric keys.

From Table 2, it can be seen that scheme in [19] and our approach provide all of the security properties, whereas the protocols in [17, 18, 20] do not ensure message integrity.

Table 2. Security comparison of the proposed protocol and existing alternatives.

	[17]	[18]	[19]	[20]	P
Mutual authentication	N	Y	Y	Y	Y
Integrity	N	N	Y	N	Y
Confidentiality	Y	Y	Y	Y	Y
Replay attack	Y	Y	Y	Y	Y

	[17]	[18]	[19]	[20]	P
Eavesdropping attack	Y	Y	Y	Y	Y
Data Modification	Y	Y	Y	Y	Y
Man-in-the-middle attack	Y	Y	Y	Y	Y

P: Our protocol

4.2. Formal Security Analysis

4.2.1. Using Scyther

We used the Scyther tool to verify that our proposed protocol was safe and robust against attacks. Security Protocol Description Language (SPDL) code for the proposed protocol is shown in Figures 2, 3 and 4, we present the results of verification, claims and automatic claims of the proposed protocol that has no attack. More information about Scyther can be found in [14, 15].

```

1./* Mutual Authentication Protocol */
2.hashfunction h;
3.usertype Timestamp;
4.const DiD, PiD;
5.const SKP-AGW, SKD-AGW, SKP-D
   :SessionKey;
6.macro m1 = DiD, PiD, nd, td, h(DiD, PiD,
   nd, td, SKD-AGW);
7.macro m2 = DiD, PiD, nd, td, h(DiD, PiD,
   nd, td, SKD-AGW), h(DiD, PiD, nd, td,
   h(DiD, PiD, nd, td, SKD-AGW, SKP-
   AGW));
8.macro m3 = np, h(np, h(DiD, PiD, nd, td,
   h(DiD, PiD, nd, td, SKD-AGW), SKP-
   AGW, SKP-AGW));
9.macro m4 = nagw, tagw, {SKP-D, h(DiD,
   PiD, nd, np, nagw, td, tagw, SKP-D, SKP-
   AGW)}SKP-AGW;
10. macro m5 = nagw, tagw, {SKP-D, h(DiD,
   PiD, nd, np, nagw, td, tagw, SKP-D,
   SKD-AGW)}SKD-AGW;
11. // The protocol description
12. protocol M-Auth(D, AGW, P)
13. {
14. role D

```

```

15. {
16. fresh td, tagw: Timestamp;
17. fresh nr, ni, nd, np, nagw: Nonce;
18. send_1(D, AGW, m1);
19. rcv_5(AGW, D, m5);
20. claim_d1(D, Secret, nd);
21. claim_d2(D, Secret, np);
22. claim_d3(D, Alive);
23. claim_d4(D, Weakagree);
24. claim_d5(D, Commit, AGW, nd,np);
25. claim_d6(D, Niagree);
26. claim_d7(D, Nisynch);
27. }
28. role AGW
29. {
30. fresh nr, ni, nd, np, nagw: Nonce;
31. fresh td, tagw: Timestamp;
32. rcv_1(D, AGW, m1);
33. send_2(AGW, P, m2);
34. rcv_3(P, AGW, m3);
35. send_4(AGW, P, m4);
36. send_5(AGW, D, m5);
37. claim_agw1(AGW, Secret, nagw);
38. claim_agw2(AGW, Secret, tagw);
39. claim_agw3(AGW, Alive);
40. claim_agw4(AGW, Weakagree);
41. claim_agw5(AGW, Commit, P, nagw);
42. claim_agw6(AGW, Niagree);
43. claim_agw7(AGW, Nisynch);
44. claim_agw8(AGW, Commit, D, nagw);
45. }
46. role P
47. {
48. fresh nr, ni, nd, np, nagw : Nonce;
49. fresh td, tagw: Timestamp;
50. rcv_2(AGW, P, m2);
51. send_3(P, AGW, m3);
52. rcv_4(AGW, P, m4);
53. claim_P1(P, Secret, nagw);
54. claim_P2(P, Secret, np);
55. claim_P3(P, Alive);
56. claim_P4(P, Weakagree);
57. claim_P5(P, Commit, D, np,nagw);
58. claim_P6(P, Niagree);
59. claim_P7(P, Nisynch);
60. }
61. /* End of Program */
62. }

```

Figure 2. SPDL code for the proposed protocol.

Scyther results : verify

Claim	Status	Comments
M_Auth D	Ok	Verified No attacks.
M_AuthD1	Ok	Verified No attacks.
M_AuthD2	Ok	Verified No attacks.
M_AuthD3	Ok	Verified No attacks.
M_AuthD4	Ok	Verified No attacks.
M_AuthD5	Ok	Verified No attacks.
M_AuthD6	Ok	Verified No attacks.
M_AuthD7	Ok	Verified No attacks.
AGW M_Authagn1	Ok	Verified No attacks.
M_Authagn2	Ok	Verified No attacks.
M_Authagn3	Ok	Verified No attacks.
M_Authagn4	Ok	Verified No attacks.
M_Authagn5	Ok	Verified No attacks.
M_Authagn6	Ok	Verified No attacks.
M_Authagn7	Ok	Verified No attacks.
M_Authagn8	Ok	Verified No attacks.
P M_AuthP1	Ok	Verified No attacks.
M_AuthP2	Ok	Verified No attacks.
M_AuthP3	Ok	Verified No attacks.
M_AuthP4	Ok	Verified No attacks.
M_AuthP5	Ok	Verified No attacks.
M_AuthP6	Ok	Verified No attacks.
M_AuthP7	Ok	Verified No attacks.

Done.

Figure 3. Verification results from the Scyther tool.

Scyther results : autoverify

Claim	Status	Comments
M_Auth D	Ok	Verified No attacks.
M_AuthD1	Ok	Verified No attacks.
M_AuthD2	Ok	Verified No attacks.
M_AuthD3	Ok	Verified No attacks.
M_AuthD4	Ok	Verified No attacks.
M_AuthD5	Ok	Verified No attacks.
M_AuthD6	Ok	Verified No attacks.
M_AuthD7	Ok	Verified No attacks.
M_AuthD8	Ok	Verified No attacks.
M_AuthD9	Ok	Verified No attacks.
M_AuthD10	Ok	Verified No attacks.
M_AuthD11	Ok	Verified No attacks.
M_AuthD12	Ok	Verified No attacks.
AGW M_AuthAGW1	Ok	Verified No attacks.
M_AuthAGW2	Ok	Verified No attacks.
M_AuthAGW3	Ok	Verified No attacks.
M_AuthAGW4	Ok	Verified No attacks.
M_AuthAGW5	Ok	Verified No attacks.
M_AuthAGW6	Ok	Verified No attacks.
M_AuthAGW7	Ok	Verified No attacks.
M_AuthAGW8	Ok	Verified No attacks.
M_AuthAGW9	Ok	Verified No attacks.
M_AuthAGW10	Ok	Verified No attacks.
M_AuthAGW11	Ok	Verified No attacks.
M_AuthAGW12	Ok	Verified No attacks.
P M_AuthP1	Ok	Verified No attacks.
M_AuthP2	Ok	Verified No attacks.
M_AuthP3	Ok	Verified No attacks.
M_AuthP4	Ok	Verified No attacks.
M_AuthP5	Ok	Verified No attacks.
M_AuthP6	Ok	Verified No attacks.
M_AuthP7	Ok	Verified No attacks.
M_AuthP8	Ok	Verified No attacks.
M_AuthP9	Ok	Verified No attacks.
M_AuthP10	Ok	Verified No attacks.

Done.

Figure 4. Autoverification using the Scyther tool.

4.2.2. Using AVISPA

We also used AVISPA to prove that our proposed protocol is safe and is robust against attacks. There are numerous research papers that have applied this approach to prove their protocols [21, 22]. AVISPA uses the High-Level Protocol Specification Language (HLPSL) specification syntax to generate a graphical on-the-fly model checker (OMFC), constraint-logic-based model checker (ATSE), and attack trace generation, to determine whether or not the authentication protocol is vulnerable to attack. The results from OMFC, ATSE and attack generation are shown in Figures 5 and 6.



Figure 5. OMFC results verified using AVISPA.



Figure 6. ATSE results verified using AVISPA.

5. Performance Analysis

Tables 3 to 5 and Figures 7 to 9 show the results for cryptographic operations cost, energy cost and time cost, respectively, and provide a comparison with the protocols in [17-20]. It can be seen that our protocol imposes lower cryptographic operations, energy and time costs than the protocols in [17-20]. Note that the process used to measure the energy consumption is derived from [23], and the method used to measure the time consumption is derived from [24]. Note that symmetric encryption uses the Advanced Encryption Standard (AES) while a one-way hash function uses Secure Hashing Algorithm (SHA1). P stands for the proposed protocol.

Table 3. Comparison of cryptographic operation cost.

	AES	SHA1	Total
[17]	7	0	7
[18]	6	0	6
[19]	5	2	7
[20]	6	0	6
P	2	5	7

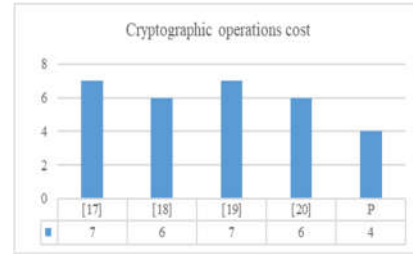


Figure 7. Comparison of cryptographic operation cost.

Table 4. Comparison of energy consumption.

	AES (1.21 μ J/byte)	SHA1 (0.76 μ J/byte)	Total
[17]	8.47	0	8.47
[18]	7.26	0	7.26
[19]	6.05	1.52	7.57
[20]	7.26	0	7.26
P	2.42	3.8	6.22

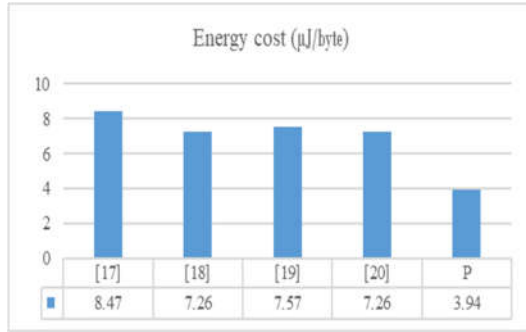


Figure 8. Comparison of energy consumption.

Table 5. Comparison of time consumption.

	AES (1.71 ms/byte)	SHA1 (1.28 ms/byte)	Total
[17]	11.97	0	11.97
[18]	10.26	0	10.26
[19]	8.55	2.56	11.11
[20]	10.26	0	10.26
P	3.42	6.4	9.82

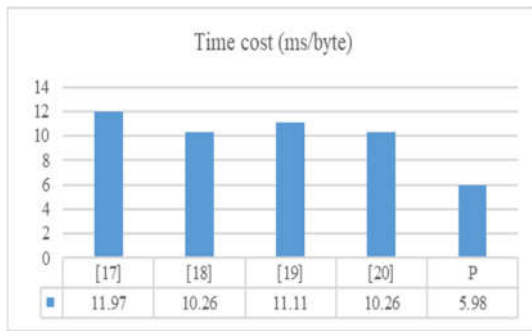


Figure 9. Comparison of time consumption.

6. Conclusion and Discussion

Protocol for WBAN devices needs to provide all essential security properties, to protect against misuse of patient health information and prevent threats from attackers. We have analyzed the security of our proposed protocol and compared it with other existing alternatives. The results of our analysis show that the proposed protocol provides security properties such as mutual authentication, integrity, confidentiality, and protection against replay, eavesdropping, data

modification and MITM attacks. Our cryptographic algorithms use only symmetric encryption and a hash function, which enhances security while creating a lightweight protocol that is more effective than other protocols. The results from the Scyther tool show that our proposed protocol ensures the most important security properties needed for a WBAN and is robust against attackers.

In future work, we will focus on developing a prototype based on the proposed protocol, in order to show that our approach is practical for real- world applications.

References

- [1] C. Techapanupreed, W. Kurutach, "Enhancing transaction security for handling accountability in electronic health records," Security and Communication Networks, 2020.
- [2] G. hamilarasu, and A. Odesile, "Securing wireless body area networks: Challenges, review and recommendations," International Conference on Computational Intelligence and Computing Research (ICCIC), pp. 1-7, 2016.
- [3] M. Kompara, and M. Hölbl, "Survey on security in intra-body area network communication," Ad Hoc Networks, vol. 70, pp. 23-43, 2018.
- [4] D. Vera, N. Costa, L. Roda-Sanchez, T. Olivares, A. Fernández-Caballero, and A. Pereira, "Body area networks in healthcare: A brief state of the art," Applied Sciences, vol. 9, no. 16, pp. 3248, 2019.
- [5] F. R. Yazdi, M. Hosseinzadeh, and S. Jabbehdari, "A review of state-of-the-art on wireless body area networks," International Journal of Advanced Computer Science and Applications, pp. 443-455, 2017.
- [6] R. A. Khan, and A. S. K. Pathan, "The state-of-the-art wireless body area sensor networks: A survey," International Journal of Distributed Sensor Networks, vol. 14, no. 4, 2018.
- [7] C. A. Tavera, J. H. Ortiz, O. I. Khalaf, D. F. Saavedra, and T. H. Aldhyani, "Wearable wireless body area networks for medical applications," Computational and Mathematical Methods in Medicine, 2021.
- [8] S. J. Hussain, M. Irfan, N. Z. Jhanjhi, K. Hussain, and M. Humayun, "Performance enhancement in wireless body area networks with secure communication," Wireless Personal Communications, vol. 116, no. 1, pp. 1-22, 2021.

- [9] C. C. Tan, H. Wang, S. Zhong, and Q. Li, "Body sensor network security: an identity-based cryptography approach," In Proceedings of the first ACM conference on Wireless network security, pp. 148-153, 2008.
- [10] K. H. Yeh, "A secure IoT-based healthcare system with body sensor networks," IEEE Access, vol. 4, pp. 10288-10299, 2016.
- [11] IEEE 802.15 WPAN Task Group 6 (TG6) Body Area Networks". IEEE Standards Association. 9 Jun 2011. Retrieved 9 Dec 2021.
- [12] IEEE Standard for Information technology-- Local and metropolitan area networks- "Specific requirements-- Part 15.1a: Wireless Medium Access Control (MAC) and Physical Layer (PHY) specifications for Wireless Personal Area Networks (WPAN)," in IEEE Std 802.15.1-2005 (Revision of IEEE Std 802.15.1-2002) , vol., no., pp.1-700, 14 June 2005, doi: 10.1109/IEEESTD.2005.96290.
- [13] Approved IEEE Draft Amendment to IEEE Standard for Information Technology-Telecommunications and Information Exchange Between Systems-Part 15.4: "Wireless Medium Access Control (MAC) and Physical Layer (PHY) Specifications for Low-Rate Wireless Personal Area Networks (LR-WPANS): Amendment to Add Alternate Phy (Amendment of IEEE Std 802.15.4)," in IEEE Approved Std P802.15.4a/D7, Jan 2007 , 2007.
- [14] C. J. Cremers, "The Scyther tool: Verification, falsification, and analysis of security protocols," In International Conference on Computer Aided Verification, Springer, Berlin, Heidelberg, pp. 414-418, 2008.
- [15] C. Thammarat, and W. Kurutach, "A lightweight and secure NFC-base mobile payment protocol ensuring fair exchange based on a hybrid encryption algorithm with formal verification," International Journal of Communication Systems, vol. 32, no. 12, 2019.
- [16] W. Stallings, L. Brown, , M. D. Bauer, and A. K. Bhattacharjee, "Computer security: principles and practice," Upper Saddle River, NJ, USA: Pearson Education, pp. 978, 2012.
- [17] R. Yan, J. Liu, and R. Sun, "An efficient authenticated key exchange protocol for wireless body area network," The Proceedings of the Third International Conference on Communications, Signal Processing, and Systems, Springer, Cham, pp. 51-58, 2015.
- [18] J. Liu, Q. Li, R. Yan, and R. Sun, "Efficient authenticated key exchange protocols for wireless body area networks," EURASIP Journal on Wireless Communications and Networking, pp. 1-11, 2015.
- [19] D. He, N. Kumar, J. Chen, C. C. Lee, N. Chilamkurti, and S. S. Yeo, "Robust anonymous authentication protocol for health-care applications using wireless medical sensor networks," Multimedia Systems, vol. 21, no. 1, pp. 49-60, 2015.
- [20] R. Vishwakarma, and R. K. Mohapatra, "A secure three-party authentication protocol for wireless body area networks," In 2017 Third International Conference on Sensing, Signal Processing and Security (ICSSS), pp. 99-103, 2017.
- [21] C. Thammarat, and C. Techapanupreeda, "A secure authentication and key exchange protocol for M2M communication," In 2021 9th International Electrical Engineering Congress (iEECON), pp. 456-459, IEEE, 2021.
- [22] C. Thammarat, and C. Techapanupreeda, "A secure mobile payment protocol for handling accountability with formal verification," In 2021 International Conference on Information Networking (ICOIN), pp. 249-254, IEEE, 2021.
- [23] N. R. Potlapally, S. Ravi, A. Raghunathan, and N. K. Jha, "A study of the energy consumption characteristics of cryptographic algorithms and security protocols," in IEEE Trans. Mobile computing, vol. 5, no. 2, pp. 128-143, 2006.
- [24] X. Zheng, L. Yang, J. Ma, G. Shi, and D. Meng, "TrustPAY: Trusted mobile payment on security enhanced ARM TrustZone platforms," in Proc. on Computers and Communication, pp. 456-462, 2016.



การแบ่งส่วนรูปภาพดอกไม้ด้วยการใช้ซาลิเอนซีแมปร่วมกับการประยุกต์ใช้ปริภูมิสี

เอชเอสวีและหน้ากากสี

Flower Image Segmentation using Saliency Map with the Application of HSV Color Space and Color Mask

ธนัญญ์ หงษ์ทอง* และ สุรณพีร์ ภูมิวิฬาสาร

Thananath Hongthong* and Suronapee Phoomvuthisarn

ภาควิชาสถิติ คณะพาณิชยศาสตร์และการบัญชี จุฬาลงกรณ์มหาวิทยาลัย

Department of Statistics, Faculty of Commerce and Accountancy, Chulalongkorn University

Received: November 13, 2021; Revised: December 09, 2021; Accepted: December 23, 2021; Published: December 28, 2021

ABSTRACT – The classification of flowers is a challenging task, due to the similarity of flowers' physical characteristics. Image segmentation techniques can simplify such details within the image background, making it possible to classify flowers efficiently. In this paper, we purpose a technique for image segmentation based on saliency map to select interested region within the image. The use of saliency map combining with the HSV color space with color mask can reduce insignificant details within the image background. Experimental results have shown that our method can select interested region and can reduce the background detail considerably. Our method can achieve 54% mean IoU (up to 13% higher than previous works), while achieving accuracy, precision, recall and F1 values at 87% when it integrates efficiency with the VGG-16 pre-trained model.

KEYWORDS: Saliency map, HSV, Color mask

บทคัดย่อ -- งานวิจัยนี้นำเสนอระบบลงทะเบียนใบหน้าและตรวจสอบนักศึกษาเข้าห้องเรียนด้วยหลักการประมวลผลภาพร่วมกับไลบรารีการเรียนรู้จำใบหน้า ในงานวิจัยนี้ได้ใช้กระบวนการทดสอบประสิทธิภาพของระบบรู้จำใบหน้าของนักศึกษาโดยการเปรียบเทียบภาพใบหน้าของนักศึกษากับภาพต้นฉบับด้วยการลงทะเบียนในระบบ จากนั้นเข้าสู่กระบวนการเตรียมภาพสำหรับการทดลอง เพื่อสร้างฐานข้อมูลภาพใบหน้าของนักศึกษา และนำมาเปรียบเทียบกับภาพใบหน้าจากฐานข้อมูล โดยทดลองใช้อัตราการสุ่มของจำนวนภาพใบหน้าที่แตกต่างกัน จากผลการทดลอง พบว่า การสุ่มบันทึกภาพใบหน้าของนักศึกษา จำนวน 100 ภาพต่อนักศึกษา 1 คน ได้ผลการทดลองที่แม่นยำที่สุด โดยระบบรู้จำใบหน้าของนักศึกษาสามารถระบุตัวตนของนักศึกษาได้ถูกต้องแม่นยำอยู่ที่ 92% และระบบนี้สามารถดูข้อมูลย้อนหลังของนักศึกษาในการเข้าห้องเรียน และสามารถนำเอาข้อมูลออกมาใช้งานในรูปแบบของไฟล์เอกสาร

คำสำคัญ: ซาลิเอนซีแมป, ปริภูมิสีเอชเอสวี, หน้ากากสี

*Corresponding Author: Thananath.h@gmail.com

1. บทนำ

การจำแนกรูปภาพดอกไม้เป็นสิ่งที่ท้าทาย เนื่องจากดอกไม้บางชนิดมีความคล้ายคลึงกันมาก อีกทั้งดอกไม้บางชนิดแม้จะมีสายพันธุ์ต่างกันแต่ก็มีลักษณะทางกายภาพที่คล้ายคลึงกันทั้งสี รูปร่าง และลักษณะภายนอก มากกว่านั้นภายในรูปภาพมักจะมียอดไม้ประกอบอื่นรวมอยู่ด้วย เช่น ใบไม้ หญ้า หิ่งฟ้าย ฯลฯ ซึ่งโดยทั่วไปการจำแนกพืชมักจะทำโดยผู้เชี่ยวชาญด้านพฤกษศาสตร์ โดยจะทำการวัดคุณลักษณะของพืชด้วยตนเอง แต่ข้อเสียของวิธีการนี้คือค่อนข้างใช้เวลานาน ยิ่งไปกว่านั้นถ้าหากรูปภาพที่ใช้มีจำนวนมาก อาจเกิดข้อผิดพลาดได้เนื่องจากพื้นหลังของรูปภาพมีความซับซ้อน เพื่อลดความซับซ้อนของภาพ เทคนิคการแบ่งส่วนรูปภาพสามารถแบ่งรูปภาพออกเป็นพื้นที่ที่มีรายละเอียดต่างกันได้อย่างชัดเจน ทำให้สามารถเลือกพื้นที่ที่สนใจภายในภาพสำหรับนำไปจำแนกประเภทดอกไม้ได้มีประสิทธิภาพมากขึ้น

โดยทั่วไปแนวคิดเกี่ยวกับการแบ่งส่วนรูปภาพจะประกอบไปด้วยวิธีพื้นฐาน 3 แบบ ได้แก่ (1) การปรับค่า Threshold (2) การสร้างโครงสร้าง (Structure) ของดอกไม้เพื่อเก็บบริเวณที่สนใจ และ (3) การเปลี่ยนปริภูมิสี (Color space) ของรูปภาพ ซึ่งการเลือกใช้เพียงวิธีการใดวิธีการหนึ่งในการแยกพื้นหน้า (Foreground) และพื้นหลัง (Background) ออกจากกันจะทำให้แบ่งส่วนรูปภาพได้ไม่แม่นยำ ถ้าหากพื้นหลังของภาพมีความซับซ้อนมาก จะทำให้แยกบริเวณที่สนใจออกจากพื้นหลังรูปภาพได้ยากมากยิ่งขึ้น ด้วยเหตุนี้จึงมีงานวิจัยที่นำเอาข้อดีของวิธีการทั้ง 3 วิธีมาผนวกกัน ด้วยการนำข้อดีของการใช้การระบุตำแหน่งวัตถุมาใช้ระบุตำแหน่งสำคัญของวัตถุภายในภาพ จากนั้นทำการลบภาพพื้นหลังที่ไม่เกี่ยวข้องด้วยการปรับความเข้มของสีภายในภาพ [1] หรือการเปลี่ยนปริภูมิสีของภาพ [2] ซึ่งงานวิจัยเหล่านี้แสดงให้เห็นว่าการที่จะระบุตำแหน่งของวัตถุให้แม่นยำจำเป็นต้องมีโครงสร้างของดอกไม้ที่ชัดเจน จึงมีงานวิจัยที่เสนอวิธีการหาโครงสร้างของดอกไม้เพื่อหาบริเวณที่สนใจในรูปภาพ ซึ่งมีทั้งการใช้ K-means clustering [3] และการใช้ Saliency map โดยเฉพาะการใช้ Saliency map เป็นสิ่งที่น่าสนใจเป็นอย่างมาก เนื่องจาก Saliency map

สามารถแสดงรายละเอียดของบริเวณที่สนใจภายในรูปภาพได้ค่อนข้างชัดเจน อีกทั้งยังช่วยลดรายละเอียดที่ไม่สำคัญภายในภาพลงได้ [4] แต่ถึงอย่างไรการใช้เพียง Saliency map เพียงอย่างเดียวยังไม่สามารถแบ่งรูปภาพได้ดีเท่าที่ควร เนื่องจากภายในรูปภาพที่ถูกครอบตัดจะยังคงเหลือบริเวณที่ไม่เกี่ยวข้องหลงเหลืออยู่ค่อนข้างมาก หากต้องการครอบตัดรูปภาพให้เหลือเพียงบริเวณที่สนใจเพื่อนำภาพไปใช้ในการจำแนกประเภทของดอกไม้จำเป็นต้องลดรายละเอียดพื้นหลังของภาพให้น้อยลง

เพื่อที่จะแก้ไขปัญหานี้ ในงานวิจัยชิ้นนี้จึงได้นำเสนอแนวคิดการแบ่งส่วนรูปภาพ โดยอิงการใช้ประโยชน์จาก Saliency map ในการเก็บรายละเอียดของบริเวณที่สนใจภายในภาพ และการใช้ประโยชน์จากปริภูมิสี HSV ผสมกับการประยุกต์ใช้ Color mask ในการช่วยลดรายละเอียดที่ไม่สำคัญภายในพื้นหลังของรูปภาพออกไป โดยงานวิจัยชิ้นนี้จะแบ่งออกเป็น 2 ขั้นตอน ขั้นตอนแรกจะเป็นขั้นตอนการแบ่งส่วนรูปภาพของดอกไม้ และในขั้นตอนที่สองจะเป็นการนำรูปภาพที่ได้มาจากขั้นตอนแรกมาใช้จำแนกประเภทของดอกไม้

จากการทดลองพบว่าวิธีที่นำเสนอสามารถเก็บบริเวณที่สนใจและลดรายละเอียดของพื้นหลังได้ค่อนข้างมาก ทำให้สามารถเลือกบริเวณที่สำคัญภายในภาพได้ดี โดยผลลัพธ์การแบ่งส่วนรูปภาพจากการวัดค่าเฉลี่ย IoU เท่ากับ 54% ซึ่งมากกว่าวิธีการที่นำเสนอโดย Yanyang ที่ให้ค่าเฉลี่ย IoU เท่ากับ 41% นอกจากนี้เมื่อนำชุดข้อมูลที่ได้นำไปจำแนกประเภทดอกไม้ด้วยแบบจำลอง VGG16 ที่ผ่านการปรับโครงสร้างพบว่าให้ผลลัพธ์ค่าความถูกต้อง ความแม่นยำ ความครบถ้วน และ F1 มากที่สุดคือ 87%

2. ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

2.1 ทฤษฎีที่เกี่ยวข้อง

Saliency map คือ รูปภาพที่ความสว่างของพิกเซลจะแสดงความเด่นของพิกเซล ซึ่งความสว่างของพิกเซลจะส่งผลโดยตรงต่อความเด่นของรูปภาพ โดย Saliency map จะได้มาจากวิธีการจำแนกความต่างในรูปภาพที่ชื่อว่า RC (Regional contrast) ซึ่งเป็นการผนวกทั้งการเปรียบเทียบความต่างทั้งหมดภายในรูปภาพ (Global contrast) และความสอดคล้องเชิงพื้นที่ (Spatial coherence) โดย Histogram based contrast หรือ HC จะเป็นตัวกำหนดค่าความโดดเด่นของพิกเซลโดยพิจารณาจากความต่างของสีของแต่ละพิกเซลในภาพ [5] ในที่นี้ค่าความเด่นของพิกเซลในรูปภาพ (Saliency value) สามารถคำนวณได้จากสมการที่ (1) และ (2) เนื่องจากพิกเซลมีสีใกล้เคียงกันจะมีค่านัยสำคัญเหมือนกัน จึงมีการจัดรูปสมการใหม่โดยการจัดกลุ่มให้สีที่คล้ายกันมาอยู่ด้วยกัน ด้วยวิธีนี้จะทำให้มองเห็นค่าความเด่นของแต่ละสีได้ดียิ่งขึ้น ดังแสดงในสมการที่ (3)

$$S(I_k) = \sum_{I_i \in I} D(I_k, I_i) \quad (1)$$

$$S(I_k) = D(I_k, I_1) + D(I_k, I_2) + \dots + D(I_k, I_N) \quad (2)$$

$$S(I_k) = S(c_l) = \sum_{j=1}^n f_j D(c_l, c_j) \quad (3)$$

โดยที่ I แทน รูปภาพ
 I_k แทน ค่าความเด่นของพิกเซลภายในภาพ
 D แทน ระยะสีระหว่างพิกเซลภายในรูปภาพ (Color distance matrix)
 c_l แทน สีของพิกเซลที่ I_k
 n แทน จำนวนสีของพิกเซลที่แตกต่างกัน
 N แทน จำนวนพิกเซลในรูปภาพ I
 f_j แทน ความน่าจะเป็นที่ c_j ปรากฏอยู่ในรูปภาพ
 $S(I_k)$ แทน ค่าความเด่นของพิกเซลในรูปภาพ

ถึงแม้ว่าจะมีการจัดกลุ่มสีที่คล้ายกันแล้ว แต่ยังมีโอกาสที่บางเฉดสียังไม่ถูกจัดกลุ่ม จึงทำให้ยังมีค่าความเด่นสุมกระจายภายในภาพ ด้วยเหตุนี้จึงมีการนิยามค่าความเด่นของสีด้วยค่าเฉลี่ยน้ำหนักของค่าความเด่นของสีที่คล้ายกัน เรียกวิธีการนี้ว่าการปรับสีพื้นที่ให้เรียบ (Color space smoothing) จากนั้น

ให้ค่าของน้ำหนักที่มากกว่าเป็นตัวแทนของค่าเฉลี่ยในแต่ละกลุ่ม ดังแสดงในสมการที่ (4)

$$S'(c) = \frac{1}{(m-1)T} \sum_{i=1}^m (T - D(c, c_i)) S(c_i) \quad (4)$$

โดยที่ $S'(c)$ แทน ค่าความเด่นของพิกเซลในรูปภาพที่มีการปรับสีพื้นที่ให้เรียบ

T แทน $\sum_{i=1}^m D(c, c_i)$

m แทน $\frac{n}{4}$

ในขั้นตอนนี้จะเป็นการอธิบายวิธีการคำนวณหาค่าความเด่นในแต่ละพื้นที่ สำหรับพื้นที่ r_k ค่าความเด่นสามารถคำนวณได้จากการวัดความต่างของสีจากพื้นที่ r_k ไปยังพื้นที่ใด ๆ ในรูปภาพ ดังแสดงในสมการที่ (5) และ (6)

$$S(r_k) = \sum_{r_i \neq r_k} w(r_i) D_r(r_k, r_i) \quad (5)$$

$$D_r(r_1, r_2) = \sum_{i=1}^{n_1} \sum_{j=1}^{n_2} f(c_{1,i}) f(c_{2,j}) D(c_{1,i}, c_{2,j}) \quad (6)$$

โดยที่ r_k แทน พื้นที่
 $w(r_i)$ แทน น้ำหนักของพื้นที่ r_i
 $f(c_{k,j})$ แทน ความน่าจะเป็นของสีที่ $c_{k,j}$ ทั่วกลางสีทั้งหมด n_k ในพื้นที่ r_k
 $S(r_k)$ แทน ค่าความเด่นของพื้นที่ในรูปภาพ
 $D_r(r_1, r_2)$ แทน ระยะห่างของเมตริกซ์สีระหว่างพื้นที่ r_1 และ พื้นที่ r_2

เพื่อที่จะเพิ่มผลกระทบเชิงพื้นที่ จึงได้มีการเพิ่มพจน์การถ่วงน้ำหนักเชิงพื้นที่ เพื่อใช้ในการคำนวณค่าความเด่นเชิงพื้นที่ โดยพจน์น้ำหนักจะคำนวณได้จากค่าเฉลี่ยระยะทางระหว่างพิกเซลในพื้นที่ r_k กับจุดศูนย์กลางของรูปภาพ ดังนั้นพื้นที่ที่มีระยะทางใกล้กับจุดศูนย์กลางของรูปภาพจะให้ค่าความเด่นที่สูง ดังแสดงในสมการที่ (7)

$$S(r_k) = w_s(r_k) \sum_{r_i \neq r_k} e^{\frac{D_s(r_k, r_i)}{-\sigma_s^2}} w(r_i) D_r(r_k, r_i) \quad (7)$$

โดยที่ $S(r_k)$ แทน ค่าความเด่นของพื้นที่ในรูปภาพ

$w_s(r_k)$ แทน พจน์การถ่วงน้ำหนักเชิงพื้นที่
 $D_s(r_k, r_i)$ แทน ระยะทางเชิงพื้นที่ระหว่างพื้นที่ r_k
 และ พื้นที่ r_i
 σ_s แทน พารามิเตอร์ (Parameter) ควบคุม
 จุดเด่นของระยะทางเชิงพื้นที่ถ่วงน้ำหนัก
 $w(r_i)$ แทน น้ำหนักของพื้นที่ r_i

2.2 งานวิจัยที่เกี่ยวข้อง

งานวิจัยที่ศึกษาเกี่ยวกับการแบ่งส่วนรูปภาพโดยทั่วไปจะใช้วิธีพื้นฐาน 3 แบบ ได้แก่ การปรับค่า Threshold การสร้างโครงร่างของคอกไม้เพื่อเก็บบริเวณที่สนใจ และการเปลี่ยนปริภูมิสี แต่งานวิจัยส่วนใหญ่มักจะใช้วิธีเหล่านี้ร่วมกันเพื่อเก็บตำแหน่งของวัตถุที่สนใจและลดรายละเอียดพื้นหลังของภาพ

การแบ่งส่วนรูปภาพด้วยวิธีการปรับค่า Threshold เป็นวิธีที่ได้รับความนิยมเป็นอย่างมาก [6] [7] [8] เนื่องจากเป็นวิธีที่ไม่ซับซ้อนและใช้เวลาน้อย แต่วิธีนี้ไม่เหมาะกับการแบ่งส่วนชุดข้อมูลรูปภาพขนาดใหญ่ เพราะการปรับค่า Threshold ให้มีค่าพอคักับชุดข้อมูลรูปภาพเป็นเรื่องที่ทำหยาบมาก จึงมีโอกาสเกิดข้อผิดพลาดได้ ด้วยเหตุนี้งานวิจัยที่ใช้วิธีนี้ [9] จึงนิยมใช้วิธีการนี้ร่วมกับการเปลี่ยนปริภูมิสี เช่น Lab เพื่อลดข้อผิดพลาดที่จะเกิดขึ้น

นอกจากการปรับค่า Threshold การแบ่งส่วนรูปภาพโดยการใช้ K-mean clustering นั้นเป็นอีกหนึ่งวิธีที่ได้รับความนิยมเช่นกัน [10] [11] [12] เนื่องจากมีอัลกอริทึมที่เข้าใจง่ายและทำงานได้เร็ว แต่การใช้วิธีนี้เพียงอย่างเดียวไม่สามารถลดรายละเอียดของพื้นหลังรูปภาพได้ งานวิจัยที่ใช้วิธีนี้ [13] [14] จึงใช้วิธีการอื่นควบคู่ไปด้วย อย่างการเปลี่ยนปริภูมิสีของรูปภาพเป็น HSV หรือ ปริภูมิสี Lab (CIELab)

เนื่องจากวิธีการที่มีกรนำเสนอในงานวิจัยส่วนมาก [15] [16] มักจะพยายามหาบริเวณของวัตถุที่สนใจภายในรูปภาพจึงได้มีงานวิจัยที่นำเสนอโดย Yangyang [17] ซึ่งทำการแบ่งส่วนรูปภาพด้วยการใช้ Saliency map ในการเก็บบริเวณของวัตถุที่สนใจ ซึ่งวิธีการนี้สามารถแสดงรายละเอียดของบริเวณที่สนใจ

ภายในรูปภาพได้ค่อนข้างชัดเจน อีกทั้งยังช่วยลดรายละเอียดที่ไม่สำคัญภายในภาพลงได้ ถึงแม้การใช้ Saliency map จะสามารถช่วยลดรายละเอียดของพื้นหลังรูปภาพได้แต่ยังไม่เพียงพอเมื่อต้องครอบคลุมบริเวณสำคัญในภาพเพื่อนำไปใช้จำแนกประเภท ด้วยเหตุนี้ในงานวิจัยของ Yangyang จึงใช้วิธีการปรับเบสสีของ Saliency map ให้เป็นเบสสีเทาแล้วจึงปรับพื้นหลังของรูปภาพให้เป็นสีดำซึ่งวิธีการนี้มีประสิทธิภาพเมื่อพื้นหลังรูปภาพมีความซับซ้อนน้อยจนถึงปานกลาง แต่ถ้าหากพื้นหลังรูปภาพมีความซับซ้อนมากขึ้นวิธีการนี้อาจจะไม่เพียงพอ ดังนั้นในงานวิจัยชิ้นนี้จึงได้นำเสนอวิธีการแบ่งส่วนรูปภาพโดยใช้ Saliency map ในการเก็บบริเวณของวัตถุที่สนใจเช่นเดียวกับงานวิจัยของ Yangyang แต่เพิ่มขั้นตอนการลดรายละเอียดของพื้นหลังคอกไม้ ด้วยการใช้ปริภูมิสี HSV และ Color mask เพื่อช่วยลดรายละเอียดของภาพที่มีพื้นหลังมีความซับซ้อนมาก เพื่อให้สามารถครอบคลุมบริเวณที่สนใจได้ดียิ่งขึ้น

3. วิธีการดำเนินการวิจัย

งานวิจัยชิ้นนี้จะแบ่งออกเป็น 2 ขั้นตอน ขั้นตอนแรกจะเป็นขั้นตอนการแบ่งส่วนรูปภาพของคอกไม้ และในขั้นตอนที่สองจะเป็นการนำรูปภาพที่ได้มาจากขั้นตอนแรกมาใช้จำแนกประเภทของคอกไม้

3.1 ขั้นตอนการแบ่งส่วนรูปภาพของคอกไม้

เนื่องจากรูปภาพคอกไม้มักจะมีองค์ประกอบอื่นที่ไม่เกี่ยวข้องอยู่ภายในรูปภาพ การจะนำรูปภาพไปใช้ในการจำแนกประเภทคอกไม้ในทันทีอาจจะทำให้แบบจำลองไม่สามารถเรียนรู้คุณลักษณะ (Feature) สำคัญของคอกไม้ได้ดี ด้วยเหตุนี้จึงจำเป็นที่จะต้องลดรายละเอียดบางส่วนภายในพื้นหลังของภาพเพื่อลดความซับซ้อนก่อนนำไปใช้ในการจำแนกประเภท โดยในงานชิ้นนี้ได้นำเสนอขั้นตอนการสกัดเพื่อหาพื้นที่สำคัญภายในรูปภาพคอกไม้ ซึ่งจะมีรายละเอียดขั้นตอนดังนี้

3.1.1 การสร้าง Saliency map

ในงานชิ้นนี้ Saliency map จะได้มาจากวิธีการจำแนกความต่างในรูปภาพที่ชื่อว่า RC ซึ่งเป็นวิธีการเดียวกับที่ใช้ในงานวิจัยของ Yangyang ดังที่นำเสนอไว้ในสมการที่ (1) ถึง (7) โดย Saliency map จะช่วยแยกวัตถุที่สนใจออกจากพื้นหลังได้และช่วยลดรายละเอียดของพื้นหลังบางส่วนภายในภาพได้ จึงทำให้ง่ายต่อการหาบริเวณที่สนใจภายในภาพ

3.1.2 การแปลงปริภูมิสีของ Saliency map จาก RGB เป็น HSV เนื่องจากความหลากหลายที่เกิดขึ้นภายในรูปภาพ ทั้งความหลากหลายของแสง เงา และจุดรบกวน (Noise) ที่เกิดขึ้นในภาพถ่าย เพื่อลดปัญหาเหล่านี้ จึงทำการเปลี่ยนปริภูมิสีของ Saliency map จาก RGB เป็นปริภูมิสี HSV สมมติ R, G และ B เป็น 3 ช่องของพิกเซล และมีค่าอยู่ในช่วง [0, 255] จากนั้นทำการหารด้วย 255 เพื่อให้มีค่าอยู่ในช่วง [0, 1] [18]

$$\text{จะได้ว่า} \quad R' = R/255 \quad (8)$$

$$G' = G/255 \quad (9)$$

$$B' = B/255 \quad (10)$$

$$\text{โดยที่} \quad Cmax = \max(R', G', B') \quad (11)$$

$$Cmin = \min(R', G, B') \quad (12)$$

$$\text{และ} \quad \Delta = Cmax - Cmin \quad (13)$$

จะได้ ค่าสี (Hue) คือ

$$H = \begin{cases} 60^\circ \times \left(\frac{G' - B'}{\Delta} \bmod 6 \right), & Cmax = R' \\ 60^\circ \times \left(\frac{B' - R'}{\Delta} + 2 \right), & Cmax = G' \\ 60^\circ \times \left(\frac{R' - G'}{\Delta} + 4 \right), & Cmax = B' \end{cases} \quad (14)$$

ค่าความอิ่มของสี (Saturation) คือ

$$S = \begin{cases} 0, & Cmax = 0 \\ \frac{\Delta}{Cmax}, & Cmax \neq 0 \end{cases} \quad (15)$$

และ ค่าความสว่างของแสง (Value)

$$V = Cmax \quad (16)$$

ด้วยปริภูมิสี HSV มีพารามิเตอร์หลายค่าได้แก่ ค่าสี ค่าความอิ่มของสี และ ค่าความสว่างของแสงของแต่ละพิกเซลภายใน

ภาพ จึงทำให้สามารถเห็นเจดสีที่แท้จริงของพิกเซลภายในรูปภาพได้ โดยปราศจากข้อกังวลเกี่ยวกับความเข้มที่มีระดับแตกต่างกัน

3.1.3 การสร้างเค้าโครงดอกไม้

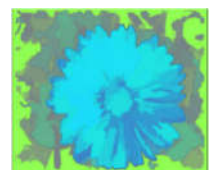
ทำการสร้าง Color mask ซึ่งได้มาจากการดูการกระจายของเจดสีของปริภูมิ RGB ในชุดข้อมูล จากนั้นทำการเลือกเจดสีที่ปรากฏมากที่สุดภายในรูปภาพจำนวน 2 สี ในงานวิจัยชิ้นนี้เลือก Color mask เจดสีน้ำเงินและสีเขียวนำมาซ้อนทับกับ Saliency map ที่ถูกแปลงสีเป็น HSV จะได้ภาพไบนารี (Binary image) เป็นรูปภาพขาวดำ ซึ่งจะแสดงเค้าโครงของดอกไม้

3.1.4 การเลือกพื้นที่ดอกไม้

คำนวณเพื่อหาพื้นที่ที่ติดกันมากที่สุดจากนั้นทำการสร้าง Bounding box เพื่อล้อมรอบบริเวณที่สนใจ จากนั้นจึงครอบตัดบริเวณที่สร้าง Bounding box และทำการปรับขนาดภาพให้เป็น 224 x 224 พิกเซลสำหรับนำไปใช้จำแนกประเภทดอกไม้

โดยตัวอย่างผลลัพธ์ที่ได้จากการประมวลผลภาพจากขั้นตอนที่ 3.1.1 ถึง 3.1.4 ได้แสดงไว้ในตารางที่ 1

ตารางที่ 1. แสดงตัวอย่างผลลัพธ์ของขั้นตอนการเลือกพื้นที่ของดอกไม้

รูปภาพดั้งเดิม	
Saliency map	
Saliency map ที่ผ่านการเปลี่ยนปริภูมิสีจาก RGB เป็น HSV	

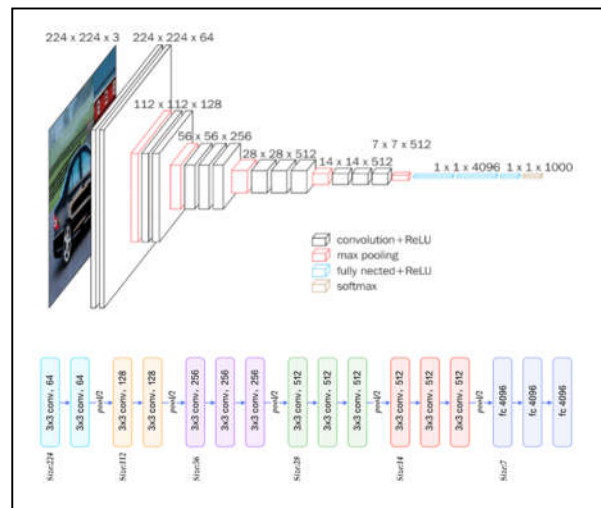
เค้าโครงของดอกไม้	
รูปภาพดอกไม้หลังจากผ่านขั้นตอนการแบ่งรูปภาพ	
พื้นที่ของดอกไม้	
ภาพดอกไม้ที่ถูกครอบตัด	

3.2 ขั้นตอนการจำแนกประเภทของเกทดอกไม้

ในการทดลองใช้ Python เป็นซอฟต์แวร์โปรแกรมและไลบรารีการเขียนโปรแกรมส่วนใหญ่เป็น Keras ที่มีแบ็คเอนด์ TensorFlow และ Scikit-learn โดยการทำงานจะประมวลผลบนหน่วยประมวลผลกราฟฟิค (GPU) แทนการประมวลผลส่วนกลาง (CPU) เนื่องจากการใช้ GPU จะเป็นประโยชน์สำหรับการเรียนรู้เชิงลึก เนื่องจากจะช่วยลดเวลาในการประมวลผล โดย GPU ที่ใช้ในการทดลองสำหรับงานวิจัยชิ้นนี้คือ NVidia K80 ชุดรูปภาพดอกไม้ที่ใช้เป็นชุดข้อมูลดอกไม้จาก Kaggle ประกอบไปด้วยรูปภาพจำนวน 4,931 รูป ชุดข้อมูลประกอบไปด้วยรูปภาพดอกไม้ 5 ประเภท ได้แก่ ดอกเดซี่ ดอกทานตะวัน ดอกทิวลิป ดอกกุหลาบ และดอกแดลลิส

ก่อนการทดลองจะทำการแบ่งข้อมูล 60% เป็นชุดข้อมูลสอน (Training set) 20% เป็นชุดข้อมูลทดสอบ (Testing set) และ 20% เป็นชุดตรวจสอบความถูกต้อง (Validation set) ก่อนนำชุดข้อมูลรูปภาพไปจำแนกประเภทด้วยแบบจำลอง ได้มีการเพิ่มจำนวนชุดข้อมูล (Data augmentation) ด้วยการสุ่มรูปภาพ

ดอกไม้มาประเภทละ 500 รูปภาพ จากนั้นทำการแปลงรูปภาพด้วยการกลับด้านรูปภาพ (Flip) หมุนรูปภาพ (Rotate) ด้วยมุม 90° ปรับค่าความสว่าง (Brightness) ที่ช่วง 0.4 และครอบตัด (Crop) จากศูนย์กลางด้วยอัตราส่วน 0.8 ด้วยวิธีการเหล่านี้ทำให้มีชุดข้อมูลเพิ่มขึ้นจากเดิมเป็น 7,931 รูป



รูปที่ 1. แสดงโครงสร้างแบบจำลอง VGG16 [19]

ในงานวิจัยชิ้นนี้เลือกใช้แบบจำลอง VGGnet ซึ่งมีจำนวน 16 ชั้น (VGG16) ด้วยการถ่ายทอดความรู้จากแบบจำลองที่ถูกสอนมาแล้ว (Transfer learning) ซึ่งแบบจำลอง VGG16 นั้นมีความแม่นยำในการทดสอบใน 5 อันดับสูงสุดอยู่ที่ 92.7% ใน ImageNet ซึ่งเป็นชุดข้อมูลที่มีรูปภาพมากกว่า 14 ล้านภาพที่ใน 1,000 คลาส [20] โดยทั่วไปแบบจำลอง VGG16 จะประกอบไปด้วยชั้นคอนโวลูชันจำนวน 5 บล็อกซึ่งเชื่อมกับชั้นเชื่อมต่อสมบูรณ์ [21] ดังแสดงในรูปที่ 1 ในงานวิจัยชิ้นนี้ได้ทำ Fine-tuning ซึ่งเป็นกระบวนการที่ใช้เมื่อมีการนำแบบจำลอง Pre-trained มาใช้สอนชุดข้อมูลของตัวเอง โดยใช้ชั้นคอนโวลูชันเป็นบล็อกแรกในขณะที่ยังบล็อกอื่นจะถูกบล็อกลบในชั้นเชื่อมต่อสมบูรณ์ (Fully connected) จะประกอบด้วยชั้นประมวลผลที่ซ่อนอยู่ (Hidden layer) จำนวน 2 ชั้น ในชั้นแรกจะประกอบไปด้วยโหนด (Node) จำนวน 256 โหนด ในชั้นที่สองประกอบไปด้วยโหนดจำนวน 64 โหนด โดยชั้นแรกและ

ชั้นที่สองจะถูก Drop out ด้วยสัดส่วน 0.65 % และ 0.3 % ตามลำดับ และใช้ Activation เป็น ReLU ทั้งสองชั้นในชั้น Output ประกอบด้วยโหนดจำนวน 5 โหนดซึ่งเป็นจำนวนประเภทของดอกไม้และใช้ Softmax เป็น Activation แบบจำลองจะถูกสอนด้วย Optimizer Adam ด้วย Learning rate 0.001 จำนวน 40 Epoch และ Batch size สำหรับเป็นชุดข้อมูลสอนจำนวน 128 ตัวอย่าง (Sample) ดังแสดงในตารางที่ 2

ตารางที่ 2. แสดงโครงสร้างแบบจำลอง CNN ที่ใช้ในงานวิจัย

Layer (type)	Function	Output shape
Input		224x224x3
VGG16	Functional	7x7x512
Flatten_1	Flatten	25088x1x1
Dense_1	Fully connected	256x1x1
Activation_3	Activation	256x1x1
Dropout_1	Drop out	256x1x1
Dense_2	Fully connected	64x1x1
Activation_4	Activation	64x1x1
Droupout_2	Drop out	64x1x1
Output	Softmax regression	5x1x1

4. ผลการทดลองและคำอธิบายรายละเอียด

การประเมินผลการทดลองจะถูกแบ่งออกเป็น 2 ขั้นตอน ขั้นตอนแรกจะเป็นการเปรียบเทียบประสิทธิภาพของการแบ่งรูปภาพ และขั้นตอนที่สองจะเป็นการเปรียบเทียบประสิทธิภาพในการจำแนกประเภทของดอกไม้ ซึ่งทั้งสองขั้นตอนจะทำการเปรียบเทียบผลลัพธ์ระหว่างวิธีที่นำเสนอโดย Yangyang และวิธีที่นำเสนอในงานวิจัยชิ้นนี้

4.1 ผลการทดลองและการเปรียบเทียบประสิทธิภาพของการแบ่งรูปภาพ

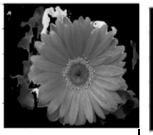
ในขั้นตอนนี้จะทำการหาผลการแบ่งส่วนรูปภาพด้วยวิธีที่นำเสนอ โดย Yangyang และวิธีที่นำเสนอในงานวิจัยชิ้นนี้ โดยจะทำการเปรียบเทียบผลลัพธ์ที่ได้กับภาพที่ถูกต้อง (Ground

truth) ด้วยการคำนวณค่า Bounding box ด้วยวิธี Intersection over Union (IoU) [22] ซึ่งสามารถหาค่าได้จากสมการที่ (16)

$$IoU = \frac{\text{true foreground} \cap \text{segmented foreground}}{\text{true foreground} \cup \text{segmented foreground}} \quad (16)$$

จากสมการสามารถอธิบายได้ว่าค่า IoU สามารถคำนวณได้จากสัดส่วนของพื้นที่ที่ซ้อนทับกันระหว่างพื้นที่ที่ถูกต้องกับพื้นที่ที่ได้จากการแบ่งส่วนรูปภาพกับพื้นที่ทั้งสองแบบรวมกัน เมื่อพิจารณาค่าเฉลี่ย IoU ของทั้ง 2 วิธี พบว่าวิธีการที่นำเสนอโดย Yangyang ให้ค่าเฉลี่ย IoU เท่ากับ 41% ในขณะที่วิธีการที่นำเสนอให้ค่าเฉลี่ย IoU เท่ากับ 54% จะเห็นได้ว่าวิธีที่นำเสนอให้ผลลัพธ์ที่ดีกว่า โดยวิธีที่นำเสนอสามารถลดรายละเอียดของพื้นหลังรูปภาพได้มากกว่า ด้วยเหตุนี้ทำให้ Bounding box สามารถล้อมรอบบริเวณที่สนใจโดยมีรายละเอียดที่ไม่เกี่ยวข้องในรูปภพน้อยกว่าวิธีการเดิม ดังแสดงในตารางที่ 3 และตาราง ที่ 4

ตารางที่ 3. แสดงตัวอย่างการแบ่งส่วนรูปภาพของทั้ง 2 วิธี

ตัวอย่าง ประเภท ของ ดอกไม้	รูปภาพตั้งต้น	วิธีการที่ นำเสนอโดย Yang Yang	วิธีการที่ นำเสนอ
เดซี่			
แดนดิไลออน			
กุหลาบ			



ตารางที่ 4. แสดงผลลัพธ์ตัวอย่างรูปภาพจากวิธีการแบ่งรูปภาพทั้ง 2 วิธี

วิธีที่นำเสนอโดย Yanyang	วิธีที่นำเสนอในงานวิจัย ชิ้นนี้

สำหรับข้อจำกัดของงานวิจัยที่นำเสนอ พบว่าวิธีการสามารถแบ่งส่วนรูปภาพได้ดีหากภายในรูปภาพมีวัตถุที่สนใจเป็นองค์ประกอบส่วนใหญ่ภายในภาพ หากบริเวณที่สนใจเป็นเพียงบริเวณเล็กภายในรูปภาพโอกาสที่จะเลือกบริเวณนั้นจะน้อยลงไป ซึ่งปัญหาเหล่านี้เป็นผลจากการใช้วิธีเลือกพื้นที่โดยการเลือกพื้นที่ที่ติดกันมากที่สุด

4.2 ผลการเปรียบเทียบประสิทธิภาพในการจำแนกประเภทของดอกไม้

ในการวัดประสิทธิภาพในการสกัดคุณลักษณะของพื้นที่ดอกไม้ด้วยวิธีการที่นำเสนอโดย Yangyang และวิธีที่นำเสนอในงานชิ้นนี้กับชุดข้อมูลดั้งเดิม ด้วยการนำชุดข้อมูลรูปภาพทั้ง 3 ชุดไปจำแนกประเภทดอกไม้โดยใช้แบบจำลอง CNN โดยจะวัดประสิทธิภาพด้วยค่าความถูกต้อง (Accuracy) ความแม่นยำ (Precision) ความครบถ้วน (Recall) และ F1 [23] ดังแสดงในสมการที่ (17) ถึง (20)

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (17)$$

$$Precision = \frac{TP}{TP+FP} \quad (18)$$

$$Recall = \frac{TP}{TP+FN} \quad (19)$$

$$F_1 = 2 \cdot \frac{Precision \cdot Recall}{Precision+Recall} \quad (20)$$

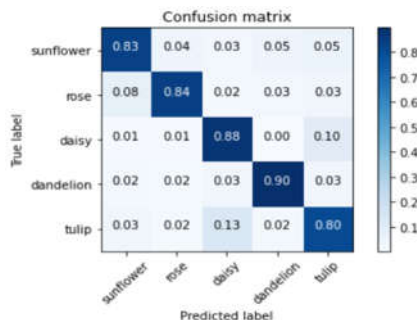
จากการทดลองพบว่าความแม่นยำในการจำแนกข้อมูลดั้งเดิมด้วยแบบจำลอง CNN ให้ความถูกต้อง 84% และความถูกต้องของชุดข้อมูลที่ได้จากวิธีการแบ่งส่วนรูปภาพด้วยวิธีของ Yangyang คือ 85% ในขณะที่ความถูกต้องของวิธีที่นำเสนอในงานชิ้นนี้คือ 87% นอกจากนี้ค่าความแม่นยำ ค่าความครบถ้วน และ F1 ของชุดข้อมูลทั้ง 3 ชุดได้แสดงผลลัพธ์ไว้ ดังแสดงในตารางที่ 5

ตารางที่ 5. แสดงค่า Precision, Recall และ F1 ที่ได้จากการจำแนกข้อมูลทั้ง 3 แบบด้วยแบบจำลอง CNN

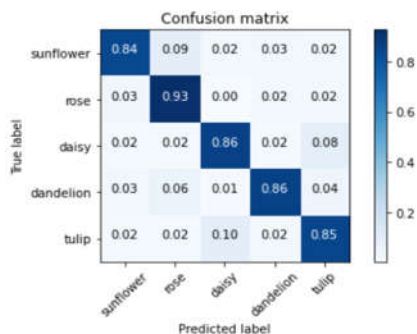
ชุดข้อมูล	ความแม่นยำ	ความครบถ้วน	F1
ชุดข้อมูลดั้งเดิม	84%	84%	84%
วิธีการแบ่งส่วนรูปภาพของ Yangyang	85%	85%	85%

ชุดข้อมูล	ความแม่นยำ	ความครบถ้วน	F1
วิธีที่นำเสนอในงานวิจัย ชิ้นนี้	87%	87%	87%

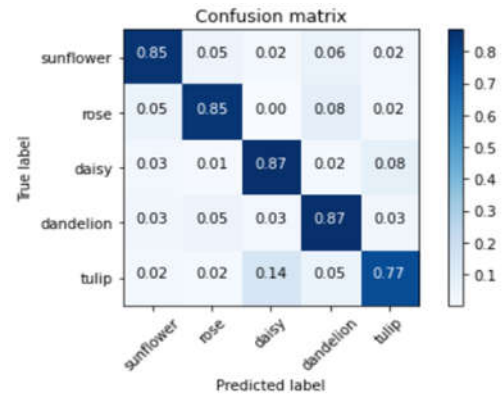
จากผลการทดลองได้แสดงให้เห็นว่าการนำชุดข้อมูลรูปภาพดอกไม้ไปผ่านขั้นตอนการแบ่งส่วนรูปภาพก่อนนำชุดรูปภาพไปจำแนกประเภทจะช่วยให้สามารถจำแนกประเภทดอกไม้ได้ดียิ่งขึ้น โดยการใช้วิธีที่นำเสนอในงานวิจัยนี้ในขั้นตอนการเตรียมรูปภาพ (Preprocessing) สามารถเพิ่มค่าความถูกต้อง ความแม่นยำ ความครบถ้วน และ F1 ได้ดีกว่าการใช้วิธีแบ่งส่วนรูปภาพของ Yang Yang อีกทั้งยังช่วยให้ค่า True positive ของของดอกไม้บางชนิดเพิ่มมากขึ้น เช่น ค่า True positive ของดอกทิวลิปที่มีค่ามากขึ้นเมื่อมีการเพิ่มขั้นตอนการเตรียมชุดข้อมูลด้วยการแบ่งรูปภาพ ดังแสดงในรูปที่ 2 ถึง 4



รูปที่ 2. แสดง Confusion matrix ของชุดข้อมูลตั้งต้น



รูปที่ 3. แสดง Confusion matrix ของวิธีการแบ่งส่วนรูปภาพของ Yangyang



รูปที่ 4. แสดง Confusion matrix ของวิธีที่นำเสนอในงานวิจัยชิ้นนี้

5. บทสรุปและการอภิปราย

ในงานวิจัยชิ้นนี้ได้นำเสนอแนวคิดการแบ่งส่วนรูปภาพ โดยอิงการใช้ประโยชน์จาก Saliency map ในการเก็บรายละเอียดของบริเวณที่สนใจภายในภาพ และการใช้ประโยชน์จากปริภูมิสี HSV ผสมกับการประยุกต์ใช้ Color mask ในการช่วยลดรายละเอียดที่ไม่สำคัญภายในพื้นหลังของรูปภาพออกไป จากการทดลองพบว่าวิธีนี้สามารถเก็บบริเวณของวัตถุที่สนใจได้ดี และสามารถลดรายละเอียดที่ไม่เกี่ยวข้องของพื้นหลังที่มีความซับซ้อนมากได้ค่อนข้างมาก ทำให้สามารถเลือกบริเวณที่สำคัญภายในภาพสำหรับนำไปใช้ในขั้นตอนการจำแนกประเภทดอกไม้ได้ค่อนข้างดี เมื่อเปรียบเทียบผลลัพธ์ระหว่างวิธีการแบ่งส่วนรูปภาพด้วยวิธีการของ Yangyang ที่ใช้วิธีการปรับเฉดสีของ Saliency map ให้เป็นเฉดสีเทาแล้วจึงปรับพื้นหลังของรูปภาพให้เป็นสีดำซึ่งสามารถลดรายละเอียดพื้นหลังของรูปภาพที่มีความซับซ้อนน้อยจนถึงปานกลางเท่านั้น จากผลการทดลองทำการวัดประสิทธิภาพการแบ่งส่วนรูปภาพวิธีการที่นำเสนอด้วยค่าเฉลี่ย IoU พบว่าค่าเฉลี่ย IoU ของวิธีการที่นำเสนอเท่ากับ 54% ซึ่งมากกว่าค่าเฉลี่ย IoU ของวิธีการของ Yangyang ซึ่งมีค่าเท่ากับ 41% และเมื่อนำชุดรูปภาพที่ไม่ผ่านขั้นตอนการแบ่งส่วนรูปภาพกับชุดข้อมูลที่

ผ่านขั้นตอนการเตรียมข้อมูลด้วยการแบ่งส่วนรูปภาพจากทั้ง 2 วิธีไปใช้จำแนกประเภทของดอกไม้ด้วยแบบจำลอง VGG16 ที่ผ่านการทำ Fine tuned พบว่าการเตรียมข้อมูลรูปภาพด้วยการแบ่งส่วนรูปภาพช่วยให้แบบจำลองสามารถจำแนกประเภทดอกไม้ได้ดีมากยิ่งขึ้น โดยผลลัพธ์แสดงให้เห็นได้จากค่าความถูกต้อง ความแม่นยำ ความครบถ้วน และ F1 ที่เพิ่มขึ้น โดยชุดข้อมูลที่แบ่งส่วนรูปภาพด้วยวิธีการที่นำเสนอช่วยให้จำแนกประเภทดอกไม้ได้ถูกต้อง 87% ในขณะที่ชุดข้อมูลที่ใช้การแบ่งส่วนรูปภาพด้วยวิธีการของ Yangyang สามารถจำแนกได้ถูกต้อง 85% ซึ่งมากกว่าชุดข้อมูลที่ไม่ผ่านขั้นตอนการแบ่งส่วนรูปภาพที่จำแนกได้ถูกต้อง 84% นอกจากนี้การเตรียมข้อมูลด้วยการแบ่งส่วนรูปภาพยังช่วยให้ค่า True positive ของดอกไม้บางชนิดเพิ่มมากขึ้น เช่น ค่า True positive ของดอกทิวลิปที่มีค่ามากขึ้น

สำหรับข้อจำกัดของงานวิจัยที่นำเสนอ ผู้วิจัยพบว่าวิธีการนี้สามารถแบ่งส่วนรูปภาพได้ดีหากภายในรูปภาพมีวัตถุที่สนใจเป็นองค์ประกอบส่วนใหญ่ภายในภาพ หากบริเวณที่สนใจเป็นเพียงบริเวณเล็กภายในรูปภาพโอกาสที่จะเลือกบริเวณนั้นก็จะน้อยลงไป ซึ่งปัญหาเหล่านี้เป็นผลจากการใช้วิธีเลือกพื้นที่โดยการเลือกพื้นที่ที่ติดกันมากที่สุด ดังนั้นสำหรับงานวิจัยที่จะขยายต่อไปในอนาคตอาจมีการนำประเด็นนี้ไปพิจารณาและศึกษาต่อเพื่อปรับปรุงงานให้ดียิ่งขึ้น

เอกสารอ้างอิง

- [1] P. Panwar, G. Gopal, and R. Kumar, "Image Segmentation using K-means clustering and Thresholding," *Image*, vol. 3, no. 05, pp. 1787-1793, 2016.
- [2] T.-W. Chen, Y.-L. Chen, and S.-Y. Chien, "Fast image segmentation based on K-Means clustering with histograms in HSV color space," in *2008 IEEE 10th Workshop on Multimedia Signal Processing*, IEEE, pp. 322-325, 2008.
- [3] N. Sabri, Z. Ibrahim, and N. N. Rosman, "K-means vs. fuzzy C-means for segmentation of orchid flowers," in *2016 7th IEEE Control and System Graduate Research Colloquium (ICSGRC)*, IEEE, pp. 82-86, 2016.
- [4] Y. Liu, F. Tang, D. Zhou, Y. Meng, and W. Dong, "Flower classification via convolutional neural network," in *2016 IEEE International Conference on Functional-Structural Plant Growth Modeling, Simulation, Visualization and Applications (FSPMA)*, IEEE, pp. 110-116, 2016.
- [5] M.-M. Cheng, N. J. Mitra, X. Huang, P. H. Torr, and S.-M. Hu, "Global contrast based salient region detection," *IEEE transactions on pattern analysis and machine intelligence*, vol. 37, no. 3, pp. 569-582, 2014.
- [6] P. Sathya and R. Kayalvizhi, "Modified bacterial foraging algorithm based multilevel thresholding for image segmentation," *Engineering Applications of Artificial Intelligence*, vol. 24, no. 4, pp. 595-615, 2011.
- [7] K. Bhargavi and S. Jyothi, "A survey on threshold based segmentation technique in image processing," *International Journal of Innovative Research and Development*, vol. 3, no. 12, pp. 234-239, 2014.
- [8] W. Wang, L. Duan, and Y. Wang, "Fast image segmentation using two-dimensional Otsu based on estimation of distribution algorithm," *Journal of Electrical and Computer Engineering*, vol. 2017, 2017.
- [9] L. T. Thanh and D. N. Thanh, "An adaptive local thresholding roads segmentation method for satellite aerial images with normalized HSV and lab color models," in *Intelligent Computing in Engineering*: Springer, pp. 865-872, 2020.

- [10] J. Yadav and M. Sharma, "A Review of K-mean Algorithm," *Int. J. Eng. Trends Technol*, vol. 4, no. 7, pp. 2972-2976, 2013.
- [11] I. Garg and B. Kaur, "Color based segmentation using K-mean clustering and watershed segmentation," in *2016 3rd International Conference on Computing for Sustainable Global Development (INDIACom)*, IEEE, pp. 3165-3169, 2016.
- [12] A. Sharma and S. Sehgal, "Image segmentation using firefly algorithm," in *2016 International Conference on Information Technology (InCITe)-The Next Generation IT Summit on the Theme-Internet of Things: Connect your Worlds*, IEEE, pp. 99-102, 2016.
- [13] X. Zheng, Q. Lei, R. Yao, Y. Gong, and Q. Yin, "Image segmentation based on adaptive K-means algorithm," *EURASIP Journal on Image and Video Processing*, vol. 2018, no. 1, pp. 1-10, 2018.
- [14] M. R. Hassan, R. R. Ema, and T. Islam, "Color image segmentation using automated K-means clustering with RGB and HSV color spaces," *Global Journal of Computer Science and Technology*, 2017.
- [15] Y. Chen, W. Xu, F. Kuang, and S. Gao, "The Study of Randomized Visual Saliency Detection Algorithm," *Computational and mathematical methods in medicine*, vol. 2013, 2013.
- [16] C. Cooley, S. Coleman, B. Gardiner, and B. Scotney, "Saliency detection and object classification," in *Proc. 19th Irish Machine Vision and Image Processing Conf.(IMVIP 2017)*, pp. 84-90, 2017.
- [17] Y. Yangyang and F. Xiang, "A Flower Image Classification Algorithm Based on Saliency Map and PCANet," *Journal of Communication and Computer*, vol. 15, pp. 14-24, 2019.
- [18] M. Deswal and N. Sharma, "A fast HSV image color and texture detection and image conversion algorithm," *International Journal of Science and Research (IJSR)*, vol. 3, no. 6, 2014.
- [19] M. Qin, Y. Xi, and F. Jiang, "A new improved convolutional neural network flower image recognition model," in *2019 IEEE Symposium Series on Computational Intelligence (SSCI)*, IEEE, pp. 3110-3117, 2019.
- [20] M. A. H. Abas, N. Ismail, A. I. M. Yassin, and M. N. Taib, "VGG16 for plant image classification with transfer learning and data augmentation," *International Journal of Engineering and Technology (UAE)*, vol. 7, pp. 90-94, 2018.
- [21] D. Theckedath and R. Sedamkar, "Detecting affect states using VGG16, ResNet50 and SE-ResNet50 networks," *SN Computer Science*, vol. 1, no. 2, pp. 1-7, 2020.
- [22] A. Najjar and E. Zagrouba, "Flower image segmentation based on color analysis and a supervised evaluation," in *2012 International Conference on Communications and Information Technology (ICCIT)*, IEEE, pp. 397-401, 2012.
- [23] A. H. Ornek and M. Ceylan, "Comparison of traditional transformations for data augmentation in deep learning of medical thermography," in *2019 42nd International Conference on Telecommunications and Signal Processing (TSP)*, IEEE, pp. 191-194, 2019.



การเปรียบเทียบความแม่นยำการจำแนกประเภทข้อมูลอนุกรมเวลา

ในปริภูมิเวกเตอร์ ระหว่างวิธีแซ็กและวิธีบอส:

กรณีศึกษา ข้อมูลคลื่นไฟฟ้าหัวใจ

The Accuracy Comparison of Time Series Classification in Vector Space between SAX and BOSS Methods: A Case Study of Electrocardiogram

นภัสสร แก้วกล้า* และ อัครินทร์ ไพบูลย์พานิช

Napatsorn Kaewkla and Akarin Phaibulpanich*

ภาควิชาสถิติ จุฬาลงกรณ์มหาวิทยาลัย

Department of Statistics, Chulalongkorn University

Received: November 25, 2021; Revised: December 17, 2021; Accepted: December 23, 2021; Published: December 28, 2021

ABSTRACT – The electrocardiogram (ECG) is an important procedure used to diagnose heart disorders. However, the ECG may contain different types of noise due to various of factors, potentially resulting in diagnostic errors. This research compares Symbolic Aggregate Approximation in Vector Space (SAXVSM) and Bag of Symbolic Fourier Approximation Symbols in Vector Space (BOSSVS) methods for classifying ECG data with noise. To choose a suitable classification algorithm for ECG5000 dataset, which is available in the Physionet database. Four types of ECG noises were simulated and then added to the data as follow: 1) Electromyography (EMG) 2) Powerline Interference 3) Baseline Wander and 4) Composite at 25%, 50% and 100% levels for the performance comparison of the ECG classification between normal and abnormal heart rhythms with SAXVSM and BOSSVS. The results show that both algorithms have similar high performance for all 13 datasets: accuracy and F1 score are 97-99%, precision is 95-99%, and recall is 97-100%, but BOSSVS has a longer running time than SAXVSM.

KEYWORDS: Time Series Classification, SAX, BOSS, Vector Space, Electrocardiogram

บทคัดย่อ -- การตรวจคลื่นไฟฟ้าหัวใจ เป็นหัตถการสำคัญที่ใช้วินิจฉัยความผิดปกติของหัวใจ แต่การตรวจวัดคลื่นไฟฟ้าหัวใจก็อาจมีสัญญาณรบกวนแบบต่าง ๆ ที่เกิดขึ้นจากหลายสาเหตุ ซึ่งอาจทำให้ผลการวินิจฉัยทางการแพทย์ผิดพลาด งานวิจัยนี้มีวัตถุประสงค์เพื่อเปรียบเทียบอัลกอริทึมสำหรับการจำแนกประเภทข้อมูลคลื่นไฟฟ้าหัวใจที่มีสัญญาณรบกวนด้วย Symbolic Aggregate Approximation in Vector Space (SAXVSM) และ Bag of Symbolic Fourier

*Corresponding Author: 6280156126@student.chula.ac.th

Approximation Symbols in Vector Space (BOSSVS) เพื่อเลือกอัลกอริทึมการจำแนกประเภทข้อมูลคลื่นไฟฟ้าหัวใจอย่างเหมาะสม โดยใช้ข้อมูลคลื่นไฟฟ้าหัวใจ ECG5000 จากฐานข้อมูล Physionet และผู้วิจัยได้จำลองสัญญาณรบกวนในคลื่นไฟฟ้าหัวใจ 4 แบบ ได้แก่ 1) Electromyography (EMG) 2) Powerline Interference 3) Baseline Wander และ 4) Composite ที่ระดับ 25% 50% และ 100% เพื่อเปรียบเทียบประสิทธิภาพของการจำแนกประเภทการเต้นของหัวใจปกติและผิดปกติด้วย SAXVSM และ BOSSVS จากการวิจัยสามารถสรุปได้ว่า สำหรับข้อมูล 13 ชุด ทั้ง SAXVSM และ BOSSVS มีประสิทธิภาพใกล้เคียงกัน โดยมีค่าความถูกต้องและคะแนน F1 อยู่ที่ 97-99% ค่าความแม่นยำอยู่ที่ 95-99% และค่าความระลึกอยู่ที่ 97-100% แต่ BOSSVS ใช้เวลาในการประมวลผลนานกว่า SAXVSM

คำสำคัญ: การจำแนกประเภทข้อมูลอนุกรมเวลา, เซ็ค, บอส, ปริภูมิเวกเตอร์, คลื่นไฟฟ้าหัวใจ

1. บทนำ

การตรวจคลื่นไฟฟ้าหัวใจ (Electrocardiogram: ECG) เป็นการวัดความต่างศักย์ของกระแสไฟฟ้าที่ไหลผ่านจุดกำเนิดไฟฟ้าไปยังเซลล์กล้ามเนื้อหัวใจ ซึ่งจะสัมพันธ์กับการบีบและคลายตัวของกล้ามเนื้อหัวใจ ดังนั้นการตรวจคลื่นไฟฟ้าหัวใจ จึงเป็นหัตถการสำคัญที่ใช้ในการวินิจฉัยภาวะผิดปกติของหัวใจ โดยจะใช้อิเล็กโทรด (Electrode) หรือขั้วไฟฟ้า ซึ่งต่ออยู่กับเครื่องบันทึกคลื่นไฟฟ้าหัวใจ มาวางไว้เหนือผิวหนังบริเวณต่าง ๆ ทำให้เกิดความต่างศักย์ไฟฟ้าระหว่างส่วนต่าง ๆ หรือที่เรียกว่าลีด (Lead) และบันทึกกระแสไฟฟ้าที่เกิดขึ้นเป็นรูปคลื่นต่าง ๆ บนกระดาษกราฟ อย่างไรก็ตามการตรวจวัดคลื่นไฟฟ้าหัวใจก็อาจมีปัจจัยหรือสัญญาณรบกวน (Noise) ซึ่งเกิดขึ้นจากหลายสาเหตุ เช่น ตำแหน่งของการติดเครื่องมือผิดพลาด การเคลื่อนไหวร่างกาย การรบกวนของสนามแม่เหล็กไฟฟ้า เป็นต้น ซึ่งสัญญาณรบกวนเหล่านี้อาจทำให้เกิดความผิดพลาดในการวินิจฉัยของแพทย์ได้ [1] ที่ผ่านมามีงานวิจัยที่นำเสนอตัวแบบการจำแนกประเภทข้อมูลอนุกรมเวลาที่มีประสิทธิภาพดี ได้แก่ Symbolic Aggregate Approximation - Vector Space Model (SAXVSM) [2, 3] และ Bag-Of-SFA-Symbols in Vector Space (BOSSVS) [4-6]

ขั้นตอนวิธี Symbolic Aggregate Approximation (SAX) และ Bag of Symbolic Fourier Approximation Symbols (BOSS) มีแนวคิดสำหรับการวิเคราะห์ข้อมูลอนุกรมเวลาด้วยการแบ่งอนุกรมเวลาออกเป็นช่วงและพยายามแปลงข้อมูลอนุกรมเวลานั้นให้เป็นลำดับของตัวอักษร เพื่อลดมิติของข้อมูล โดยใช้ Piecewise Aggregate Approximation (PAA) และ Discrete Fourier Transform (DFT) ตามลำดับ ในขณะเดียวกันก็ยังคงเก็บคุณลักษณะสำคัญของอนุกรมเวลาเดิมไว้ นอกจากนี้ยังไม่ค่อยมี

ผลกระทบต่อสัญญาณรบกวน แต่ยังไม่มีการวิจัยใช้สองวิธีนี้เพื่อศึกษาประสิทธิภาพการจำแนกประเภทคลื่นไฟฟ้าหัวใจภาวะหัวใจห้องล่างเต้นผิดจังหวะแบบ PVC ที่มีสัญญาณรบกวนแบบต่าง ๆ ผู้วิจัยจึงประยุกต์ใช้ SAX และ BOSS ร่วมกับปริภูมิเวกเตอร์ (Vector Space) ในการจำแนกประเภทข้อมูลและจำลองสัญญาณรบกวนในคลื่นไฟฟ้าหัวใจ 4 แบบ คือ 1) Electromyography (EMG) 2) Powerline Interference 3) Baseline Wander และ 4) Composite Noise เพื่อเปรียบเทียบประสิทธิภาพการจำแนกประเภทของทั้งสองวิธีนี้

2. งานวิจัยและทฤษฎีที่เกี่ยวข้อง

Kaya & Pehlivan (2015) [7] เปรียบเทียบประสิทธิภาพ การจำแนกประเภทข้อมูลคลื่นไฟฟ้าหัวใจในภาวะหัวใจเต้นผิดจังหวะ ที่เกิดจากกระแสไฟฟ้าออกจากหัวใจห้องล่างแบบ Premature Ventricular Contraction (PVC) โดยใช้ข้อมูลคลื่นไฟฟ้าหัวใจจาก The MIT-BIH Arrhythmia Database แล้วสุ่มจังหวะการเต้นของหัวใจทั้งหมด 7,000 จังหวะ ประกอบด้วยปกติ (3,500 จังหวะ) และ PVC (3,500 จังหวะ) แล้วผ่านขั้นตอนการกรองสัญญาณรบกวนด้วยค่ามัธยฐาน (Median Filtering) จากนั้นได้เปรียบเทียบวิธีการลดมิติข้อมูล (Feature Reduction) ด้วย 3 วิธี คือ 1) การวิเคราะห์องค์ประกอบหลัก (Principal Components Analysis: PCA) 2) การวิเคราะห์องค์ประกอบอิสระ (Independent Component Analysis: ICA) 3) การทำแผนที่โยงก่อร่างตัวเอง (Self-Organizing Maps: SOM) แล้วเปรียบเทียบประสิทธิภาพการจำแนกประเภทด้วย 4 ตัวแบบ คือ 1) โครงข่ายประสาทเทียม (Neural Networks: NN)

2) การหาเพื่อนบ้านที่ใกล้ที่สุด (K-Nearest Neighbour: K-NN)
3) ซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machines: SVM)
และ 4) ต้นไม้ตัดสินใจ (Decision Tree: DT) ในภาพรวมตัวแบบ K-NN ใช้เวลาในการจำแนกประเภทน้อยและมีประสิทธิภาพการจำแนกประเภทดีที่สุด

Senin & Malinchik (2013) [3] นำเสนอตัวแบบ Symbolic Aggregate Approximation - Vector Space Model (SAX-VSM) แนวคิดหลักของตัวแบบนี้ คือ การแบ่งอนุกรมเวลาออกเป็น อนุกรมเวลาย่อย (Sliding Window) และใช้ Piecewise Aggregate Approximation (PAA) เพื่อลดมิติของข้อมูลอนุกรมเวลา แล้วจึงใช้แซ็ค (SAX) เพื่อแปลงข้อมูลอนุกรมเวลาให้เป็นรูปแบบกลุ่มตัวอักษรและใช้ตัวแบบปริภูมิเวกเตอร์ (Vector Space Model) เพื่อจำแนกประเภทข้อมูลอนุกรมเวลาจากรูปแบบกลุ่มตัวอักษรเหล่านี้ ในงานวิจัยนี้ได้เปรียบเทียบอัตราความคลาดเคลื่อนการจำแนกประเภทด้วย 5 ตัวแบบ คือ 1) INN-Euclidean 2) INN-Dynamic Time Wrapping (DTW) 3) Fast Shapelets และ 4) Bag of Patterns และ 5) SAX-VSM โดยใช้ข้อมูลจาก UCR Time Series ทั้งหมด 19 ชุด แสดงให้เห็นว่า มีข้อมูล 12 ชุดจากทั้งหมด 19 ชุด ที่ SAX-VSM ให้อัตราความคลาดเคลื่อน (Error Rate) ต่ำที่สุด

Schäfer (2014, 2015) [4-6] นำเสนอตัวแบบ Bag-Of-SFA-Symbols in Vector Space (BOSS VS) มีแนวคิดมาจากการแปลงสัมประสิทธิ์ฟูเรียร์แบบไม่ต่อเนื่อง (Discrete Fourier Transform: DFT) จากงานวิจัยได้ทำการเปรียบเทียบประสิทธิภาพและเวลาสะสมที่ใช้ในการเรียนรู้ (Training) และทดสอบ (Testing) โดยทดลองกับข้อมูลอนุกรมเวลาจาก UCR Time Series และข้อมูลอื่นๆ รวมทั้ง 91 ชุด พบว่า BOSS ใช้เวลาการประมวลผลสะสมน้อยและประสิทธิภาพการจำแนกประเภทโดยเฉลี่ยดีกว่า เมื่อเทียบกับ INN-Dynamic Time Wrapping (DTW) ในงานวิจัยระบุว่า DFT มีหลักการ คือ แปลงข้อมูลอนุกรมเวลาที่เป็นสัญญาณจากโดเมนเวลา (Time Domain) ไปเป็น โดเมนความถี่ (Frequency Domain) ทำให้สามารถลดสัญญาณรบกวนได้

คลื่นไฟฟ้าหัวใจ (Electrocardiogram: ECG) เป็นข้อมูลอนุกรมเวลา (Time Series) ซึ่งประกอบด้วย $T = (t_1, t_2, \dots, t_n)$ คือ ลำดับของค่าข้อมูลที่ถูกบันทึกตามช่วงเวลา โดยที่ $n \in N$

สัญญาณรบกวนในคลื่นไฟฟ้าหัวใจ คือ สัญญาณผิดปกติที่เกิดขึ้นในขณะที่บันทึกคลื่นไฟฟ้าหัวใจ ซึ่งอาจส่งผลกระทบต่อประสิทธิภาพของการตรวจจับสัญญาณการเต้นของหัวใจ โดย

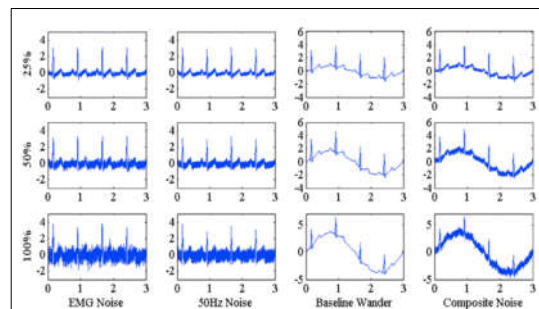
สัญญาณรบกวนที่มักพบในคลื่นไฟฟ้าหัวใจและนำมาศึกษาในงานวิจัยนี้มี 4 แบบ [8] ได้แก่

1) Electromyographic Noise (EMG) เป็นสัญญาณรบกวนที่เกิดจากการหดเกร็งของกล้ามเนื้อหรือการเคลื่อนไหวของร่างกายอย่างกะทันหัน ความถี่ใน EMG ทับซ้อนกันมาก ทำให้ตรวจจับยาก

2) Powerline Interference (50 Hz) เป็นสัญญาณรบกวนความถี่สูง โดยทั่วไปจะมีความถี่ประมาณ 50 เฮิรตซ์ มีลักษณะเป็นสัญญาณรบกวนแบบคลื่นไซน์ อาจจะมีฮาร์โมนิก (Harmonic) จำนวนหนึ่ง กล่าวคือ กระแสหรือแรงดันในรูปสัญญาณคลื่นไซน์ของสัญญาณหรือปริมาณในคาบใด ๆ ที่มีค่าเป็นจำนวนเท่าของความถี่หลักมูล (Fundamental Frequency) สาเหตุหลักเกิดจากการรบกวนคลื่นแม่เหล็กไฟฟ้าจากสายไฟ (Electromagnetic Field : EMF) หรือจากเครื่องจักรในระยะใกล้ๆ การต่อสายดินที่ไม่ดี เป็นต้น

3) Baseline Wander เป็นสัญญาณรบกวนความถี่ต่ำ มีลักษณะรบกวนจากแนวแกน x สูงขึ้นหรือลดต่ำลงจากแนวเส้นฐาน (Baseline) สาเหตุหลักเกิดจากการหายใจ การดิ้นขยับไฟฟ้าไม่ดี เป็นต้น

4) Composite Noise เป็นสัญญาณรบกวนที่เกิดจากการเกิด



ร่วมกันของสัญญาณรบกวนทั้ง 3 แบบที่กล่าวมา

รูปที่ 1. แสดงตัวอย่างสัญญาณรบกวนแต่ละแบบ

ที่ระดับ 25% 50% และ 100%

แหล่งที่มา : “Arrhythmia ECG Noise Reduction by Ensemble Empirical Mode Decomposition” โดย Chang, K. M., 2010

3. Symbolic Aggregate Approximation (SAX)

Symbolic Aggregate Approximation [2, 3] เป็นวิธีการแปลงข้อมูลอนุกรมเวลาให้เป็นรูปแบบของกลุ่มตัวอักษรเพื่อลดมิติของข้อมูล โดยมีขั้นตอนดังนี้

3.1 การทำข้อมูลให้เป็นมาตรฐาน (Z-normalization)

กำหนดให้ อนุกรมเวลา $T=(t_1, t_2, \dots, t_n)$ คือ ลำดับของค่าข้อมูลที่ถูกบันทึกตามช่วงเวลา โดยที่ $n \in N$ [9]

$$T_{norm} = \frac{T - \mu}{\sigma} \quad (1)$$

โดยที่ T คือ ค่าข้อมูลอนุกรมเวลา
 μ คือ ค่าเฉลี่ยของอนุกรมเวลา
 σ คือ ส่วนเบี่ยงเบนมาตรฐานของอนุกรมเวลา

3.2 การทำหน้าต่างบานเลื่อน (Sliding Windows)

การแบ่งอนุกรมเวลา $T=(t_1, t_2, \dots, t_n)$ ที่มีความยาว n ออกเป็นหน้าต่างขนาด w คงที่ จะได้

$$windows(T, w) = \left\{ \begin{matrix} S_{1:w} & , & S_{2:w} & , & \dots & , & S_{n-w+1:w} \\ (t_1, t_2, \dots, t_w) & (t_2, t_3, \dots, t_{w+1}) \end{matrix} \right\} \quad (2)$$

โดยที่ $i = 1, 2, \dots, n - w + 1$
 $S_{i:w}$ คือ หน้าต่างบานเลื่อนที่ i ความยาว w
 i คือ ลำดับของหน้าต่างแต่ละบาน
 ดังนั้น หน้าต่างสองบานที่ติดกันจะทับซ้อนกัน w ตำแหน่ง [6]

3.3 Piecewise Aggregate Approximation (PAA)

การลดมิติของข้อมูล (Dimension Reduction) ด้วย Piecewise Aggregate Approximation (PAA) เป็นวิธีการลดมิติของข้อมูลโดยการแบ่งข้อมูลออกเป็นช่วงย่อย ๆ ส่วนละเท่า ๆ กัน และคำนวณค่าเฉลี่ยของแต่ละส่วน เพื่อเป็นตัวแทนของข้อมูลในแต่ละส่วน

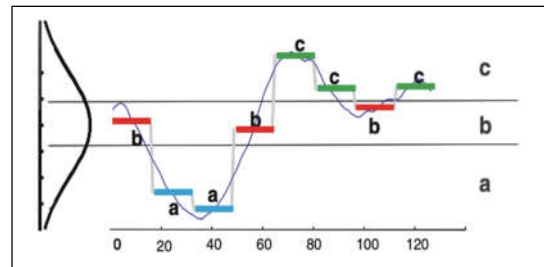
กำหนดให้ อนุกรมเวลามีความยาว w จุดเวลา จะแบ่งข้อมูลออกเป็น p ส่วน และหาค่าเฉลี่ยของข้อมูลในแต่ละส่วน [2]

$$\bar{c}_i = \frac{P}{w} \sum_{j=\frac{w}{p}(i-1)+1}^{\frac{w}{p}i} C_j \quad (3)$$

โดยที่ $i = 1, \dots, p$ และ $j = 1, \dots, w$
 $\bar{C}_i$ คือ ค่าเฉลี่ยของอนุกรมเวลาส่วนที่ i
 C_j คือ ค่าของข้อมูลอนุกรมเวลาจุดที่ j

3.4 การแบ่งช่วงข้อมูล(Discretization) ด้วยวิธีแซ็ค(SAX)

การแบ่งช่วงข้อมูล (Discretization) ด้วยวิธีแซ็ค (SAX) เป็นการแปลงค่าเฉลี่ยของข้อมูลในแต่ละส่วนที่ได้จากขั้นตอน PAA ให้เป็นตัวอักษร โดยจะต้องกำหนดจำนวนตัวอักษร (alphabet size) ที่ใช้ และแบ่งจำนวนช่วงของข้อมูลให้เท่ากับจำนวนตัวอักษร โดยที่พื้นที่ใต้กราฟในแต่ละช่วงจะต้องมีขนาดเท่า ๆ กัน จะได้จุดแบ่ง (Breakpoint: β) แต่ละช่วง ด้วยค่ามาตรฐาน (Z-Score) โดยกำหนดความยาวเท่ากับ 8 และจำนวนตัวอักษรเท่ากับ 3 (A, B, C)



รูปที่ 2. แสดงตัวอย่างการแปลงข้อมูลอนุกรมเวลา

โดยวิธีแซ็ค (SAX) แหล่งที่มา : “Experiencing SAX: a novel symbolic representation of time series” โดย Lin et al., 2007

3.5 การลดขนาดข้อมูลด้วย Numerosity Reduction

Numerosity Reduction จะนับรูปแบบของตัวอักษรที่ซ้ำกันในหน้าต่างบานเลื่อนที่ติดกันเพียงครั้งเดียวและจะนับซ้ำอีกครั้งก็ต่อเมื่อ พบรูปแบบของตัวอักษรเดิมที่เกิดขึ้นอีกในหน้าต่างลำดับอื่น ๆ ที่ไม่ติดกัน [2]

4. Bag of Symbolic Fourier Approximation

Symbols (BOSS)

การประมาณฟูเรียร์เชิงสัญลักษณ์ Symbolic Fourier Approximation (SFA) [4-6] เป็นอีกวิธีหนึ่งในการแปลงข้อมูลอนุกรมเวลาให้เป็นรูปแบบตัวอักษร โดยมี 2 ขั้นตอนแรก คือ การทำข้อมูลให้เป็นมาตรฐาน (Z-Normalization) และการทำหน้าต่างบานเลื่อน (Sliding Windows) ตามที่ได้กล่าวมาในขั้นตอนของแซ็ค (SAX) และมีขั้นตอนต่อไปดังนี้

4.1 การประมาณค่าสัมประสิทธิ์การแปลงฟูเรียร์แบบไม่ต่อเนื่อง

การประมาณค่าการแปลงฟูเรียร์แบบไม่ต่อเนื่อง (Discrete Fourier Transform: DFT) เป็นการประมาณค่าสัมประสิทธิ์ของอนุกรมฟูเรียร์จากข้อมูลสัญญาณรายจุดหรือสัญญาณไม่ต่อเนื่อง (Discrete) รายคาบตั้งแต่ 0 ถึง N-1 ที่แบ่งมาจากสัญญาณต่อเนื่อง เพื่อที่จะแปลงสัญญาณโดเมนเวลาให้เป็นโดเมนความถี่ [10]

$$X(m) = \sum_{n=0}^{N-1} x(n) e^{-j \frac{2\pi n m}{N}} \quad (4)$$

โดยที่ $n, m = 0, 1, \dots, N-1$ และ $j = \sqrt{-1}$

$x(n)$ คือ ลำดับของตัวเลขจากสัญญาณไม่ต่อเนื่อง

$X(m)$ คือ ผลการแปลงฟูเรียร์แบบไม่ต่อเนื่องของ $x(n)$

จากสมการของออยเลอร์ (Euler's Equation)

$$e^{j\theta} = \cos(\theta) + j \sin(\theta) \quad (5)$$

$$\text{โดยที่ } \theta = \frac{2\pi n m}{N}$$

ถ้าแทนสมการที่ (5) ลงในสมการที่ (4) จะได้

$$X(m) = \sum_{n=0}^{N-1} x(n) \left[\cos\left(\frac{2\pi n m}{N}\right) - j \sin\left(\frac{2\pi n m}{N}\right) \right] \quad (6)$$

จะได้ว่า $X(m)$ เป็นค่าสัมประสิทธิ์การแปลงฟูเรียร์ ซึ่งจะอยู่ในรูปของจำนวนเชิงซ้อน ประกอบด้วย (ส่วนจริง(m) , ส่วนจินตภาพ(m)) โดยที่ $m = 0, 1, 2, \dots, N-1$

4.2 การทำควอนไทซ์ (Quantisation)

การทำควอนไทซ์ (Quantisation) เป็นการกำหนดช่วงของข้อมูลที่จะแปลงค่าประมาณสัมประสิทธิ์การแปลงฟูเรียร์แบบไม่ต่อเนื่อง (DFT) เป็นตัวอักษร (alphabet) ด้วยเทคนิค The Multiple Coefficient Binning (MCB) [4] โดยมีการเตรียมข้อมูลดังนี้

กำหนดให้ เมทริกซ์ A ที่มาจากขั้นตอนการแปลงฟูเรียร์แบบไม่ต่อเนื่อง (DFT) ของชุดข้อมูลเรียนรู้ (Training Data) แสดงได้สมการที่ (7)

$$A = \begin{pmatrix} \text{DFT}(T_1) \\ \text{DFT}(T_2) \\ \vdots \\ \text{DFT}(T_N) \end{pmatrix} = \begin{pmatrix} \text{real}_{1,1} & \text{img}_{1,1} & \dots & \text{real}_{1,\frac{l}{2}} & \text{img}_{1,\frac{l}{2}} \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ \text{real}_{i,1} & \text{img}_{i,1} & \dots & \text{real}_{i,\frac{l}{2}} & \text{img}_{i,\frac{l}{2}} \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ \text{real}_{N,1} & \text{img}_{N,1} & \dots & \text{real}_{N,\frac{l}{2}} & \text{img}_{N,\frac{l}{2}} \end{pmatrix} = (C_1, C_2, \dots, C_l) \quad (7)$$

โดยที่ $i = 1, \dots, N$ $k = 1, \dots, \frac{l}{2}$ และ $j = 1, \dots, l$

T_i คือ ชุดข้อมูลเรียนรู้ (Training Data) แถวที่ i

real_{ik} คือ ค่าจริงของค่าสัมประสิทธิ์ฟูเรียร์ตัวที่ k ของข้อมูลเรียนรู้ตัวที่ i

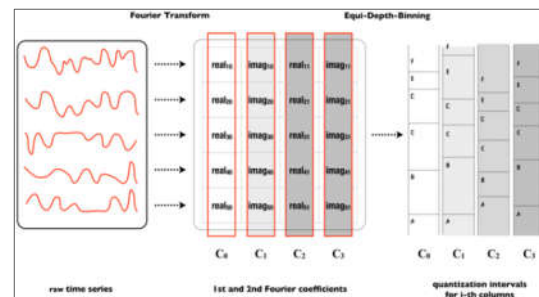
img_{ik} คือ ค่าจินตภาพของค่าสัมประสิทธิ์ฟูเรียร์ตัวที่ k ของข้อมูลเรียนรู้ตัวที่ i

C_j คือ ค่าสัมประสิทธิ์ฟูเรียร์คอลัมน์ที่ j

The Multiple Coefficient Binning (MCB) และมีขั้นตอนดังนี้
ขั้นตอนที่ 1 เรียงค่าของข้อมูลในแต่ละคอลัมน์ C_j

ขั้นตอนที่ 2 กำหนดจำนวนตัวอักษรเท่ากับ a และแบ่งข้อมูลออกเป็น a ส่วน โดยที่จำนวนข้อมูลในแต่ละส่วนเท่า ๆ กัน (equi-depth binning bins) จะได้ จุดแบ่ง (breakpoints) สำหรับคอลัมน์ C_j

ขั้นตอนที่ 3 กำหนดลำดับตัวอักษรประจำตำแหน่งในแต่ละส่วนในคอลัมน์ C_j จะได้ตัวอย่างการทำ The Multiple Coefficient Binning (MCB) แสดงได้ดังรูปที่ 3.



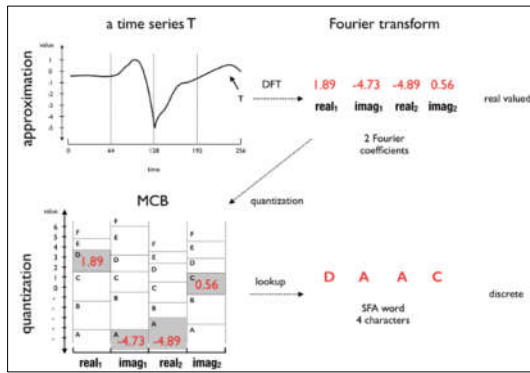
รูปที่ 3. แสดงขั้นตอน The Multiple Coefficient Binning (MCB)

แหล่งที่มา: “Human Activity Recognition Based on Symbolic Representation Algorithms for Inertial Sensors”

โดย Wesllen Sousa Lima et al., 2018

4.3 Symbolic Fourier Approximation

การแปลงฟูเรียร์เชิงสัญลักษณ์ (Symbolic Fourier Approximation) เป็นการแปลงค่าประมาณสัมประสิทธิ์ฟูเรียร์ (DFT) ให้เป็นตัวอักษร ด้วย Multiple Coefficient Binning (MCB)



รูปที่ 4. แสดงตัวอย่าง Symbolic Fourier Approximation
แหล่งที่มา: “Bag-Of-SFA-Symbols in Vector Space
(BOSS VS)” โดย Schäfer, 2015

5. แบบจำลองปริภูมิเวกเตอร์ (Vector Space Model

: VSM)

แบบจำลองปริภูมิเวกเตอร์ (Vector Space Model: VSM) เป็นวิธีหนึ่งในการแทนเอกสารที่ไม่มีโครงสร้าง (Unstructured Text Document) ด้วยแบบจำลองปริภูมิเวกเตอร์ โดยกำหนดให้เอกสารแต่ละฉบับเปรียบเสมือนเวกเตอร์ของค่าขนาดของเวกเตอร์ขึ้นอยู่กับจำนวนของคำที่ปรากฏอยู่ในเอกสารฉบับนั้น [11]

5.1 การหาความถี่ของคำและการผกผันความถี่ในเอกสาร

(Term Frequency-Inverse Document Frequency: TF-IDF) [5]

$$tf_{t,T} = \begin{cases} 1 + \log(B_T(t)) & ; \text{if } B_T(t) > 0 \\ 0 & ; \text{otherwise} \end{cases} \quad (8)$$

โดยที่ T คือ อนุกรมเวลา
 $B_T(t)$ คือ ความถี่ของคำ (t) ของอนุกรมเวลา T

$$idf_{t,C} = \begin{cases} 1 + \log(\sum_{T \in C} B_T(t)) & ; \text{if } \sum_{T \in C} B_T(t) > 0 \\ 0 & ; \text{otherwise} \end{cases} \quad (9)$$

โดยที่ $\sum_{T \in C} B_T(t)$ คือ ความถี่ของคำ (t) ของอนุกรมเวลา T ที่อยู่ในประเภท C

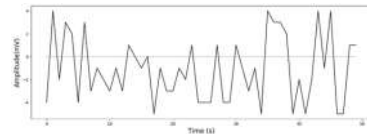
$$idf_{t,C} = \log(1 + \frac{|CLASSES|}{|\{C | T \in C \cap B_T(t) > 0\}|}) \quad (10)$$

โดยที่ $|\{C | T \in C \cap B_T(t) > 0\}|$ คือ จำนวนประเภท C ที่พบคำ T
 $|CLASSES|$ คือ จำนวนประเภท C ทั้งหมด

$tf \times idf(t,C)$ ใช้วัดความสำคัญของคำ (t) ในประเภท (C) แต่จะถูกลดทอนด้วยความสำคัญของคำ (t) ในประเภท (C) ทั้งหมด สามารถคำนวณได้ดังสมการที่ (10)

$$tf \times idf(t,C) = 1 + \log(\sum_{T \in C} B_T(t)) \times \log(1 + \frac{|CLASSES|}{|\{C | T \in C \cap B_T(t) > 0\}|}) \quad (11)$$

(1) การแปลงข้อมูลให้เป็นมาตรฐาน (Z-normalization)



(2) การทำหน้าต่างบานเลื่อน (Sliding windows)

$$windows(T, w) = \left\{ \begin{matrix} S_{1,w} & , & S_{2,w} & , & \dots & , & S_{n-w+1,w} \\ (t_1, t_2, \dots, t_w) & (t_2, t_3, \dots, t_{w+1}) \end{matrix} \right\}$$

(3) การแปลงฟูเรียร์เชิงสัญลักษณ์ด้วย Symbolic Aggregate Approximation (SAX) หรือ Bag of Symbolic Fourier Approximation: BOSS

$$S_{1,w} = (\text{aacbbabc}, \text{aacbbabc}, \dots, \text{baccbcaa}, \text{ccaccbaa})$$

$$S_{2,w} = (\text{baccbabc}, \text{acccccbc}, \dots, \text{baccbaba}, \text{baccbaba}, \text{ccaccbaa})$$

: $S_{n-w+1,w} = (\text{bacbbaa}, \text{cbcabbb}, \text{cbcabbb}, \dots, \text{abccaacc}, \text{accbbabc}, \text{bbcbbaaca})$				
(4) การลดขนาดข้อมูลด้วย Numerosity Reduction				
$S_{1,w} = (\text{aacbbabc}, \dots, \text{bacbcaa}, \text{ccacbaa})$				
$S_{2,w} = (\text{bacbbabc}, \text{acccccbc}, \dots, \text{baccaba}, \text{ccacbaa})$				
:				
$S_{n-w+1,w}$ $= (\text{bacbbaa}, \text{cbcabbb}, \dots, \text{abccaacc}, \text{accbbabc}, \text{bbcbbaaca})$				
(5) การนับความถี่ของรูปแบบตัวอักษรต่าง ๆ Bag Of SFA				
Symbol: BOSS หรือ Bag of Symbolic Fourier				
Approximation: BOSS				
Window	aacbbabc	abceaaa	bbceaaa	...
1	0	10	3	...
2	1	0	2	...
:	:	:	:	...
n	...	...	...	...
(6)				

รูปที่ 5. แสดงสรุปขั้นตอนของ SAX และ BOSS

5.2 ค่าความเหมือนโคไซน์ (Cosine similarity)

ค่าความเหมือนโคไซน์ (Cosine similarity) คือ การหาความคล้ายคลึงด้วยองศา เป็นการวัดความเหมือนของเวกเตอร์ 2 เวกเตอร์ว่าไปในทิศทางเดียวกันหรือไม่ โดยตัดขนาด (Magnitude) ของเวกเตอร์ออกไป [12]

$$\text{similarity}(Q, C) = \frac{\bar{Q} \cdot \bar{C}}{\|\bar{Q}\| \cdot \|\bar{C}\|} \quad (12)$$

$$\text{similarity}(Q, C) = \frac{\sum_{t \in Q} tf(t, Q) \cdot (tf \times idf(t, C) + 1)}{\sqrt{\sum_{t \in Q} (tf(t, Q))^2} \sqrt{\sum_{t \in C} (tf \times idf(t, Q))^2}} \quad (13)$$

$$\text{label}(Q) = \arg \max_{C \in \text{CLASSES}} (\text{similarity}(Q, C)) \quad (14)$$

โดยที่ $\text{similarity}(Q, C)$ คือ ค่าความเหมือนโคไซน์ของอนุกรมเวลา Q กับประเภท C

$\text{label}(Q)$ คือ ผลลัพธ์หรือประเภทของอนุกรม

เวลา Q $\sum_{t \in Q} tf(t, Q)$ คือ ความถี่ของค่า t ใน Q

$tf \times idf(t, C)$ คือ ความถี่ของค่าและ

การหักเหความถี่ (TF-IDF) ของค่าในประเภท C

6. การจำลองสัญญาณรบกวนในคลื่นไฟฟ้าหัวใจ

ในงานวิจัยนี้จะใช้วิธีการจำลองสัญญาณรบกวนในคลื่นไฟฟ้าหัวใจจาก งานวิจัย K. M. Chang, 2010 [8]

ค่าจำกัดความในการจำลองสัญญาณรบกวน

ค่าจำกัดความที่ 1: Function คือ ฟังก์ชันในการจำลองของสัญญาณรบกวน

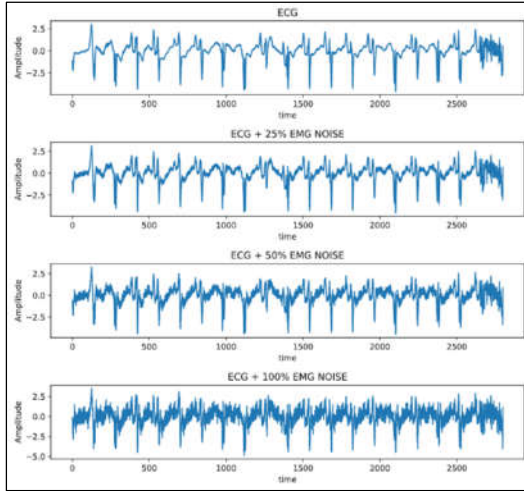
ค่าจำกัดความที่ 2: Vpp คือ ระยะขจัดตั้งแต่จุดสูงสุดถึงจุดต่ำสุดของคลื่นไฟฟ้าหัวใจปกติ

ค่าจำกัดความที่ 3: Frequency คือ ความถี่ของสัญญาณรบกวน

ค่าจำกัดความที่ 4: Reduced Ratio คือ อัตราส่วนสัญญาณต่อสัญญาณรบกวน

ตารางที่ 1. แสดงการจำลองสัญญาณรบกวน

Noise Type	Function	Vpp	Frequency	Reduced Ratio
EMG	Normal	5 mV	-	$\frac{1}{8}$
Powerline Interference (50Hz)	Sine	5 mV	50 Hz	$\frac{1}{4}$
Baseline Wander	Sine	5 mV	0.333 Hz	-
Composite	$0.5 \times (\text{EMG} + \text{Powerline Interference}) + \text{Baseline}$			



รูปที่ 6. แสดงตัวอย่างคลื่นไฟฟ้าหัวใจ 20 จังหวะ ที่ปกติและที่มีสัญญาณรบกวนแบบ EMG ระดับ 25% 50% และ 100%

7. เครื่องมือและวิธีการ

งานวิจัยนี้จะใช้ข้อมูล ECG5000 ซึ่งประกอบด้วยจังหวะการเต้นของหัวใจทั้งหมด 5,000 จังหวะ ความยาว 140 จุดเวลา และจัดประเภทการเต้นของหัวใจที่ปกติและผิดปกติไว้ทั้งหมด 5 ประเภท แต่ด้วยสัดส่วนข้อมูลประเภทที่ 2 3 และ 4 มีน้อยเกินไป ดังนั้นในงานวิจัยนี้จะใช้ข้อมูลในประเภท 0 (จังหวะการเต้นของหัวใจปกติ) และประเภท 1 (ภาวะหัวใจห้องล่างเต้นผิดจังหวะแบบพีวีซี R-ON-T) ทำให้เหลือข้อมูลทั้งหมด 4,686 จังหวะ และจำลองสัญญาณรบกวน 4 แบบ ได้แก่ 1) EMG Noise 2) Powerline Interference 3) Baseline Wander และ 4) Composite Noise และปรับระดับของสัญญาณรบกวนที่ 25% 50% 100% และไม่มีสัญญาณรบกวน เพื่อเปรียบเทียบประสิทธิภาพการจำแนกประเภทข้อมูลคลื่นไฟฟ้าหัวใจระหว่าง SAXVSM และ BOSSVS

งานวิจัยนี้ใช้ข้อมูลทั้งหมด 4,686 จังหวะ แต่ละจังหวะมีความยาว 140 จุดเวลา ดังตารางที่ 2.

ตารางที่ 2. แสดงตัวอย่างข้อมูลคลื่นไฟฟ้าหัวใจ

No.	t_1	...	t_{140}	Class
1	-0.6566696	...	0.2257076	1
2	0.3973526	...	0.52316457	0
3	-1.862747	...	-0.6100227	0

...	...	...	...	...
4686	-1.7946848	...	-1.2108684	0

สำหรับชุดที่ใช้ศึกษา ประกอบด้วย ข้อมูลคลื่นไฟฟ้าหัวใจที่มีจำลองสัญญาณรบกวน 4 แบบ ได้แก่ 1) EMG Noise 2) Powerline Interference 3) Baseline Wander และ 4) Composite Noise แต่ละแบบปรับระดับของสัญญาณรบกวนไว้ที่ 25% 50% 100% จะได้ชุดข้อมูลที่ใช้ศึกษา ดังตารางที่ 3.

ตารางที่ 3. แสดงชุดข้อมูลที่ใช้ศึกษา

ชุดข้อมูล	EMG Noise	Powerline Interference	Baseline Wander	Composite Noise
1	-	-	-	-
2	25%	-	-	-
3	50%	-	-	-
4	100%	-	-	-
7	-	25%	-	-
8	-	50%	-	-
9	-	100%	-	-
...	...	...	...	...
13	-	-	-	100%

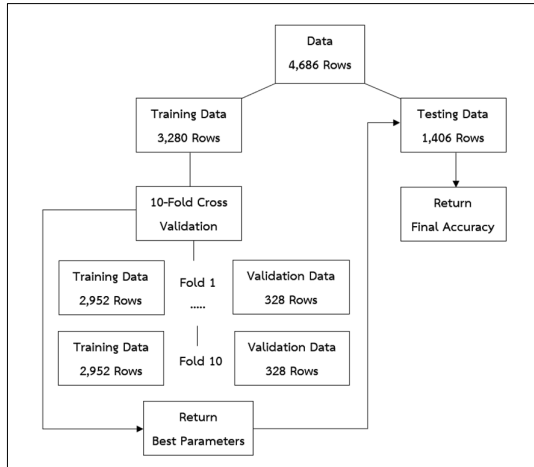
7.1 การแบ่งข้อมูล

จากข้อมูลทั้งหมดจะถูกแบ่งข้อมูลออกเป็น 2 ส่วน คือ ข้อมูลชุดเรียนรู้ (Training Data) สัดส่วน 70% (3,280 ตัวอย่าง) และข้อมูลชุดทดสอบ (Testing Data) สัดส่วน 30% (1,406 ตัวอย่าง) สำหรับข้อมูลชุดเรียนรู้ (Training Data) สัดส่วน 70% จะนำไปทำ 10-Folds Cross Validation เพื่อหาพารามิเตอร์ที่ดีที่สุด ส่วนข้อมูลชุดทดสอบ (Testing Data) สัดส่วน 30% จะนำไปทดสอบความแม่นยำ

7.2 การเลือกพารามิเตอร์

เริ่มจากการกำหนดขอบเขตล่างและขอบเขตบนสำหรับพารามิเตอร์แต่ละตัว จากนั้นจะสุ่มเลือกชุดของพารามิเตอร์เพื่อหาค่าความแม่นยำโดยเฉลี่ยของพารามิเตอร์ชุดนั้น ๆ สำหรับ

SAX และ BOSS โดยจะใช้ชุดข้อมูลเรียนรู้ (Training Data) ทำ K-fold Cross Validation เพื่อแบ่งข้อมูลออกเป็นชุดข้อมูลเรียนรู้ (Training Data) และชุดข้อมูลตรวจสอบ (Validation Data)



รูปที่ 7. แสดงแผนภาพการแบ่งข้อมูล

ตารางที่ 4. แสดงพารามิเตอร์

พารามิเตอร์	ความหมาย
Window Size (w)	ขนาดหน้าต่างบานเลื่อน
Word Size (p)	ความยาวของตัวอักษร
Alphabet Size (a)	จำนวนตัวอักษร

สำหรับในแต่ละส่วน (fold) มีขั้นตอนย่อย ดังนี้

ขั้นตอนที่ 1 ใช้ชุดข้อมูลเรียนรู้ (Training Data) หาความถี่ของรูปแบบตัวอักษร (Bag of Patterns) โดยใช้วิธีของ SAX และ BOSS

ขั้นตอนที่ 2 สร้างเมทริกซ์ความถี่ของคำ (Term Frequency Matrix)

ขั้นตอนที่ 3 คำนวณความถี่ของคำและการหักล้างความถี่ในเอกสาร (tf-idf) ของรูปแบบต่าง ๆ ในแต่ละประเภท (Class)

ขั้นตอนที่ 4 ใช้ชุดข้อมูลตรวจสอบ (Validation Data) หาความถี่ของรูปแบบตัวอักษร ตามขั้นตอนของ SAX และ BOSS

ขั้นตอนที่ 5 คำนวณค่าความเหมือนโคไซน์ (Cosine Similarity) ระหว่าง tf-idf จากขั้นตอนที่ 3) กับความถี่ของตัวอักษรในชุดข้อมูลตรวจสอบ (Validation Data) จากขั้นตอนที่ 4) หาค่าความเหมือนโคไซน์ (Cosine Similarity) เพื่อจำแนกประเภทข้อมูล (Class) ให้กับชุดข้อมูลนี้ ถ้าประเภทใดมีค่า

ความเหมือนโคไซน์ (Cosine Similarity) สูงที่สุด จะทำนายว่าเป็นประเภทนั้น (Predicted Class)

ขั้นตอนที่ 6 คำนวณค่าความแม่นยำ (Accuracy) ระหว่างค่าจริง (Actual Class) กับค่าทำนาย (Predicted Class) ของชุดข้อมูลตรวจสอบ (Validation Data)

ขั้นตอนที่ 7 หาค่าเฉลี่ยของค่าความแม่นยำทั้ง K รอบ สำหรับพารามิเตอร์ชุดนั้น ๆ และเลือกชุดของพารามิเตอร์ที่ดีที่สุดที่ให้ค่าเฉลี่ยค่าความแม่นยำสูงที่สุด เพื่อไปใช้ทดสอบประสิทธิภาพในข้อมูลชุดทดสอบ

7.3 การทดสอบประสิทธิภาพ

ใช้ชุดของพารามิเตอร์ที่ได้จากขั้นตอนที่ 7 ไปทดสอบกับชุดข้อมูลทดสอบ (Testing Data) และคำนวณค่าความแม่นยำ (Accuracy) ได้จาก Confusion Matrix สำหรับ SAX และ BOSS

ตารางที่ 5. แสดง Confusion Matrix

	ค่าทำนาย (Predicted Class)	
ค่าจริง (Actual Class)	Positive	Negative
Positive	True Positive (TP)	False Negative (FN)
Negative	False Positive (FP)	True Negative (TN)

โดยที่ True Positive (TP) คือ จำนวนข้อมูลที่ถูกทำนายว่า “จริง” และมีค่าเป็น “จริง”

True Negative (TN) คือ จำนวนข้อมูลที่ถูกทำนายว่า “ไม่จริง” และมีค่า “ไม่จริง”

False Positive (FP) คือ จำนวนข้อมูลที่ถูกทำนายว่า “จริง” แต่มีค่าเป็น “ไม่จริง”

False Negative (FN) คือ จำนวนข้อมูลที่ถูกทำนายว่า “ไม่จริง” แต่มีค่าเป็น “จริง”

ค่าความถูกต้อง (Accuracy) สามารถคำนวณได้ดังสมการที่ 15

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (15)$$

ค่าความแม่นยำ (Precision) สามารถคำนวณได้ดังสมการที่ 16

$$Precision = \frac{TP}{TP + FP} \quad (16)$$

ค่าความระลึก (Recall) สามารถคำนวณได้ดังสมการที่ 17

$$Recall = \frac{TP}{TP + FN} \quad (17)$$

คะแนน F1 (F1-Score) สามารถคำนวณได้ดังสมการที่ 18

$$F1-Score = \frac{2 \times (Precision \times Recall)}{Precision + Recall} \quad (18)$$

8. ผลการวิเคราะห์ข้อมูล

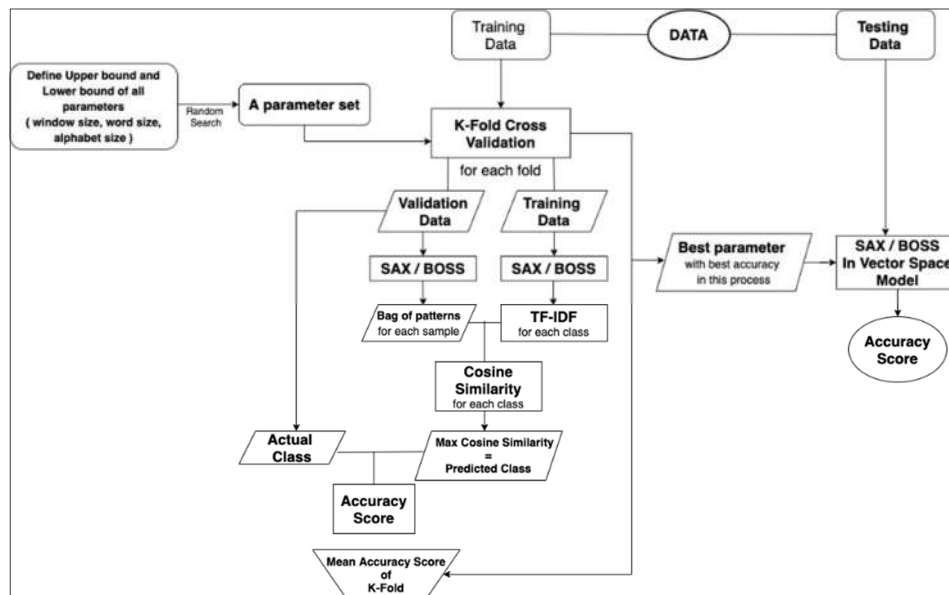
การเปรียบเทียบประสิทธิภาพการจำแนกประเภทข้อมูลอนุกรมเวลาระหว่างวิธี SAXVSM และ BOSSVS สำหรับข้อมูลคลื่นไฟฟ้าหัวใจที่มีสัญญาณรบกวนแบบต่าง ๆ และที่ระดับต่าง ๆ เริ่มต้นจากการจำลองการเพิ่มสัญญาณรบกวนในคลื่นไฟฟ้าหัวใจและแบ่งข้อมูลออกเป็นชุดเรียนรู้และชุดทดสอบสัดส่วน 70% และ 30% ตามลำดับ และใช้ข้อมูลชุดเรียนรู้เพื่อหาพารามิเตอร์ที่ดีที่สุดด้วยการสุ่มชุดพารามิเตอร์ทั้งหมด 30 ชุดแล้วใช้ 10-Folds Cross Validation เลือกชุดพารามิเตอร์ที่ให้ค่าความแม่นยำสูงที่สุดไปทดสอบกับข้อมูลชุดทดสอบ (Testing Data) สามารถสรุปได้ว่า ทั้ง SAX และ BOSS มีประสิทธิภาพใกล้เคียงกัน

8.1 การเปรียบเทียบค่าความถูกต้อง (Accuracy)

สำหรับข้อมูลคลื่นไฟฟ้าหัวใจที่ไม่มีสัญญาณรบกวน ทั้งสองตัวแบบให้ค่าความถูกต้อง (Accuracy) อยู่ที่ประมาณ 99% แต่เมื่อเพิ่มสัญญาณรบกวนเข้าไปทำให้ค่าความถูกต้อง (Accuracy) ของทั้งสองตัวแบบลดลงเพียงเล็กน้อยแต่ยังคงทำงานได้ดีใกล้เคียงกัน มีค่าความถูกต้อง (Accuracy) ประมาณ 97-99%

Noise Type	Level	SAXVSM	BOSSVS
None	0	0.99075	0.99075
baseline	25	0.98435	0.99075
	50	0.98791	0.99289
	100	0.99289	0.98009
composite	25	0.99004	0.99147
	50	0.99289	0.99004
	100	0.99218	0.98720
emg	25	0.99431	0.99004
	50	0.99147	0.99360
	100	0.98933	0.99218
powerline	25	0.98578	0.99147
	50	0.98933	0.98862
	100	0.97937	0.98933

รูปที่ 9. แสดง Heat Map เปรียบเทียบค่าถูกต้องระหว่าง SAXVSM และ BOSSVS



รูปที่ 8. แสดงกระบวนการทำงาน

8.2 การเปรียบเทียบคะแนน F1 (F1-Score)

สำหรับข้อมูลคลื่นไฟฟ้าหัวใจที่ไม่มีสัญญาณรบกวน ทั้งสองตัวแบบให้คะแนน F1 (F1 Score) อยู่ที่ประมาณ 99% แต่เมื่อเพิ่มสัญญาณรบกวนเข้าไปทำให้คะแนน F1 (F1 Score) ทั้งสองตัวแบบลดลงเพียงเล็กน้อยแต่ยังคงทำงานได้ดีใกล้เคียงกัน มีคะแนน F1 (F1 Score) ประมาณ 97-99%

Noise Type	Level	SAXVSM	BOSSVS
None	0	0.98751	0.98749
baseline	25	0.97921	0.98758
	50	0.98389	0.99046
	100	0.99044	0.97297
composite	25	0.98654	0.98846
	50	0.99040	0.98854
	100	0.98949	0.98273
emg	25	0.99228	0.98854
	50	0.98846	0.99132
	100	0.98562	0.98941
powerline	25	0.98942	0.98846
	50	0.98570	0.98470
	100	0.97262	0.98554

รูปที่ 10. แสดง Heat Map เปรียบเทียบคะแนน F1
ระหว่าง SAXVSM และ BOSSVS

8.4 การเปรียบเทียบค่าความระลึก (Recall)

สำหรับข้อมูลคลื่นไฟฟ้าหัวใจที่ไม่มีสัญญาณรบกวน ทั้งสองตัวแบบให้ค่าความระลึก (Recall) อยู่ที่ประมาณ 99% แต่เมื่อเพิ่มสัญญาณรบกวนเข้าไปทำให้ค่าความระลึก (Recall) ทั้งสองตัวแบบลดลงเพียงเล็กน้อยแต่ยังคงทำงานได้ดีใกล้เคียงกัน มีค่าความระลึก (Recall) ประมาณ 97-100%

Noise Type	Level	SAXVSM	BOSSVS
None	0	0.99037	0.98844
baseline	25	0.99037	0.99615
	50	1.00000	1.00000
	100	0.99807	0.97110
composite	25	0.99044	0.99037
	50	0.99472	0.98844
	100	0.99807	0.98651
emg	25	0.99037	0.98844
	50	0.99037	0.99037
	100	0.99037	0.99037
powerline	25	0.99037	0.99037
	50	0.99615	0.99229
	100	0.99229	0.98459

รูปที่ 12. แสดง Heat Map เปรียบเทียบค่าความระลึก
ระหว่าง SAXVSM และ BOSSVS

8.3 การเปรียบเทียบค่าความแม่นยำ (Precision)

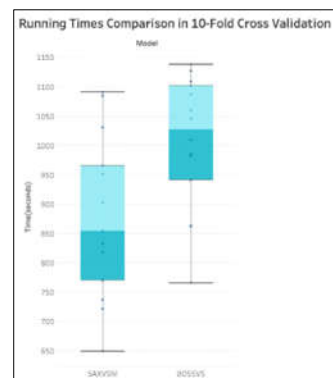
สำหรับข้อมูลคลื่นไฟฟ้าหัวใจที่ไม่มีสัญญาณรบกวน ทั้งสองตัวแบบให้ค่าความแม่นยำ (Precision) อยู่ที่ประมาณ 98% แต่เมื่อเพิ่มสัญญาณรบกวนเข้าไปทำให้ค่าความแม่นยำ (Precision) ทั้งสองตัวแบบลดลงเพียงเล็กน้อยแต่ยังคงทำงานได้ดีใกล้เคียงกัน มีค่าความแม่นยำ (Precision) ประมาณ 95-99%

Noise Type	Level	SAXVSM	BOSSVS
None	0	0.99467	0.98654
baseline	25	0.96104	0.97917
	50	0.96828	0.98110
	100	0.96292	0.97485
composite	25	0.98464	0.98656
	50	0.98652	0.98464
	100	0.98106	0.97897
emg	25	0.99420	0.98464
	50	0.98854	0.99228
	100	0.98092	0.98846
powerline	25	0.97164	0.98656
	50	0.97547	0.97723
	100	0.95370	0.98040

รูปที่ 11. แสดง Heat Map เปรียบเทียบค่าความแม่นยำ
ระหว่าง SAXVSM และ BOSSVS

8.5 การเปรียบเทียบเวลาในการสอนตัวแบบ

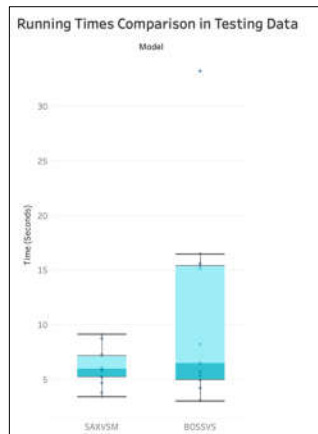
สำหรับการเปรียบเทียบเวลาที่ใช้นั้นขั้นตอนการสอนตัวแบบด้วย 10-Folds Cross Validation เมื่อทดสอบกับ SAXVSM และ BOSSVS เพื่อหาพารามิเตอร์ที่ดีที่สุดสำหรับข้อมูลแต่ละชุด จะเห็นว่า BOSSVS ใช้เวลาประมาณ 950-1,100 วินาที หรือ 15-18 นาที ต่อข้อมูล 1 ชุด ในขณะที่ SAXVSM ใช้เวลาประมาณ 650-950 วินาที หรือ 10-16 นาที ต่อข้อมูล 1 ชุด



รูปที่ 13. แสดง Boxplot เปรียบเทียบเวลาที่ใช้นั้นขั้นตอนการ
สอนตัวแบบ ระหว่าง SAXVSM และ BOSSVS

8.6 การเปรียบเทียบเวลาในการทดสอบตัวแบบ

สำหรับการเปรียบเทียบเวลาที่ใช้ในการทดสอบตัวแบบสำหรับ Testing Data ระหว่าง SAXVSM และ BOSSVS จะเห็นว่า BOSSVS ใช้เวลานานกว่า SAXVSM และมีการกระจายตัวมากกว่า อยู่ที่ประมาณ 5-15 วินาที ต่อข้อมูล 1 ชุด ส่วน SAXVSM ใช้เวลาอยู่ที่ประมาณ 5-8 วินาที ต่อข้อมูล 1 ชุด



รูปที่ 14. แสดง Box plot เปรียบเทียบเวลาการทดสอบตัวแบบระหว่าง SAXVSM และ BOSSVS

9. บทสรุปและการอภิปราย

การเปรียบเทียบประสิทธิภาพของการจำแนกประเภทข้อมูลอนุกรมเวลาระหว่าง SAXVSM และ BOSSVS โดยใช้ข้อมูลคลื่นไฟฟ้าหัวใจเป็นกรณีศึกษา และทำการเพิ่มสัญญาณรบกวนแบบ EMG Powerline Baseline และ Composite ที่ระดับ 25% 50% และ 100% สรุปได้ว่า โดยภาพรวมทั้ง 2 ตัวแบบมีประสิทธิภาพดีใกล้เคียงกัน ทั้งในแง่ของความถูกต้อง คะแนน F1 ค่าความแม่นยำ และค่าความระลึก โดยเฉลี่ยอยู่ที่ 98-99% สำหรับข้อมูลทั้ง 13 ชุด ทั้งข้อมูลคลื่นไฟฟ้าหัวใจที่ยังไม่ได้มีการเติมสัญญาณรบกวนหรือมีการเพิ่มสัญญาณรบกวนแบบต่างๆ อย่างไรก็ดี เมื่อเปรียบเทียบเวลาที่ใช้ในการประมวลผลเนื่องจาก SAXVSM มีกระบวนการที่ซับซ้อนน้อยกว่าจะใช้เวลาน้อยกว่า BOSSVS

สำหรับตัวแบบ SAXVSM เมื่อพิจารณาถึงค่าความถูกต้อง (Accuracy) คะแนน F1 (F1 Score) และค่าความแม่นยำ (Precision) ที่ทดสอบกับข้อมูลที่มีสัญญาณรบกวนแบบ Powerline ที่ระดับ 100% ทำให้ตัววัดประสิทธิภาพทั้ง 3 ตัวข้างต้นน้อยที่สุด ซึ่ง

สัญญาณรบกวนประเภทนี้เป็นคลื่นความถี่สูง เกิดจากการรบกวนของสนามแม่เหล็กไฟฟ้า ถูกจำลองมาจากฟังก์ชันคลื่นไซน์ (Sine Wave) ยิ่งที่ระดับสัญญาณรบกวนที่สูงขึ้น อาจมีส่วนทำให้ระยะขจัดที่แกว่งจากแนวเส้นฐาน (Amplitude) ผิดเพี้ยนไปจากรูปแบบของคลื่นไฟฟ้าหัวใจปกติ ทำให้ตัวแบบทำนายผิดพลาดได้ อย่างไรก็ดี เมื่อพิจารณาในส่วนของการเวลาในการประมวลผลทั้งขั้นตอนการสอนและการทดสอบตัวแบบ SAXVSM รวดเร็วกว่า BOSSVS อย่างเห็นได้ชัด

สำหรับตัวแบบ BOSSVS เมื่อพิจารณาถึงค่าความถูกต้อง (Accuracy) คะแนน F1 (F1 Score) ค่าความแม่นยำ (Precision) และค่าความระลึก (Recall) ที่ทดสอบกับข้อมูลที่มีสัญญาณรบกวนแบบ Baseline ที่ระดับ 100% ทำให้ตัววัดประสิทธิภาพทั้ง 4 ตัวข้างต้นน้อยที่สุด ซึ่งสัญญาณรบกวนประเภทนี้เกิดขึ้นได้จากสาเหตุทางกายภาพ เช่น การเคลื่อนไหวของร่างกาย ความชื้นหรือเหงื่อที่ผิวหนัง การเคลื่อนที่ของอิเล็กโทรดหรือขั้วไฟฟ้าที่ติดบนผิวหนังและเชื่อมเข้ากับเครื่องบันทึกคลื่นไฟฟ้าหัวใจ ถึงแม้จะมีความถี่ต่ำแต่ก็พบสัญญาณรบกวนประเภทนี้ได้ง่าย สัญญาณรบกวนประเภทนี้ถูกจำลองมาจากฟังก์ชันคลื่นไซน์ (Sine Wave) ยิ่งมีระดับสัญญาณรบกวนที่สูงขึ้น อาจมีส่วนทำให้ระยะขจัดที่แกว่งจากแนวเส้นฐาน (Amplitude) ผิดเพี้ยนไปจากรูปแบบของคลื่นไฟฟ้าหัวใจปกติ ทำให้ตัววัดประสิทธิภาพทั้ง 4 ตัวให้ผลออกมาน้อยที่สุด เมื่อเทียบกับสัญญาณรบกวนแบบอื่น ๆ เมื่อพิจารณาค่าความระลึกสำหรับสัญญาณรบกวนแบบ Baseline ที่ระดับ 50% พบว่าตัวแบบไม่มีการทำนายผิดพลาด แต่เมื่อระดับ 100% ค่าความระลึก ลดลงน้อยที่สุด เมื่อเทียบกับสัญญาณรบกวนแบบอื่น ๆ จึงเป็นข้อควรระวังหากตรวจจับสัญญาณรบกวนประเภทนี้ได้ ก็ควรใช้เทคนิคการกรองสัญญาณรบกวนประเภทนี้ออกเนื่องจากตัวแบบอาจจะยังทำงานได้ไม่ดีเท่าที่ควร

สำหรับเวลาที่ใช้ในการประมวลผลสำหรับตัวแบบ SAXVSM และ BOSSVS ในขั้นตอนการเลือกพารามิเตอร์โดยใช้ 10-Folds Cross Validation สำหรับข้อมูลในแต่ละชุด ยังมีการเพิ่มสัญญาณรบกวนในระดับที่สูงขึ้นทั้ง SAXVSM และ BOSSVS จะใช้เวลาในการสอนตัวแบบน้อยลง เนื่องจากในขั้นตอนการสร้างเมทริกซ์น้ำหนักค่าด้วยการหาความถี่ของค่าและการผกผันความถี่ในเอกสาร (TF-IDF) ยิ่งข้อมูลที่มีระดับสัญญาณรบกวนสูงขึ้น ทำให้ตัวแบบเจอรูปแบบของสัญญาณรบกวนมาบดบังรูปแบบที่แท้จริงของคลื่นไฟฟ้าหัวใจ ทำให้ได้

จำนวนรูปแบบตัวอักษรระหว่างประเภทที่ปกติและประเภทผิดปกติเข้ากันมากขึ้น และความหลากหลายของรูปแบบตัวอักษรน้อยลง จึงทำให้เมตริกซ์มีขนาดเล็กลง เมื่อไปคำนวณค่าความเหมือนโคไซน์ (Cosine Similarity) จึงทำได้รวดเร็วขึ้น

โดยสรุปจากการวิจัยนี้ ทั้งสองตัวแบบนี้ใช้เทคนิคการลดมิติของอนุกรมเวลาด้วยการแบ่งอนุกรมเวลาเป็นอนุกรมเวลาย่อย ๆ และพยายามแปลงให้เป็นลำดับของตัวอักษรที่สามารถอธิบายได้ด้วย Piecewise Aggregate Approximation (PAA) สำหรับ SAX และ Discrete Fourier Transform (DFT) สำหรับ BOSS จะได้ความสำคัญของรูปแบบตัวอักษรนั้น ๆ จึงต้องพิจารณาความถี่ของรูปแบบของอักษรนั้น แต่จะถูกลดทอนลงด้วยสัดส่วนของรูปแบบตัวอักษรที่อยู่ในประเภท (Class) อื่น ๆ ด้วย และหาค่าความเหมือนโคไซน์ เพื่อใช้ในการจำแนกประเภทข้อมูลต่อไป ผลการวิจัยพบว่าทั้ง 2 ตัวแบบเหมาะสำหรับการจำแนกประเภทข้อมูลอนุกรมเวลา และมีประสิทธิภาพดีใกล้เคียงกัน แต่ SAXVSM ใช้เวลาการประมวลผลน้อยกว่า

10. ข้อเสนอแนะ

งานวิจัยนี้ใช้ข้อมูลคลื่นไฟฟ้าหัวใจจากฐานข้อมูล Physionet เพื่อเปรียบเทียบประสิทธิภาพของตัวแบบจำแนกประเภทข้อมูลอนุกรมเวลา ซึ่งให้ผลการเปรียบเทียบประสิทธิภาพของทั้ง 2 วิธียังไม่ชัดเจนเท่าที่ควร อาจเพราะปริมาณข้อมูลที่นำมาศึกษาน้อยเกินไป ในการวิจัยครั้งต่อไปควรเพิ่มปริมาณข้อมูลที่ใช้ศึกษาให้มากขึ้น รวมทั้งศึกษาข้อมูลด้านอื่น ๆ เพิ่มเติมด้วย อีกทั้งงานวิจัยนี้ทำการเปรียบเทียบประสิทธิภาพของตัวแบบจำแนกประเภทข้อมูลอนุกรมเวลาด้วยวิธีการ 2 วิธี ได้แก่ SAXVSM และ BOSSVS เท่านั้น ในการวิจัยครั้งต่อไปอาจจะเพิ่มการเปรียบเทียบการจำแนกประเภทข้อมูลอนุกรมเวลาด้วยเทคนิคอื่น ๆ เพิ่มเติม

เอกสารอ้างอิง

- [1] Thailand Online Hospital. "การตรวจคลื่นไฟฟ้าหัวใจ (Electrocardiography)," 2021. [Online]. Available: <https://bit.ly/3EiRCUa>. [Accessed April 15, 2021].
- [2] J. Lin, E. Keogh, L. Wei, and S. Lonardi, "Experiencing SAX: a novel symbolic representation of time series," *Data Mining and Knowledge Discovery*, vol. 15, no. 2, pp. 107-144, doi: 10.1007/s10618-007-0064-z, 2007.
- [3] P. Senin and S. Malinchik, "Sax-vsm: Interpretable time series classification using sax and vector space model," in *2013 IEEE 13th international conference on data mining*, IEEE, pp. 1175-1180, 2013.
- [4] P. Schäfer, "Bag-Of-SFA-Symbols in Vector Space (BOSS VS)," Zuse-Institut Berlin (ZIB), 2015.
- [5] P. Schäfer, "Scalable time series similarity search for data analytics," doi: 10.18452/17338, 2015.
- [6] P. Schäfer, "The BOSS is concerned with time series classification in the presence of noise," *Data Mining and Knowledge Discovery*, vol. 29, no. 6, pp. 1505-1530, doi: 10.1007/s10618-014-0377-7, 2014.
- [7] Y. Kaya and H. Pehlivan, "Classification of premature ventricular contraction in ECG," *Int J Adv Comput Sci Appl*, vol. 6, no. 7, pp. 34-40, doi: 10.14569/IJACSA.2015.060706, 2015.
- [8] K. M. Chang, "Arrhythmia ECG noise reduction by ensemble empirical mode decomposition," *Sensors (Basel)*, vol. 10, no. 6, pp. 6063-80, doi: 10.3390/s100606063, 2010.
- [9] P. Senin. "Z-normalization of time series," 2021. [Online]. Available: https://jmotif.github.io/saxvsm_site/morea/algorithm/znorm.html. [Accessed April 21, 2021].
- [10] Souspace. "รู้จัก Discrete Fourier Transform และ Fast Fourier Transform," 2021. [Online]. Available: <https://www.youtube.com/watch?v=a03-5mr5yn4>. [Accessed April 21, 2021].
- [11] V. V. Raghavan and S. M. Wong, "A critical analysis of vector space model for information retrieval," *Journal of the American Society for information Science*, vol. 37, no. 5, pp. 279-287, doi: 10.1002/(SICI)1097-4571(198609)37:5<279::AID-ASII>3.0.CO;2-Q, 1986.
- [12] N. Srikong. "Cosine similarity," 2021. [Online]. Available: https://medium.com/@srikong_n/cosine-similarity-f1f9a962ddc5 [Accessed April 21, 2021].



ปัจจัยที่มีผลต่อความตั้งใจใช้แอปพลิเคชันธุรกรรมทางการเงินในกลุ่มผู้สูงอายุในช่วง

สถานการณ์การแพร่ระบาดของเชื้อโควิด-19

Factors Affecting Elderly's Intention to use Financial Services Application During the COVID-19 Pandemic

ภักจิรา สุคัสติย์* และ ปราโมทย์ ลื่อนาม

Pakjira Sukassathit and Pramote Luenam*

สาขาบริหารเทคโนโลยีสารสนเทศ คณะสถิติประยุกต์

สถาบันบัณฑิตพัฒนบริหารศาสตร์

IT Management Program, Graduate School of Applied Statistics,

National Institute of Development Administration

Received: November 18, 2021; Revised: December 12, 2021; Accepted: December 23, 2021; Published: December 28, 2021

ABSTRACT – The purpose of this research was to determine the factors that affect to elderly's intention to use financial services application during the COVID-19 pandemic. In this regard, 400 questionnaires collected from elders, who have experience with financial services application, were analyzed. The descriptive statistics used in this research were employed to summarize the characteristics of the sample and included frequency and percentages. The data were analyzed through Structural Equation Model to test the hypotheses and derive model fit. The model showed a goodness-of-fit with $P\text{-value} = 0.087$, $\chi^2/df = 1.142$, $RMSEA = 0.019$, $CFI = 0.997$ and $TLI = 0.994$. According to analysis results, the relative advantage, the compatibility, and the Trialability have a positive influence on the perceived usefulness and also have an indirect affect to elderly's intention to use financial services application through the perceived usefulness. The results also show the complexity and the observability have a positive influence on the perceived ease of use and also have an indirect affect to elderly's intention to use financial services application through the perceived ease.

KEYWORDS: Financial Service Application, Elderly Adults, Technology Acceptance Model, Diffusion of innovation theory, COVID-19

บทคัดย่อ -- การวิจัยครั้งนี้มีวัตถุประสงค์เพื่อศึกษาอิทธิพลต่อความตั้งใจใช้แอปพลิเคชันธุรกรรมทางการเงินในกลุ่มผู้สูงอายุในช่วงสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 กลุ่มตัวอย่างเป็นผู้สูงอายุที่ได้มีการใช้งานแอปพลิเคชันธุรกรรมทางการเงินออนไลน์ 400 คน โดยใช้แบบสอบถามออนไลน์เป็นเครื่องมือในการเก็บรวบรวมข้อมูล สถิติเชิงพรรณนาได้นำมาใช้เพื่อสรุปลักษณะของประชากร ประกอบด้วยค่าความถี่ และ ค่าร้อยละ เมื่อนำข้อมูลมาวิเคราะห์ด้วยโมเดลสมการโครงสร้าง เพื่อทดสอบสมมติฐาน และหาความสอดคล้องกับข้อมูลเชิงประจักษ์ โมเดลแสดงความเหมาะสม

*Corresponding Author: pakjira.suk@stu.nida.ac.th

ของด้วยค่า $P\text{-value} = 0.087$, $\chi^2/df = 1.142$, $RMSEA = 0.019$, $CFI = 0.997$ และ $TLI = 0.994$ ผลการศึกษาพบว่า คุณลักษณะประโยชน์เชิงเปรียบเทียบ คุณลักษณะที่เข้าถึงได้ และคุณลักษณะสามารถทดลองใช้ได้ มีอิทธิพลเชิงบวกต่อการรับรู้ประโยชน์ที่เกิดจากการใช้ และมีอิทธิพลทางอ้อมต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินออนไลน์ของผู้สูงอายุผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ ผลการศึกษายังพบว่าคุณลักษณะความยุ่งยากซับซ้อน และคุณลักษณะสามารถสังเกตมีอิทธิพลเชิงบวกต่อการรับรู้ถึงความง่ายในการใช้งาน และมีอิทธิพลทางอ้อมต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินออนไลน์ของผู้สูงอายุผ่านการรับรู้ถึงความง่ายในการใช้งาน

คำสำคัญ: แอปพลิเคชันธุรกรรมทางการเงิน, ผู้สูงอายุ, ตัวแบบการยอมรับเทคโนโลยี, ทฤษฎีการแพร่กระจายนวัตกรรม, โควิด-19

1. บทนำ

ปัจจุบันการทำธุรกรรมทางการเงินเป็นสิ่งที่เข้ามามีบทบาท เกี่ยวข้องกับการดำเนินชีวิตประจำวันของทุกคน ไม่ว่าจะเป็น การโอนเงินจากบัญชีธนาคารหนึ่งไปอีกธนาคารหนึ่ง การจ่าย บิลค่าสินค้าและบริการ การชำระค่าใช้จ่ายบัตรเครดิต ชื้อ ขาย หรือสับเปลี่ยนหน่วยลงทุน เป็นต้น และด้วยปัจจุบันระบบ เครือข่ายอินเทอร์เน็ตและเทคโนโลยี ที่มีการพัฒนาให้ผู้ใช้เข้าถึง ได้ง่ายขึ้นด้วยอุปกรณ์พกพาเคลื่อนที่ ได้แก่ โทรศัพท์มือถือ แท็บเล็ต เป็นต้น บวกกับการที่ธนาคารสถาบันการเงินต่าง ๆ มีการสร้างแอปพลิเคชันธุรกรรมทางการเงิน (Financial Services Application) เพื่อทำธุรกรรมได้ตลอดเวลา 24 ชั่วโมง ในทุก สถานที่ ที่ไม่ต้องเดินทางไปดำเนินการที่สาขาเหมือนแต่ก่อน และ จากสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 ทำให้การ ออกไปทำธุรกรรมที่สาขาของธนาคารเป็นเรื่องที่ยากขึ้น จาก การที่มีผู้ใช้บริการที่สาขาเป็นจำนวนมาก ซึ่งอาจมีความเสี่ยงที่ เชื้อจะแพร่กระจายในที่ ๆ มีผู้คนอยู่เป็นจำนวนมาก แอปพลิเคชัน ธุรกรรมทางการเงินสามารถช่วยประหยัดเวลา ไม่ต้อง ออกไปทำธุรกรรมที่สาขาที่จะเพิ่มความเสี่ยงให้แก่ตนเอง อีกทั้ง ยังเพิ่มความสามารถของแอปพลิเคชัน ให้อำนวยความสะดวก ให้กับลูกค้ามากขึ้น คอบไลฟสไตล์การใช้ชีวิตของ ผู้ใช้งานที่หลากหลายในยุคปัจจุบัน เช่น การแนะนำร้านอาหาร การบริจาคเพื่อการกุศล ชื้อบัตรชมภาพยนตร์ การติดตามดูความ เคลื่อนไหวของเงินฝาก บัตรเครดิต กองทุน สินเชื่อ การลงทุน การจองซื้อหุ้น เป็นต้น แต่เนื่องด้วยผู้ใช้งานมีทุกเพศ ทุกวัย ซึ่ง ผู้สูงอายุเป็นวัยที่ชอบการใช้ในรูปแบบเดิม ไม่ชอบที่จะ ปรับเปลี่ยนการใช้ชีวิตหรือเรียนรู้เทคโนโลยีใหม่ ๆ ซึ่งผู้สูงอายุ ส่วนใหญ่มักคิดว่า ประโยชน์ที่ได้จากการใช้แอปพลิเคชัน

ธุรกรรมทางการเงิน อาจไม่ตรงกับความต้องการของตัวเอง และ ที่สำคัญคือความกลัว ทั้งในความยุ่งยากในการเรียนรู้ และใช้งาน ความเสี่ยงในด้านการรักษาความปลอดภัย และความกลัวที่จะ เกิดข้อผิดพลาดในการใช้เทคโนโลยี [1], [2], [3]

ในขณะที่ฐานลูกค้าในกลุ่มผู้สูงอายุมีแนวโน้มที่จะ กลายเป็นลูกค้ากลุ่มใหญ่ของธนาคารพาณิชย์ และสถาบัน การเงิน เนื่องจากโครงสร้างประชากรของประเทศ ที่มีการ เปลี่ยนแปลงไปสู่สังคมที่มีคนวัยทำงานลดลง และมีผู้สูงวัย จำนวนมากขึ้น ทำให้เกิดการเปลี่ยนแปลงทั้งในแง่ของสัดส่วน ฐานลูกค้า ความต้องการ และความคาดหวังในผลิตภัณฑ์ จึงเป็น ความท้าทายของสถาบันการเงินที่จะต้องการปรับเปลี่ยนแบบ การให้บริการ ตลอดจนการสร้างโมเดลธุรกิจใหม่ ๆ ให้สามารถ ตอบสนองความต้องการเฉพาะทางของกลุ่มลูกค้าผู้สูงอายุ การ ทำความเข้าใจในพฤติกรรมด้านธุรกรรมการเงินของผู้สูงอายุ และการศึกษาอุปสรรคที่ทำให้ผู้สูงอายุเกิดความกังวลในการทำ ธุรกรรมทางการเงินในรูปแบบออนไลน์ จึงเป็นสิ่งที่สำคัญ จำเป็น โดยเฉพาะในช่วงสถานการณ์การแพร่ระบาดของเชื้อ โควิด-19 ที่ทำให้สภาพแวดล้อมการใช้ชีวิตที่เปลี่ยนไป และอาจมี ผลทำให้พฤติกรรมของผู้สูงอายุต้องปรับเปลี่ยนตาม กล่าวคือ อาจจะมีคามสนใจในการทำธุรกรรมรูปแบบออนไลน์เพิ่มขึ้น

งานวิจัยนี้มีวัตถุประสงค์ของการศึกษา ดังนี้ (1) เพื่อศึกษา ปัจจัยด้านการเผยแพร่นวัตกรรมที่มีอิทธิพลทางตรงต่อการรับรู้ ประโยชน์ที่เกิดจากการใช้ และ การรับรู้ถึงความง่ายในการใช้ งาน (2) เพื่อศึกษาปัจจัยด้านการเผยแพร่นวัตกรรมที่มีอิทธิพล ทางอ้อมต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทาง การเงินออนไลน์ของผู้สูงอายุในช่วงสถานการณ์การแพร่ระบาด

ของเชื้อโควิด-19 ทั้งนี้เพื่อผลลัพธ์ที่ได้มากำหนดแนวทางส่งเสริมให้ผู้สูงอายุหันมาทำธุรกรรมทางออนไลน์มากขึ้น

2. แนวคิด

องค์การอนามัยโลก [4] กำหนดเกณฑ์อายุตามสภาพของการมีอายุเพิ่มขึ้น โดยแบ่งเป็น ช่วงที่ 1 ผู้สูงอายุ (Elderly) มีอายุระหว่าง 60-74 ปี ช่วงที่ 2 คนชรา (Old) มีอายุระหว่าง 75-90 ปี ช่วงที่ 3 คนชรามาก (Very Old) มีอายุ 90 ปีขึ้นไป สำหรับผู้สูงอายุในประเทศไทย ตามพระราชบัญญัติบำเหน็จบำนาญข้าราชการ พ.ศ. 2494 ได้กำหนดนิยามผู้มีอายุ 60 ปี ขึ้นไปให้เป็นผู้เกษียณจากการทำงาน ในความรู้สึกรู้สึกของคนที่ไปผู้สูงอายุมักถูกมองว่า อยู่ในภาวะที่ร่างกายเกิดการเปลี่ยนแปลงไปทางเสื่อมถอยมากขึ้น มีโรคประจำตัว อ่อนแอช่วยเหลือตัวเองไม่ได้ ไม่มีรายได้ เป็นภาระของภาครัฐที่ต้องจัดหาสวัสดิการ และญาติ ผู้ใกล้ชิดที่ต้องคอยดูแล ในเรื่องการใช้ชีวิตประจำวัน อย่างไรก็ตามในเรื่องการทำธุรกรรมด้านต่าง ๆ โดยเฉพาะด้านธุรกรรมด้านการเงิน เป็นสิ่งที่ผู้สูงอายุยังคงต้องดำเนินการอยู่ตลอดเวลาไม่สามารถหลีกเลี่ยงได้ แม้จะอยู่ในวัยที่พ้นจากทำงานประจำแล้วก็ตาม ในประเทศที่พัฒนาแล้ว เช่น กลุ่มสหภาพยุโรปจะมีผู้สูงอายุเป็นประชากรกลุ่มใหญ่ที่มีการใช้จ่ายเงินจำนวนมาก และมีส่วนสำคัญในการขับเคลื่อนเศรษฐกิจ โดยเฉพาะการใช้จ่ายในเรื่องค่าบริการทางการแพทย์ และการพักผ่อนหย่อนใจเนื่องจากปัญหาสุขภาพ และการที่มีเวลามากกว่าคนในวัยทำงาน [5] ทั้งนี้ผู้สูงอายุเหล่านี้จะต้องมีการปรับตัว ให้เข้ากับยุคปัจจุบันที่การทำธุรกรรมทางการเงินเน้นไปในทางออนไลน์มากขึ้น

ธนาคารอิเล็กทรอนิกส์ (e-Banking) เป็นระบบการให้บริการทางการเงินออนไลน์ (Online Financial Services) ของสถาบันการเงินผ่านช่องทางอิเล็กทรอนิกส์ อาศัยเทคโนโลยีและอินเทอร์เน็ตในการให้บริการลูกค้า [6] เช่น การฝากเงิน การถอนเงิน การโอนเงิน หรือการสอบถาม แจ้งยอดทางการเงิน ธนาคารอิเล็กทรอนิกส์ช่วยในเรื่องการทำธุรกรรมที่รวดเร็ว

สะดวกสบาย ใช้งานได้ในทุกสถานที่ ทุกเวลา และประหยัดทรัพยากร ผ่านช่องทางในการให้บริการหลักสองรูปแบบ ได้แก่ (1) การธนาคารอินเทอร์เน็ต (Internet Banking) หมายถึงบริการที่ลูกค้าสามารถทำธุรกรรมทางการเงิน ผ่านเว็บแอปพลิเคชันของธนาคารบนอุปกรณ์ใช้งาน เช่น คอมพิวเตอร์ส่วนบุคคล โน้ตบุ๊ก ที่มีการเชื่อมต่อกับอินเทอร์เน็ต (2) การธนาคารโมบาย (Mobile

Banking) หมายถึงบริการที่ลูกค้าสามารถทำธุรกรรมทางการเงินบนอุปกรณ์เคลื่อนที่ [7] เช่น สมาร์ทโฟน (Smartphone) แท็บเล็ต (Tablet) ผ่านโมบายแอปพลิเคชัน (Mobile Application) ของธนาคารที่ได้ติดตั้งบนโทรศัพท์มือถือ เชื่อมต่อระบบเครือข่ายอินเทอร์เน็ตไปสู่บริการการธนาคารเคลื่อนที่ ที่ธนาคารเปิดให้บริการผ่านระบบเครือข่ายของโทรศัพท์เคลื่อนที่ ได้แก่ระบบ GPRS (General Packet Radio Service), EDGE (Enhanced Data Rate for Global Evolution), 4G, 5G หรือผ่านระบบเครือข่ายอินเทอร์เน็ตแบบไร้สาย (Wireless LAN) เป็นต้น

จากการทบทวนวรรณกรรมเกี่ยวกับทฤษฎีการแพร่กระจายนวัตกรรม พบว่านิยามของ “นวัตกรรม” คือ การนำวิธีการใหม่ ๆ มาปฏิบัติหลังจากได้ผ่านการทดลองหรือ ได้รับการพัฒนาแล้ว โดยมีขั้นตอน ดังนี้ ขั้นตอนการคิดค้น (Innovation) ขั้นตอนการพัฒนา (Development) และขั้นนำไปปฏิบัติจริง ซึ่งทำให้เกิดความแตกต่างจากการปฏิบัติเดิม ๆ ที่เคยปฏิบัติมา [8]

โดยความหมายของ การแพร่กระจาย (Diffusion) คือ กระบวนการซึ่งของนวัตกรรมถูกสื่อสารผ่านช่องทางในช่วงเวลาหนึ่งระหว่างสมาชิกต่าง ๆ ที่อยู่ในระบบสังคม [9] ทฤษฎีการแพร่กระจายนวัตกรรม (Diffusion of Innovation Theory: DOI) เป็นทฤษฎีพื้นฐานทางสังคมวิทยา กล่าวคือ เป็นกระบวนการที่ นวัตกรรมได้มีการแพร่กระจายจากแหล่งประดิษฐ์ใหม่ ๆ โดยผ่านสื่อทางใดทางหนึ่งไปสู่สังคมให้บรรลุตามวัตถุประสงค์ ซึ่งอาจมีผลต่อการเปลี่ยนแปลงทางสังคม [9] จากทฤษฎีการแพร่กระจายนวัตกรรมของ Roger [10] ซึ่งคุณลักษณะที่มีผลต่อการยอมรับของธนาคารดิจิทัล มีทั้งหมด 5 ประการ ดังนี้ (1) คุณลักษณะประโยชน์เชิงเปรียบเทียบ (Relative Advantage: RA) ความหมายคือ เมื่อเกิดการรับรู้ว่านวัตกรรมดีกว่ามีประโยชน์กว่าวิธีการปฏิบัติเดิม ๆ เช่น สะดวกกว่า รวดเร็วกว่า มีผลตอบแทนที่ดีกว่าอื่น ๆ เป็นต้น หรือหากเห็นว่าได้รับประโยชน์มากกว่าเสียประโยชน์ ก็จะทำให้เกิดการยอมรับในนวัตกรรม หรือมีแนวโน้มในการยอมรับมากขึ้น (2) คุณลักษณะที่เข้ากันได้ (Compatibility: CPT) ความหมายคือ การที่ผู้รับนวัตกรรมรู้สึกหรือคิดว่าเข้ากันได้หรือไปด้วยกันได้กับค่านิยมที่เป็นอยู่เดิม ถ้านวัตกรรมใดมีลักษณะสอดคล้องกับความคิดเดิม ๆ ก็จะทำการยอมรับมีแนวโน้มสูงขึ้น จากประสบการณ์ในอดีตตลอดจนความต้องการของผู้รับความคิดใหม่ ๆ การเข้ากันได้ของนวัตกรรมกับสิ่งต่าง ๆ ทำให้ผู้ยอมรับรู้สึกมั่นใจและไม่ต้องเสี่ยงภัยมาก ทำให้เกิดความรู้สึกที่มี

ความหมายมากขึ้น (3) คุณลักษณะความยุ่งยากซับซ้อน (Complexity: CPC) ความหมายคือ หากนวัตกรรมมีความสลับซับซ้อน ยุ่งยากต่อการทำความเข้าใจและนำไปใช้งานในเชิงเปรียบเทียบ ก็จะทำให้เกิดการยอมรับน้อยลง ผู้ที่นำนวัตกรรมเหล่านั้นมาใช้เมื่อมีความยุ่งยากก็จะทำให้เกิดการต่อต้าน จะเห็นว่าปัจจัยความยุ่งยากซับซ้อนนี้มีความสัมพันธ์ผกผันกับกับอัตราการยอมรับ ถ้านวัตกรรมมีความซับซ้อนมาก อัตราการยอมรับจะลดลง ตรงกันข้ามหากนวัตกรรมมีความซับซ้อนน้อยอัตราการยอมรับก็จะเพิ่มขึ้น ดังนั้นในการออกแบบฮาร์ดแวร์หรือซอฟต์แวร์ จึงควรเน้นความเป็นมิตรกับผู้ใช้ (User-friendly) เพื่อให้ประสบความสำเร็จในการยอมรับการใช้งาน [11] (4) คุณลักษณะสามารถทดลองใช้ได้ (Trialability: TA) ความหมายคือ การนำเอาส่วนนวัตกรรมส่วนย่อย ๆ ไปทดลองใช้โดยใช้ระยะเวลาไม่มากนัก เมื่อประสบความสำเร็จตามที่ต้องการก็จะทำให้เกิดการยอมรับมากขึ้นในนวัตกรรมนั้น ๆ และหากเพิ่มการสร้างหรือนำเสนอในรูปแบบใหม่ (Reinvent) ก็จะสามารถทำให้เกิดการยอมรับนวัตกรรมได้เร็วมากขึ้น [12] (5) คุณลักษณะที่สังเกตเห็นได้ (Observability: OB) ความหมายคือ ระดับของการสังเกตเห็นผลที่เกิดขึ้นจากนวัตกรรม นั่นคือยิ่งทำให้สามารถมองเห็นผลจากนวัตกรรมจากบุคคลอื่นได้มากเท่าไรก็จะทำให้เกิดการยอมรับมากขึ้นเท่านั้น ดังนั้นการมีต้นแบบที่ดี (Role Model) การสังเกตได้จากคนรอบข้าง (Peer Observation) จึงเป็นปัจจัยจุดใจหลัก ที่สร้างให้เกิดการยอมรับและแพร่กระจายเทคโนโลยี

จากผลการวิเคราะห์ความสัมพันธ์ของปัจจัยดังกล่าว ผู้วิจัยมองว่าคุณลักษณะทั้ง 5 ประการ น่าจะส่งผลต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินของผู้สูงอายุ โดยที่คุณลักษณะประโยชน์เชิงเปรียบเทียบ คุณลักษณะที่เข้ากันได้ และ คุณลักษณะความยุ่งยากซับซ้อน มีอิทธิพลทางตรงต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ และนำไปสู่ความตั้งใจในการใช้งาน ส่วนคุณลักษณะสามารถทดลองใช้ได้ และ คุณลักษณะที่สังเกตเห็นได้ มีอิทธิพลทางตรงต่อการรับรู้ความง่ายในการใช้งาน และนำไปสู่ความตั้งใจในการใช้งาน

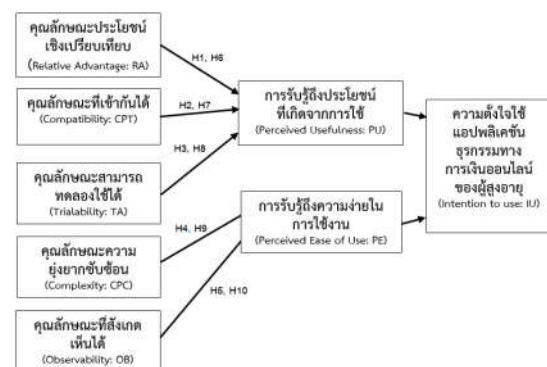
จากการศึกษาเกี่ยวกับที่มาของ ตัวแบบการยอมรับเทคโนโลยี (The Technology Acceptance Model: TAM) เป็นทฤษฎีที่คิดค้นโดย Davis, Bagozzi และ Warshaw [13] TAM เป็นโมเดลที่พัฒนามาจากทฤษฎีการกระทำอย่างมีเหตุผล (Theory of Reasoned Action: TRA) ที่เสนอโดย Fishbein และ

Ajzen [14] ตามแนวคิดที่ว่า มนุษย์โดยปกติแล้วเป็นผู้มีเหตุผล พฤติกรรมของปัจเจกบุคคลจึงไม่ได้เกิดขึ้น โดยขาดการพิจารณามาก่อน ดังนั้นการที่บุคคลจะมีหรือไม่มีพฤติกรรมอย่างใดอย่างหนึ่งนั้น จะเกิดจากความตั้งใจและมีเหตุผล

สำหรับแนวคิดของตัวแบบการยอมรับเทคโนโลยี หรือ TAM เกิดจากการนำทฤษฎีการกระทำอย่างมีเหตุผล มาผนวกกับแนวคิดพื้นฐานการยอมรับเทคโนโลยี ซึ่งเป็นวิทยานิพนธ์ปริญญาเอกของ Davis [15] เพื่อสร้างเป็นตัวแบบการยอมรับเทคโนโลยี ที่ใช้สำหรับอธิบายพฤติกรรมของผู้ใช้เทคโนโลยีสารสนเทศ โดยการประเมินระดับของการรับรู้ของผู้ใช้ที่มีต่อระบบ โดย TAM จะเน้นการศึกษาเกี่ยวกับ ปัจจัยต่าง ๆ ที่ส่งผลต่อการยอมรับหรือการตัดสินใจที่จะใช้เทคโนโลยีหรือ

นวัตกรรมใหม่ ปัจจัยหลักที่ส่งผลโดยตรงต่อการยอมรับเทคโนโลยีหรือนวัตกรรมของผู้ใช้ประกอบด้วย “การรับรู้ถึงความง่ายในการใช้งาน” (Perceived Ease of Use: PE) และ “การรับรู้ถึงประโยชน์ที่เกิดจากการใช้” (Perceived Usefulness: PU) ซึ่งผู้วิจัยเห็นว่า ปัจจัยหลักทั้งสองส่งผลต่อความตั้งใจในการใช้เทคโนโลยี (Intention to Use: IU) ของผู้สูงอายุ ดังนั้นหากผู้สูงอายุสามารถรับรู้ได้ว่าแอปพลิเคชันธุรกรรมทางการเงิน มีประโยชน์ต่อตัวเอง และการใช้งานไม่ได้ยุ่งยากเกินความสามารถ ก็น่าจะเกิดการยอมรับการใช้งานแอปพลิเคชันธุรกรรมทางการเงินที่มากขึ้น

จากการที่ผู้วิจัยได้ศึกษาแนวคิด และทฤษฎีที่เกี่ยวข้อง สามารถกำหนดกรอบแนวความคิดของการวิจัย ดังแสดงในรูปที่ 1



รูปที่ 1. กรอบแนวความคิดงานวิจัย

สมมติฐานที่ 1 (H1) : คุณลักษณะประโยชน์เชิงเปรียบเทียบ (RA) มีอิทธิพลต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU)

สมมติฐานที่ 2 (H2) : คุณลักษณะที่เข้าถึงได้ (CPT) มีอิทธิพลต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU)

สมมติฐานที่ 3 (H3) : คุณลักษณะสามารถทดลองใช้ได้ (TA) มีอิทธิพลต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU)

สมมติฐานที่ 4 (H4) : คุณลักษณะความยุ่งยากซับซ้อน (CPC) มีอิทธิพลต่อการรับรู้ถึงความง่ายในการใช้งาน (PE)

สมมติฐานที่ 5 (H5) : คุณลักษณะที่สังเกตเห็นได้ (OB) มีอิทธิพลต่อการรับรู้ถึงความง่ายในการใช้งาน (PE)

สมมติฐานที่ 6 (H6) : คุณลักษณะประโยชน์เชิงเปรียบเทียบ (RA) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU)

สมมติฐานที่ 7 (H7) : คุณลักษณะที่เข้าถึงได้ (CPT) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU)

สมมติฐานที่ 8 (H8) : คุณลักษณะสามารถทดลองใช้ได้ (TA) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU)

สมมติฐานที่ 9 (H9) : คุณลักษณะความยุ่งยากซับซ้อน (CPC) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ถึงความง่ายในการใช้งาน (PE)

สมมติฐานที่ 10 (H10) : คุณลักษณะที่สามารถสังเกตเห็นได้ (OB) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ถึงความง่ายในการใช้งาน (PE)

งานวิจัยนี้เป็นงานวิจัยเชิงปริมาณ (Quantitative Research) โดยใช้แบบสอบถามเป็นเครื่องมือเก็บรวบรวมข้อมูลจากกลุ่มตัวอย่าง ในช่วงระหว่างมิถุนายน-กรกฎาคม 2564 ประชากรคือกลุ่มผู้สูงอายุ ที่มีอายุ 60 ปีขึ้นไป และเคยใช้แอปพลิเคชันธุรกรรมทางการเงิน ผู้วิจัยใช้โปรแกรม G*Power ในการคำนวณขนาดของกลุ่มตัวอย่าง โดยกำหนดค่าเพาเวอร์ (1-β) เท่ากับ 0.95 ค่าอัลฟา (α) เท่ากับ 0.05 จำนวนตัวแปรทำนายเท่ากับ 8 ตัวแปรขนาดของอิทธิพล (Effect Size) เท่ากับ 0.033 ผลที่ได้คือขนาดของกลุ่มตัวอย่างเท่ากับ 396 ตัวอย่าง ซึ่งผู้วิจัยได้ปรับจำนวนขนาดตัวอย่างแบบสอบถามเพิ่มเติมเป็น 400 ตัวอย่าง

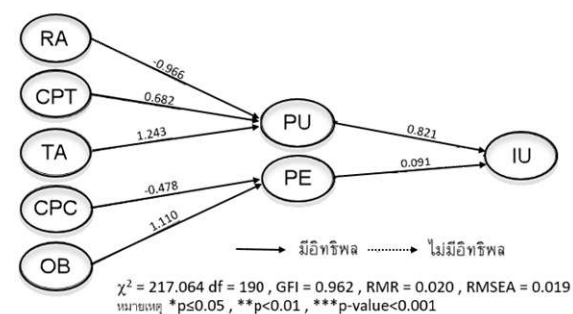
การวัดคุณภาพเครื่องมือแบบสอบถาม ทำโดยการวัดความเที่ยงตรง (Valid) โดยนำไปให้ผู้ทรงคุณวุฒิ 3 ท่าน แล้วนำมาพิจารณาค่าความสอดคล้องระหว่างข้อคำถามกับวัตถุประสงค์ หรือเนื้อหาด้วยเทคนิค IOC (Index of Item Objective Congruent) พบว่าทุกข้อมีค่าผ่านเกณฑ์มาตรฐานที่ยอมรับได้

แล้วจึงนำมาวัดค่าความเที่ยง (Stability) กับกลุ่มตัวอย่างที่มีลักษณะใกล้เคียงกับกลุ่มตัวอย่างที่ต้องการจะเก็บข้อมูลจำนวน 30 ชุด ได้ค่าสัมประสิทธิ์แอลฟา (Alpha Coefficient) ของครอนบาค (Cronbach) เท่ากับ 0.92 ถือได้ว่าแบบสอบถามที่ใช้มีความน่าเชื่อถือได้

สถิติที่ใช้สำหรับงานวิจัยนี้ประกอบด้วย (1) สถิติเชิงพรรณนา ได้แก่ ค่าความถี่ (Frequency) และ ค่าร้อยละ (Percentage) (2) สถิติเชิงอนุมาน ได้แก่ สถิติการวิเคราะห์โมเดลสมการโครงสร้าง (Structural Equation Modeling: SEM) เป็นพื้นฐานในการทดสอบสมมติฐาน และวิเคราะห์อิทธิพลของตัวแปรอิสระ ตัวแปรตาม และตัวแปรกึ่งกลาง (Mediator) ในการทดสอบสมมติฐานจะทำการทดสอบที่ระดับนัยสำคัญทางสถิติ 0.05 การวิเคราะห์โมเดลสมการโครงสร้างอาศัยโปรแกรมสำเร็จรูป AMOS ซึ่งเป็นโปรแกรมที่ให้ความแม่นยำสูงในการพิสูจน์ข้อสมมติฐานของงานวิจัย และอธิบายระดับความสัมพันธ์ของตัวแปร

3. ผลการทดลองและคำอธิบายรายละเอียด

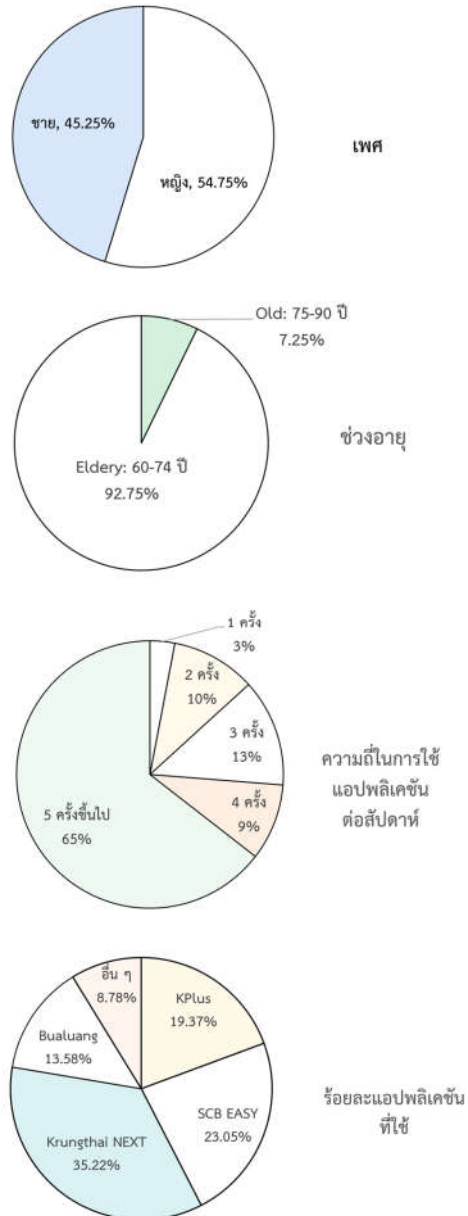
เมื่อนำข้อมูลมาวิเคราะห์ด้วยโมเดลสมการโครงสร้างเพื่อทดสอบสมมติฐาน และหาความสอดคล้องกับข้อมูลเชิงประจักษ์ โมเดลแสดงความเหมาะสมของด้วยค่า P-value = 0.087, $\chi^2/df = 1.142$, RMSEA = 0.019, CFI = 0.997 และ TLI = 0.994 ดังแสดงในรูปที่ 2 ส่วนผลการวิเคราะห์อิทธิพลเชิงสาเหตุภายในโมเดลแสดงในตารางที่ 1



รูปที่ 2. สัมประสิทธิ์เส้นทางมาตรฐานของโมเดล

ภาพรวมลักษณะสำคัญของตัวอย่างดังแสดงในรูปที่ 3 พบว่าตัวอย่างส่วนใหญ่เป็นเพศหญิงจำนวน 219 คน คิดเป็นร้อยละ 54.75 มีตัวอย่างที่อยู่ในช่วงผู้สูงอายุ (Elderly) ระหว่าง 60-74 ปี จำนวน 371 คน คิดเป็นร้อยละ 92.75 มีอายุอยู่ในช่วงคนชรา

(Old) ระหว่าง 75-90 ปี จำนวน 29 คนคิดเป็นร้อยละ 7.25 มีความถี่ในการใช้งานแอปพลิเคชัน 5 ครั้งต่อสัปดาห์คิดเป็นร้อยละ 65 และมีการใช้แอปพลิเคชัน Krungthai NEXT ของธนาคารกรุงไทย สูงที่สุดคิดเป็นร้อยละ 35.22 รองลงมาเป็นแอปพลิเคชัน SCB EASY ของธนาคารไทยพาณิชย์ และ KPlus ของธนาคารกสิกรไทย คิดเป็นร้อยละ 23.06 และ 19.38 ตามลำดับ



รูปที่ 3. ภาพรวมลักษณะสำคัญของตัวอย่าง

ตารางที่ 1. ผลการวิเคราะห์อิทธิพลเชิงสาเหตุภายในโมเดล

ตัวแปรตาม	อิทธิพล	ตัวแปรต้น						
		RA	CPT	TA	CPC	OB	PU	PE
PU	ทางตรง	-0.966 ***	0.682 *	1.243 ***	-	-	-	-
	ทางอ้อม	-	-	-	-	-	-	-
PE	ทางตรง	-	-	-	-0.478 *	1.110 ***	-	-
	ทางอ้อม	-	-	-	-	-	-	-
IU	ทางตรง	-	-	-	-	-	0.821 ***	0.091 *
	ทางอ้อม	-0.793 ***	0.560 ***	1.021 ***	-0.043 ***	0.101 ***	-	-

หมายเหตุ * $p \leq 0.05$, ** $p \leq 0.01$, *** $p \leq 0.001$

จากตารางที่ 1 สามารถตอบวัตถุประสงค์ของการวิจัย
จากสมมติฐานของการวิจัยได้ดังแสดงในตารางที่ 2

ตารางที่ 2. สรุปผลการทดสอบสมมติฐานในการวิจัย

สมมติฐานงานวิจัย	ค่าอิทธิพล	ค่า P	ผลการทดสอบ
H1: คุณลักษณะประโยชน์เชิงเปรียบเทียบ (RA) มีอิทธิพลต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU)	0.966	< 0.0001	ยอมรับ
H2: คุณลักษณะที่เข้าถึงได้ (CPT) มีอิทธิพลต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU)	0.682	0.043	ยอมรับ
H3: คุณลักษณะสามารถทดลองใช้ได้ (TA) มีอิทธิพลต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU)	1.243	< 0.0001	ยอมรับ
H4: คุณลักษณะความยุ่งยากซับซ้อน (CPC) มีอิทธิพลต่อการรับรู้ถึงความง่ายในการใช้งาน (PE)	-0.478	0.039	ยอมรับ
H5: คุณลักษณะที่สังเกตเห็นได้ (OB) มีอิทธิพลต่อการรับรู้ถึงความง่ายในการใช้งาน (PE)	1.110	< 0.0001	ยอมรับ
H6: คุณลักษณะประโยชน์เชิงเปรียบเทียบ (RA) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU)	-0.793	< 0.0001	ยอมรับ
H7: คุณลักษณะที่เข้าถึงได้ (CPT) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU)	0.560	< 0.0001	ยอมรับ
H8: คุณลักษณะสามารถทดลองใช้ได้ (TA) มีอิทธิพลกับความตั้งใจในการใช้	1.021	< 0.0001	ยอมรับ

สมมติฐานงานวิจัย	ค่า อิทธิพล	ค่า P	ผลการ ทดสอบ
เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU)			
H9: คุณลักษณะความพึงพอใจซับซ้อน (CPC) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ถึงความง่ายในการใช้งาน (PE)	-0.043	< 0.0001	ยอมรับ
H10: คุณลักษณะที่สามารถสังเกตได้ (OB) มีอิทธิพลกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ถึงความง่ายในการใช้งาน (PE)	0.101	< 0.0001	ยอมรับ

4. สรุปและการอภิปราย

จากการศึกษาพฤติกรรมกรรมการทำธุรกรรมทางการเงินของผู้สูงอายุในช่วงสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 เนื่องจากสภาพแวดล้อมอาจทำให้ผู้สูงอายุมีแนวโน้มการทำธุรกรรมทางการเงินที่เปลี่ยนไปจากเดิมและอาจมีความสนใจในการทำธุรกรรมทางการเงินออนไลน์มากขึ้น ซึ่งในการศึกษาครั้งนี้ประกอบไปด้วยสมมติฐานการวิจัย 10 ข้อ ดังแสดงในตารางที่ 2

จากตารางผลการทดสอบสมมติฐาน ตามกรอบแนวคิดงานวิจัย พบว่าคุณลักษณะประโยชน์เชิงเปรียบเทียบ (RA) มีอิทธิพลทางตรงต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU) ซึ่งมีความสอดคล้องกับแนวคิดของ Roger [10] ทั้งนี้น่าจะเป็นเพราะ ผู้สูงอายุเห็นว่าเมื่อเปรียบเทียบกับกรรมการทำธุรกรรมในรูปแบบดั้งเดิม การใช้แอปพลิเคชันธุรกรรมทางการเงินสามารถช่วยประหยัดเวลา ไม่มีค่าใช้จ่ายในการเดินทาง สามารถทำธุรกรรมได้ตลอดเวลา 24 ชั่วโมง ในทุกสถานที่ นอกจากนั้นในช่วงการเกิดสถานการณ์แพร่ระบาดของเชื้อโควิด-19 การเดินทางไปทำธุรกรรมที่สาขาจากการที่มีผู้ใช้บริการที่สาขาเป็นจำนวนมาก อาจมีความเสี่ยงที่จะติดเชื้อที่แพร่กระจายในระหว่างการทำธุรกรรมได้ ทั้งนี้สถาบันการเงินหรือ ผู้พัฒนาแอปพลิเคชันธุรกรรมทางการเงิน นอกจากควรเพิ่มความสามารถของระบบให้สูงขึ้นแล้ว ยังควรเพิ่มการประชาสัมพันธ์ ในลักษณะเปรียบเทียบประโยชน์ผ่านสื่อรูปแบบต่าง ๆ ที่เข้าถึงกลุ่มเป้าหมาย เพื่อสร้างการรับรู้ถึงประโยชน์ที่มากกว่าเดิมให้กับผู้สูงอายุ โดยอาจเปรียบเทียบให้เห็นถึงลำบากในการทำธุรกรรมแบบดั้งเดิม กับความสะดวกในการทำธุรกรรมทางการเงินของผู้สูงอายุผ่านแอปพลิเคชัน เช่น การถอนเงินแบบไม่ใช้บัตรที่ตู้เอทีเอ็ม ซึ่งแม้แต่การสอดบัตรเข้าตู้

เอทีเอ็ม ก็อาจจะเป็นเรื่องยากลำบากสำหรับผู้สูงอายุที่มีกล้ามเนื้อที่อ่อนแรง [16]

สำหรับปัจจัยคุณลักษณะที่เข้าถึงได้ (CPT) พบว่ามีอิทธิพลทางตรงต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU) ซึ่งมีความสอดคล้องกับแนวคิดของ Roger [10] จากผลการวิจัยจะเห็นได้ว่า หากผู้สูงอายุสามารถรู้สึกได้ว่าการใช้แอปพลิเคชันธุรกรรมทางการเงินเข้ากันได้ หรือไปด้วยกันได้กับการใช้ชีวิตประจำวันที่เป็นอยู่เดิม ไม่ต่างไปจากการไปทำธุรกรรมทางการเงินด้วยตนเองที่สาขา โดยเฉพาะในสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 เช่นนี้ ก็จะมีส่วนช่วยให้ผู้สูงอายุรับรู้ถึงประโยชน์ของการใช้ที่มากขึ้น สถาบันการเงินควรมีนโยบายสร้างความเข้ากันได้ของแอปพลิเคชันธุรกรรมทางการเงินกับรูปแบบการใช้ชีวิตของผู้สูงอายุ โดยเฉพาะในการใช้งานของผู้สูงอายุที่มีปัญหาด้านสภาพร่างกาย เช่น ความสามารถในการมองเห็น การใช้มือในการพิมพ์ หรือกดปุ่มรับสนับอุปกรณ์โมบายที่มีขนาดเล็ก ก็นับเป็นอุปสรรคที่สำคัญในการใช้งาน ดังนั้นจึงควรออกแบบให้แอปพลิเคชันสามารถอำนวยความสะดวกในการเข้าใช้งาน โดยอาจประยุกต์ใช้เทคโนโลยีไบโอเมตริก เช่น การสแกนรูม่านตา สแกนลายนิ้วมือ หรือเทคโนโลยีการรู้จำใบหน้า รู้จำเสียง ก็น่าจะเป็นการช่วยลดปัญหาในการใช้รหัสผ่าน ทั้งนี้จะช่วยทำให้ผู้สูงอายุรู้สึกมั่นใจ เป็นการยกระดับการให้บริการที่มอบความสะดวกสบาย ความปลอดภัยให้กับผู้ใช้ และทำให้เกิดความรู้สึกว่าแอปพลิเคชันมีประโยชน์มากขึ้น

คุณลักษณะสามารถทดลองใช้ได้ (TA) พบว่ามีอิทธิพลทางตรงต่อการรับรู้ถึงประโยชน์ที่เกิดจากการใช้ (PU) ซึ่งสอดคล้องกับแนวคิดของ Roger [10] จากผลการวิจัยจะเห็นว่า หากผู้สูงอายุได้มีโอกาสการทดลองใช้แอปพลิเคชัน โดยเฉพาะในช่วงสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 นี้ ซึ่งโดยส่วนใหญ่การใช้งานจะใช้เวลาเพียงไม่นาน ก็จะทำให้ผู้สูงอายุเข้าใจและเห็นประโยชน์ในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินมากขึ้น ธนาคารมีนโยบายในการเปิดโอกาสให้ผู้สูงอายุ ได้ทดลองใช้แอปพลิเคชันธุรกรรมทางการเงินในระยะเวลาสั้น ๆ เช่น การใช้นโยบายธนาคารที่เคลื่อนที่ได้ (Mobile Branch) [16] โดยให้พนักงานของธนาคารออกไปบริการลูกค้าตามชุมชน แทนการรอลูกค้าที่ธนาคาร และให้พนักงานให้คำแนะนำการใช้งานกับลูกค้าสูงอายุ อาจจะทำให้ผู้ใช้กับบางโมดูลที่จำเป็นและไม่ยากนัก เมื่อประสบ

ความสำเร็จตามที่ต้องการ ก็จะทำให้เกิดการรับรู้ถึงประโยชน์ที่เกิดจากการใช้มากขึ้น

คุณลักษณะความยุ่งยากซับซ้อน (CPC) พบว่ามีอิทธิพลทางตรงต่อการรับรู้ถึงความง่ายในการใช้งาน (PE) ซึ่งมีความสอดคล้องกับแนวคิดของ Roger [10] จากผลการวิจัยจะเห็นว่า หากผู้สูงอายุได้ใช้งานและรับรู้ถึงว่าการใช้งานที่ง่าย ไม่มีความยุ่งยากซับซ้อนและสามารถเรียนรู้การใช้งานได้อย่างง่าย จึงทำให้ผู้สูงอายุหันมาใช้งานแอปพลิเคชันธุรกรรมทางการเงินออนไลน์ในสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 มากขึ้น ดังนั้นผู้พัฒนาแอปพลิเคชันธุรกรรมทางการเงินควรให้ความสำคัญในการออกแบบซอฟต์แวร์ เน้นความเป็นมิตรกับผู้ใช้ (User-friendly) เพื่อให้ประสบความสำเร็จในการยอมรับการใช้งาน Thasuwat [17] แนะนำว่าการเน้นสื่อบนหน้าจอ หรือข้อความที่แสดงกันจะช่วยทำให้ผู้สูงอายุเห็นได้อย่างชัดเจน หน้าจอควรลดความสว่างลงเพื่อป้องกันความเมื่อยล้าของสายตา เนื่องจากผู้สูงอายุมีข้อจำกัดในการจำ จึงควรมีระบบช่วยเหลือให้คำอธิบายส่วนต่าง ๆ และให้การให้คำแนะนำแบบโต้ตอบเพื่อช่วยเหลือระหว่างการใช้งาน หรือเมื่อเกิดปัญหาในการใช้งาน

คุณลักษณะที่สังเกตเห็นได้ (OB) พบว่ามีอิทธิพลทางตรงต่อการรับรู้ถึงความง่ายในการใช้งาน (PE) ซึ่งมีความสอดคล้องกับแนวคิดของ Roger [10] จากงานวิจัยจะเห็นว่า หากผู้สูงอายุสามารถเห็นผลลัพธ์จากการให้บริการแอปพลิเคชันธุรกรรมทางการเงินของบุคคลอื่น ๆ ในช่วงการแพร่ระบาดของเชื้อโควิด-19 ก็จะทำให้ผู้สูงอายุมีความเข้าใจรับรู้ถึงความง่ายในการใช้งานมากขึ้น ธนาคารหรือสถาบันการเงินอาจใช้กลยุทธ์การโฆษณาประชาสัมพันธ์เฉพาะกลุ่ม การสร้างต้นแบบที่ดี (Role Model) ที่ใช้ผู้สูงอายุเป็นผู้นำเสนอ และการออกแบบโปสเตอร์ต้นแบบที่ดีเพื่อสร้างแรงจูงใจให้ผู้สูงอายุ เพื่อนำไปสู่การรับรู้ถึงความง่ายในการใช้งาน

คุณลักษณะประโยชน์เชิงเปรียบเทียบ (RA) พบว่ามีอิทธิพลทางอ้อมกับความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU) จากงานวิจัยจะเห็นว่า หากผู้สูงอายุรับรู้การเปรียบเทียบ ว่าการใช้งานแอปพลิเคชันธุรกรรมทางการเงินมีข้อดีว่าการไปทำธุรกรรมทางการเงินที่สาขาอย่างไร โดยเฉพาะในสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 นั้น ก็จะมีอิทธิพลทางอ้อมต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงิน โดยผ่านการรับรู้ถึงประโยชน์

ของการใช้งานแอปพลิเคชันธุรกรรมทางการเงินออนไลน์ ซึ่งสอดคล้องกับงานวิจัยของ Davis, Bagozzi และ Warshaw [13] ว่า หากผู้สูงอายุรับรู้ว่าการใช้งานแอปพลิเคชันธุรกรรมทางการเงินนั้นดีกว่าการไปทำธุรกรรมที่สาขาอย่างไร ผู้สูงอายุก็จะมีความตั้งใจในการใช้งานเช่นกัน

คุณลักษณะที่เข้าถึงได้ (CPT) พบว่ามีอิทธิพลกับทางอ้อมต่อความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU) จากผลการวิจัยจะเห็นว่า หากผู้สูงอายุคิดว่าการใช้งานแอปพลิเคชันธุรกรรมทางการเงินมีความสอดคล้องกับการใช้งานในชีวิตประจำวัน ภายใต้สถานการณ์การแพร่ระบาดของเชื้อโควิด-19 ก็จะมีอิทธิพลทางอ้อมให้ผู้สูงอายุเกิดความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินมากขึ้น ซึ่งได้สอดคล้องกับงานวิจัย Davis, Bagozzi และ Warshaw [13] ว่าหากผู้สูงอายุได้รับรู้ว่าการใช้งานแอปพลิเคชันธุรกรรมทางการเงิน สามารถทำธุรกรรมทางการเงินได้เหมือนกับการไปทำธุรกรรมทางการเงินที่สาขา และไม่ทำให้เกิดความเสี่ยงในการติดเชื้อโควิด-19 ผู้สูงอายุก็จะมีความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินมากขึ้น

คุณลักษณะสามารถทดลองใช้ได้ (TA) พบว่ามีอิทธิพลทางอ้อมต่อความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU) จากผลการวิจัยจะเห็นว่าหากผู้สูงอายุได้มีการทดลองใช้งานแอปพลิเคชันธุรกรรมทางการเงินในช่วงการแพร่ระบาดของสถานการณ์เชื้อโควิด-19 ซึ่งใช้ระยะเวลาทดลองเพียงไม่นาน แต่หากประสบความสำเร็จตามที่ต้องการ ก็จะส่งผลต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงิน ผลจากการศึกษาสอดคล้องกับงานวิจัยของ Davis, Bagozzi และ Warshaw [13] หากผู้สูงอายุได้ทดลองใช้แอปพลิเคชันธุรกรรมทางการเงิน และรู้สึกถึงประโยชน์ก็จะมีความตั้งใจในการใช้งาน และมีการยอมรับมากขึ้น

คุณลักษณะความยุ่งยากซับซ้อน (CPC) พบว่ามีอิทธิพลทางอ้อมต่อความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ถึงความง่ายในการใช้งาน (PE) จากผลการวิจัยจะเห็นว่าในสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 จึงทำให้ผู้สูงอายุมีการใช้งานแอปพลิเคชันธุรกรรมทางการเงินเพิ่มขึ้น และการจากการทดลองใช้นั้น ผู้สูงอายุได้เห็นว่าการใช้งานแอปพลิเคชันธุรกรรมทางการเงินไม่ใช่เรื่องที่ยาก จึงส่งผลต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินมากขึ้น จากงานวิจัยของ Davis, Bagozzi และ Warshaw [13] หากผู้สูงอายุรับรู้ว่าการ

ใช้งานแอปพลิเคชันไม่ใช่เรื่องยุ่งยากอีกต่อไป และในสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 เช่นนี้ การใช้งานแอปพลิเคชันธุรกรรมทางการเงิน นับว่าช่วยลดการแพร่ระบาดของเชื้อโควิด-19 และลดความเสี่ยงในการไปทำธุรกรรมที่สาขา ผู้สูงอายุก็จะมีการยอมรับการใช้งานและตั้งใจใช้งานแอปพลิเคชันธุรกรรมทางการเงินมากขึ้น

คุณลักษณะที่สามารถสังเกตได้ (OB) มีอิทธิพลกับทางอ้อมต่อความตั้งใจในการใช้เทคโนโลยี (IU) ผ่านการรับรู้ถึงความง่ายในการใช้งาน (PE) จากผลการวิจัยจะเห็นว่าหากผู้สูงอายุสามารถเห็นผลสำเร็จจากการใช้บริการแอปพลิเคชันธุรกรรมทางการเงินของบุคคลอื่น ๆ โดยเฉพาะจากต้นแบบที่ดีก็จะส่งผลต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงิน ซึ่งสอดคล้องกับ Davis, Bagozzi และ Warshaw [13] ผู้สูงอายุจะต้องสังเกตเห็นถึงการใช้งานแอปพลิเคชันการเงิน ของบุคคลรอบข้างก็จะส่งผลให้เกิดความตั้งใจในการใช้เทคโนโลยีผ่านการรับรู้ถึงความง่ายในการใช้งาน

เมื่อเปรียบเทียบระหว่างการรับรู้ประโยชน์ที่เกิดจากการใช้ (PU) กับการรับรู้ความง่ายในการใช้งาน (PE) ที่มีต่อความตั้งใจในการใช้เทคโนโลยี (IU) พบว่าทั้งสองปัจจัยมีอิทธิพลต่อความตั้งใจในการใช้เทคโนโลยีของผู้สูงอายุ แต่หากพิจารณาถึงระดับของอิทธิพล การรับรู้ประโยชน์ที่เกิดจากการใช้จะมีผลต่อต่อความตั้งใจในการใช้เทคโนโลยีสูงกว่าการรับรู้ความง่ายในการใช้งานมาก ทั้งนี้อาจจะเป็นเพราะในยุคปัจจุบันผู้สูงอายุมีการรับรู้ถึงความง่ายในการใช้งานแอปพลิเคชันยุคใหม่มากขึ้น ด้วยเหตุที่ว่าสถาบันการเงินต่าง ๆ มีโครงการ และมาตรการเชิงรุกเพื่อประชาสัมพันธ์ เผยแพร่ความรู้ ให้คำแนะนำวิธีการใช้ เพิ่มพูนทักษะการใช้งานแอปพลิเคชันการเงินให้กับผู้สูงอายุมากขึ้น เนื่องจากสถาบันการเงินเห็นถึงความสำคัญของผู้ใช้งานกลุ่มนี้ ที่กำลังกลายเป็นลูกค้ากลุ่มใหญ่ของธนาคารต่อไป นอกจากนั้นยังเป็นผลมาจากการที่อุปกรณ์สมาร์ตโฟนปัจจุบันได้มีการออกแบบ พัฒนารูปแบบการใช้งาน และปรับปรุงส่วนต่อประสาน รูปแบบการอินเทอร์เฟซให้ใช้งานให้รองรับการใช้งานของผู้มีปัญหาด้านสภาพร่างกาย ทั้งสายตา การได้ยิน และการสัมผัสได้ดีขึ้นกว่ายุคก่อนมาก

ดังนั้นจากวัตถุประสงค์ของการวิจัยที่ได้ตั้งไว้ว่า (1) เพื่อศึกษาปัจจัยด้านการเผยแพร่ข่าวสารที่มีอิทธิพลทางตรงต่อการรับรู้ประโยชน์ที่เกิดจากการใช้ และ การรับรู้ถึงความง่ายในการใช้งาน ผลจากการศึกษาพบว่าปัจจัยที่มีอิทธิพลทางตรงต่อ

การรับรู้ประโยชน์ที่เกิดจากการใช้ ประกอบด้วย คุณลักษณะประโยชน์เชิงเปรียบเทียบ คุณลักษณะที่เข้ากันได้ และคุณลักษณะสามารถทดลองใช้ได้ ส่วนปัจจัยที่มีอิทธิพลทางตรงต่อการรับรู้ถึงความง่ายในการใช้งาน ประกอบด้วย คุณลักษณะความยุ่งยากซับซ้อน และคุณลักษณะที่สังเกตเห็นได้ (2) เพื่อศึกษาปัจจัยด้านการเผยแพร่ข่าวสารที่มีอิทธิพลทางอ้อมต่อความตั้งใจในการใช้งานแอปพลิเคชันธุรกรรมทางการเงินออนไลน์ของผู้สูงอายุในช่วงสถานการณ์การแพร่ระบาดของเชื้อโควิด-19 พบว่า ปัจจัยที่มีอิทธิพลกับทางอ้อมต่อความตั้งใจในการใช้เทคโนโลยี ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ ประกอบด้วย คุณลักษณะประโยชน์เชิงเปรียบเทียบ คุณลักษณะที่เข้ากันได้ และคุณลักษณะสามารถทดลองใช้ได้ ส่วนปัจจัยที่มีอิทธิพลกับทางอ้อมต่อความตั้งใจในการใช้เทคโนโลยี ผ่านการรับรู้ประโยชน์ที่เกิดจากการใช้ ประกอบด้วย คุณลักษณะความยุ่งยากซับซ้อน และคุณลักษณะที่สังเกตเห็นได้

การวิจัยในครั้งนี้ เป็นการศึกษาพฤติกรรมของผู้สูงอายุกับความตั้งใจในการใช้แอปพลิเคชันธุรกรรมทางการเงินในช่วงระหว่างการแพร่ระบาดของเชื้อโควิด-19 อย่างไรก็ตามเป็นที่น่าสนใจว่า หากผ่านพ้นสถานการณ์การแพร่ระบาดนี้ไปแล้ว สิ่งต่าง ๆ ที่เกิดขึ้นเหล่านี้จะมีความยั่งยืนหรือไม่ หรือพฤติกรรมของผู้สูงอายุจะมีการเปลี่ยนในอนาคตในลักษณะใด มีแนวโน้มอย่างไร ข้อเสนอแนะในการวิจัยครั้งต่อไปนอกจากการศึกษาแนวโน้มของพฤติกรรมดังกล่าว การวิเคราะห์หาความแตกต่างของการยอมรับการใช้แอปพลิเคชันธุรกรรมทางการเงินในแต่ละช่วงวัย แยกตาม 3 กลุ่มหลักของผู้สูงอายุ ที่ประกอบด้วย ผู้สูงอายุ คนชรา และคนชรามาก รวมทั้งการทำความเข้าใจในผลกระทบของ ปัจจัยบุคคล (Personal Factors) เช่น ค่านิยมและรูปแบบการดำรงชีวิต ปัจจัยทางสังคม (Societal Factors) เช่น บทบาทและสถานะ (Roles and Status) ชั้นสังคม (Social Class) ล้วนเป็นสิ่งที่น่าสนใจศึกษาในกลไกและรูปแบบผลกระทบที่มีต่อพฤติกรรมของผู้สูงอายุ

เอกสารอ้างอิง

- [1] Mariano, J., Marques, S., Ramos, R. M., Gerardo, F., Lage da Cunha, C., Girenko, A., Alexandersson, J., Stree, B., Lamanna, M., Lorenzatto, M., Mikkelsen, P. L., Bundgård-Jørgensen, U., Rêgo, S., & Vries, H. "Too old for

- technology? Stereotype threat and technology use by older adults,” *Behaviour & Information Technology*, 2021.
- [2] Raymundo, T., & Santana, C. (2014). “Fear and the use of technological devices by older people,” *Gerontechnology*, Vol. 13 (2), 2014.
- [3] Giacomo, D. D., Ranieri, J., D’Amico, M., Guerra, F., & Passafiume, D., “Psychological barriers to digital living in older adults: computer anxiety as predictive mechanism for technophobia,” *Behavioral Sciences*, Vol. 9(9), pp. 96, 2019.
- [4] World Health Organization. “Men, ageing and health: achieving health across the life span,” *World Health Organization*, 2000. [Online]. Available: <https://apps.who.int/iris/handle/10665/66941>
- [5] Hirankasi, P., “Online banking services for the elderly in the digital era,” October 2020. [Online]. Available: <https://www.krungsri.com/th/research/research-intelligence/ri-elders-20>. [Accessed Sep. 20, 2021].
- [6] Maeroufi, F., Nouri, H., & Iman, Z. K., “Examination of the Role of Factors Influencing the Acceptance of E-banking,” *Journal of Applied Sciences*, Vol. 15(6), pp. 845-849, 2015.
- [7] Navavongsathian, A., Vongchavalitkul, B., & Limsarun, T., “Causal factors affecting mobile banking services acceptance by customers in Thailand,” *The Journal of Asian Finance, Economics and Business*, Vol. 7(11), pp. 421–428, November 2020.
- [8] Hughes, L., “Some Limits to Freedom,” *Philosophical Investigations*, Vol. 15(4), pp. 329-345, 1992.
- [9] Rogers, E. M., “*Diffusion of innovations*. (3rd ed.),” New York: The Free Press, 1995.
- [10] Rogers, E. M., “*Diffusion of Innovations*. (5th ed.),” New York: Free Press A Division of Simon & Schuster, Inc, 2003.
- [11] Martin, M. H., “Factors influencing faculty adoption of Web-based courses in teacher education programs within the State University of New York,” *Doctoral dissertation, Virginia Polytechnic Institute and State University*, 2003.
- [12] Sahin, I., “Detailed review of rogers’ diffusion of innovations theory and educational technology-related studies based on rogers’ theory,” *The Turkish Online Journal of Educational Technology*, Vol. 5(2), 2006.
- [13] Davis, F., Bagozzi, R. P., & Warshaw, P. R. “User acceptance of computer technology: A comparison of two theoretical models,” *Management Science*, Vol. 35(8), pp. 982–1003, 1989.
- [14] Fishbein, M., & Ajzen, I., *Belief, attitude, intention and behavior: An introduction to theory and research*. Reading, MA: Addison-Wesley, 1975.
- [15] Davis, F. “A technology acceptance model for empirically testing new end-user information systems: theory and results,” *Unpublished Doctoral dissertation, MIT Sloan School of Management*, 1985.
- [16] Choeykumhang, S. “What will a bank for the elderly look like in the future?”, November 2021. [Online]. Available: <https://www.scb.co.th/th/personal-banking/stories/elder-banking.html>. [Accessed Sep. 20, 2021].
- [17] Thasuwan, N., “Elderly Friendly User Interfaces in Sweden,” *Veridian e-Journal, Silpakorn University*, September-December, Vol. 11(3), 2018.



การสร้างตัวแบบจำแนกเอพิโทปของเซลล์มะเร็ง

Cancer Epitope Classification

มานนท์ บุญบางยาง* และ ศราวุธ นนท์ศิริ

Manon Boonbangyang* and Sarayut Nonsiri

ห้องปฏิบัติการวิจัยปัญญาประดิษฐ์และอินเทอร์เน็ตของสรรพสิ่ง

หลักสูตรเทคโนโลยีสารสนเทศ คณะเทคโนโลยีสารสนเทศ, สถาบันเทคโนโลยีไทย-ญี่ปุ่น

แขวงสวนหลวง เขตสวนหลวง กรุงเทพมหานคร ประเทศไทย 10250

Information of Technology, Faculty of Information of Technology, Thai-Nichi Institute of Technology,

Suan Luang, Bangkok Thailand 10250

Received: November 25, 2021; Revised: December 17, 2021; Accepted: December 23, 2021; Published: December 28, 2021

ABSTRACT – Cancer is a leading cause of death in the world. In 2020 World Health Organization (WHO) reported that approximately 10 million deaths caused by cancer and will increase for the coming years. This research paper aims to study the prediction of cancer epitope using machine learning for classifying between cancer cell surface and epitope on healthy cell surface. The comparison between the different machine learning algorithms is presented. This work can help to training T-cell for recognizing cancer cell and release enzyme to kill cancer cell (Targeted Therapy). The experiment results shown that imbalance data the model from Support Vector Machine (SVM) calculated based on Dipeptide Composition (DPC) feature achieved the best accuracy of 79 % Sensitivity 16% and Specificity 100% on test dataset. While balance data with SMOTE Random Forest (RF) calculated based on Dipeptide Composition (DPC) feature achieved the best accuracy of 80% Sensitivity 28% and Specificity 96% on the same test dataset. In conclusion, Support Vector Machine (SVM) and Random Forest (RF) calculated based on Dipeptide Composition (DPC) feature can employ these models for predicting the cancer epitope in imbalance dataset and balanced dataset.

KEYWORDS: Epitope, Machine Learning, Precision Medicine, Cancer vaccine, SMOTE

บทคัดย่อ -- มะเร็งเป็นสาเหตุการเสียชีวิตลำดับต้น ๆ ของโลก โดยองค์การอนามัยโลก (WHO)[1,2] รายงานในปี 2020 มีผู้เสียชีวิตจากมะเร็ง ประมาณ 10 ล้านรายและคาดการณ์ว่าจะเพิ่มมากขึ้นในปีถัดไปโดยผู้ป่วยส่วนใหญ่อยู่ในกลุ่มประเทศที่มีรายได้ปานกลางไปจนถึงรายได้ต่ำเมื่อตรวจพบมักเป็นระยะที่รุนแรง มะเร็งบางชนิดหากตรวจพบได้เร็วก็มีโอกาสรักษาให้หายขาดได้การรักษาในปัจจุบันทำได้โดยการฉายแสง การผ่าตัด การใช้เคมีบำบัด ไปจนถึงการรักษาด้วย วัคซีนที่มีความจำเพาะสูง [2]งานวิจัยชิ้นนี้ได้นำเสนอตัวแบบทำนายเอพิโทป (Epitope) ของแอนติเจน (Antigen) บนผิวเซลล์ของผู้ป่วยมะเร็งเทียบกับเอพิโทปบนผิวเซลล์ของผู้ที่มีสุขภาพดีข้อดีของการทำนายเอพิโทปเซลล์มะเร็งได้จะช่วยให้สามารถนำไปฝึกเซลล์ภูมิคุ้มกัน (T-cell) เพื่อเป็นวัคซีนรักษามะเร็งแบบจำเพาะเจาะจง (Precision Medicine) ในงานวิจัยนี้ทดลองใช้คุณลักษณะ (Feature) จำนวน 5 คุณลักษณะวิธีจำแนกข้อมูลแบบ binary classes จำนวน 7 ขั้นตอนวิธี ผลการทดสอบ

*Corresponding Author: bo.manon_st@tni.ac.th

ประสิทธิภาพพบว่า ข้อมูลที่ยังไม่ได้ปรับสมดุล ตัวแบบที่ได้จาก คุณลักษณะองค์ประกอบกระดุมโน้ต (DPC) ที่ใช้วิธี จำแนก ซัพพอร์ทเวกเตอร์แมชชีน (SVM) สามารถทำนายข้อมูลทดสอบมีค่าความแม่นยำสูงสุด 79% ค่าความไว 16% และค่าความจำเพาะ 100% ในข้อมูลทดสอบขณะที่ ข้อมูลที่ปรับสมดุลด้วย เทคนิค SMOTE สมดุล ตัวแบบที่ได้จาก คุณลักษณะองค์ประกอบกระดุมโน้ต (DPC) ที่ใช้วิธีจำแนกป่าสุ่ม (RF) มีค่าความแม่นยำสูงสุดที่ 80% ค่าความไว 28% และค่าความจำเพาะ 96% ในข้อมูลทดสอบ จากผลการทดสอบข้างต้นแสดงให้เห็นว่าคุณลักษณะองค์ประกอบกระดุมโน้ต (DPC) เมื่อใช้ร่วมกับ วิธีจำแนก ซัพพอร์ทเวกเตอร์แมชชีน(SVM) หรือ วิธีจำแนกป่าสุ่ม (RF) สามารถนำมาใช้ทำนายเอพิโทปของเซลล์มะเร็งได้ทั้งในข้อมูลที่ยังไม่ได้ปรับสมดุลและปรับสมดุลแล้วตามลำดับ

คำสำคัญ: เอพิโทป, การเรียนรู้ของเครื่อง, การรักษาแบบแม่นยำและจำเพาะ, วัคซีนมะเร็ง, SMOTE

1. บทนำ

มะเร็งเป็นสาเหตุการเสียชีวิตลำดับต้น ๆ ของโลกโดยองค์การอนามัยโลก (WHO) รายงาน ในปี 2020 มีผู้เสียชีวิตจากมะเร็งประมาณ 10 ล้านรายคาดการณ์ว่าจะเพิ่มมากขึ้นในปีถัดไปโดยผู้ป่วยส่วนใหญ่อยู่ในกลุ่มประเทศที่มีรายได้ปานกลางไปจนถึงรายได้ต่ำ [1]

มะเร็งมักจะถูกตรวจพบเมื่อเป็นระยะสุดท้ายแต่เมื่อเริ่มหลายชนิดสามารถรักษาให้หายได้หากตรวจพบในระยะเริ่มแรกการรักษาในปัจจุบันสามารถทำได้โดยการฉายแสง การผ่าตัด การใช้เคมีบำบัด ไปจนถึงการรักษาด้วยวัคซีนที่มีความจำเพาะสูง ประสิทธิภาพการรักษามักขึ้นอยู่กับร่างกายสามารถกำจัดเซลล์มะเร็งบางส่วนได้แต่เซลล์มะเร็งเกิดจากการกลายพันธุ์ของเซลล์ปกติ ซึ่งมีแอนติเจนบนผิวเซลล์ส่วนใหญ่คล้ายเซลล์ปกติทำให้เซลล์ภูมิคุ้มกันเข้าใจผิดว่าเป็นเซลล์ปกติจึงไม่ปล่อยสารเพื่อกำจัดเซลล์มะเร็งได้

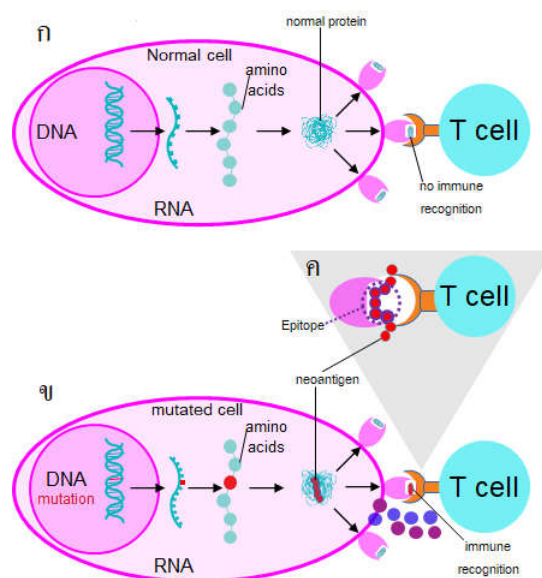
หากสามารถทำนายได้ว่าเอพิโทปบนแอนติเจนของเซลล์มะเร็งมีลำดับ กรดอะมิโนอย่างไรจะเป็นประโยชน์ในการพัฒนาและกระตุ้นเซลล์ภูมิคุ้มกันให้สามารถตรวจจับและกำจัดเซลล์มะเร็งได้อย่างแม่นยำและมีประสิทธิภาพ

งานวิจัยนี้มุ่งเน้นหาคุณลักษณะและตัวแบบที่เหมาะสมสำหรับทำนาย เอพิโทปเพื่อนำมาใช้สำหรับพัฒนาวัคซีนรักษามะเร็ง

2. ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

แอนติเจน หมายถึง สิ่งแปลกปลอมหรือสารที่ถูกสร้างโดยร่างกายหรือรับมาจากภายนอก ร่างกาย ประกอบขึ้นจากกรดอะมิโน

ในหลายชนิดเป็นสายมีทั้งแบบสั้นที่เรียกว่า เปปไทด์จนถึงสายยาวเป็นโปรตีนซึ่งมีคุณสมบัติหลากหลายขึ้นอยู่กับชนิดของกรดอะมิโนดังแสดงในตารางที่ 1 โดยปกติทีเซลล์ทำหน้าที่เป็นภูมิคุ้มกันจะไม่มีปฏิกิริยากับแอนติเจนที่ร่างกายผลิต แอนติเจนบนผิวของเซลล์ปกติที่บนผิวของเซลล์มะเร็งจะพบ แอนติเจน ที่เกิดการกลายพันธุ์ที่เรียกว่านี้โอแอนติเจนหากทีเซลล์ตรวจพบจะทำปฏิกิริยากับนี้โอแอนติเจนและปล่อยอินเตอร์เฟอรอนแกมมาเพื่อทำลายเซลล์มะเร็งเนื่องจากนี้โอแอนติเจนมีอัตราการพบน้อยมากบนผิวของเซลล์มะเร็งทำให้ทีเซลล์ตรวจพบแต่แอนติเจนที่ร่างกายผลิตจึงไม่ทำ



รูปที่ 1. แสดง (ก) เซลล์ปกติ (ข) เซลล์กลายพันธุ์ (ค) เอพิโทปบนนี้โอแอนติเจน

ปฏิกิริยาเป็นสาเหตุให้เซลล์มะเร็งไม่ถูกทำลายและแพร่กระจายอย่างรวดเร็วโดยบางส่วนของแอนติเจน ที่ทำปฏิกิริยาแบบจำเพาะกับแอนติบอดีจนเกิดการตอบสนองของ ภูมิคุ้มกันของร่างกายเรียกว่า Antigenic determinant หรือ เอพิโทป ตำแหน่งที่แอนติบอดีเข้าจับกับเอพิโทป เรียกว่า พาราโทป

2.1 การเรียนรู้ของเครื่อง (Machine Learning)

การทำนายแอนติเจนของเซลล์มะเร็งเป็นการเรียนรู้แบบมีผู้สอน (Supervised Learning) แบบ 2 กลุ่มคือใช่แอนติเจนหรือไม่ใช่แอนติเจน (binary classes) วิธีจำแนก (Classification Algorithm) ที่นำมาใช้งานวิจัยนี้มีจำนวน 7 ขั้นตอนวิธีดังนี้

2.1.1 ต้นไม้ตัดสินใจ (Decision Tree)

ต้นไม้ตัดสินใจเป็นเทคนิคที่สร้างโมเดลการจำแนกข้อมูลในรูปแบบของต้นไม้ประกอบด้วยโหนดราก (Root node) กิ่ง (Branch) และโหนดใบ (Leaf node) ซึ่งจะพิจารณาจากคุณลักษณะและค่าต่าง ๆ ที่กำหนดให้ โดยผลลัพธ์ของโครงสร้างต้นไม้จะแสดงขั้นตอนวิธีการตัดสินใจที่มนุษย์สามารถเข้าใจได้ สำหรับขั้นตอนการคำนวณสามารถพิจารณาได้จากค่า Gini impurity หรือค่าเอนโทรปี (Entropy) ในการเลือกคุณลักษณะและค่าที่เหมาะสมในการสร้างโมเดลโดยสามารถแสดงได้จากสมการที่ (1) และ (2)

$$Gini\ Impurity = 1 - \sum_{i=1}^n p_i^2 \quad (1)$$

$$Entropy = - \sum_{i=1}^n p_i \log_2(p_i) \quad (2)$$

เมื่อกำหนดให้ p_i คือค่าความน่าจะเป็นของกลุ่ม i

ในงานวิจัยนี้ใช้เทคนิคต้นไม้ตัดสินใจประเภท Classification and Regression Trees (CART) [4] เนื่องจากเหมาะสำหรับการจำแนกแบบไบนารี (Binary classification) ซึ่งจะใช้ Gini impurity ในการพิจารณาเลือกคุณลักษณะและค่าที่เหมาะสมเพื่อนำมาสร้างกฎ

2.1.2 เกรเดียนท์บูตติ้ง (Gradient Boosting)

เกรเดียนท์บูตติ้งเป็นเทคนิคที่พัฒนามาจากเทคนิคต้นไม้ตัดสินใจเพื่อนำมาใช้แก้ปัญหาการวิเคราะห์การถดถอยและการจำแนกข้อมูลได้หลักการคือจะทำการการสร้างต้นไม้ตัดสินใจ

ตามจำนวนรอบที่กำหนดโดยในแต่ละรอบจะมีประสิทธิภาพที่สูงขึ้นเนื่องจากการเรียนรู้ข้อผิดพลาดจากการทำนายรอบก่อนหน้า [5]

2.1.3 เคนเนียร์สเนเบอร์ (K-Nearest Neighbors)

เคนเนียร์สเนเบอร์เป็นเทคนิคการจำแนกข้อมูลประเภทหนึ่ง (Classification) โดยที่ไม่มีข้อมูลชุดที่ใช้ในการเรียนรู้ (Training data) แต่จะใช้วิธีการสุ่มเปรียบเทียบข้อมูลที่สนใจกับข้อมูลอื่น ๆ รอบข้างว่ามีความคล้ายคลึงมากน้อยเพียงใดตามจำนวนข้อมูลที่ล้อมรอบ (K) จำนวนหากพบว่ามีความคล้ายคลึงกันมากข้อมูลนั้นจะถูกจัดให้อยู่ในกลุ่มเดียวกันโดยการวัดความแตกต่างของข้อมูลจะใช้การวัดระยะทางแบบยูคลิด (Euclidean distance) สามารถแสดงได้ดังสมการที่ (3)

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (3)$$

เมื่อ $d(x, y)$ คือระยะทางระหว่างข้อมูล x และข้อมูล y ในปริภูมิ n มิติ

2.1.4 การถดถอยโลจิสติก (Logistic Regression)

การถดถอยโลจิสติกเป็นเทคนิคการวิเคราะห์สถิติเชิงคุณภาพ โดยมีตัวแปรตามเป็นตัวแปรเชิงคุณภาพสามารถแบ่งการวิเคราะห์เป็นสองแบบ แบบแรกการวิเคราะห์การถดถอยโลจิสติกที่ใช้กับตัวแปรที่แบ่งเป็น 2 กลุ่มย่อยมีค่าเป็น 0 กับ 1 และการวิเคราะห์การถดถอยโลจิสติกทุกกลุ่มใช้กับตัวแปรที่มีมากกว่า 2 กลุ่มเพื่อทำนายโอกาสที่จะเกิดเหตุการณ์ที่สนใจซึ่งตัวแปรที่นำมาใช้กับทั้งสองแบบต้องมีความสัมพันธ์กันตามเกณฑ์ของ Burns and Grove (1993) ค่า r ไม่เกิน 0.65 หรือตามเกณฑ์ของ Stevens (1996) ค่า r ไม่เกิน 0.80

สำหรับการคำนวณความน่าจะเป็นของการเกิดเหตุการณ์ y สามารถพิจารณาได้จากฟังก์ชันของตัวแปรทำนายดังสมการที่ (4)

$$p(y) = \frac{1}{1 + e^{-f(x)}} \quad (4)$$

เมื่อกำหนดให้

$p(y)$ คือความน่าจะเป็นของการเกิดเหตุการณ์ y

e คือ exponential function โดยให้ $e = 2.71828$

$f(x)$ คือฟังก์ชันของตัวแปรทำนาย

2.1.5 เนอ์ฟเบย์ (Naïve Bayes)

เนอ์ฟเบย์เป็นเทคนิคการแบ่งกลุ่มข้อมูลโดยใช้หลักความน่าจะเป็น (Probability) ตามทฤษฎีบทของเบย์ (Naïve Bayes Theorem) โดยกำหนดค่าความน่าจะเป็น (Probability) ให้อยู่ระหว่าง 0 ถึง 1 ถ้าค่าความน่าจะเป็นเข้าใกล้หนึ่งหมายถึงโอกาสที่จะเกิดเหตุการณ์ที่

ทำนายมีมากในทางกลับกันหากมีค่าความน่าจะเป็นเข้าใกล้ 0 แสดงว่าโอกาสที่จะเกิดเหตุการณ์ที่ทำนายน้อยมากแต่ค่าความน่าจะเป็นเท่ากับ 0 หมายความว่าเหตุการณ์ที่ทำนายจะไม่มีโอกาสเกิดขึ้นได้สำหรับวิธีการคำนวณสามารถแสดงได้จากสมการที่ (5) ดังนี้

$$P(C|A) = \frac{P(A|C)P(C)}{P(A)} \quad (5)$$

เมื่อกำหนดให้

$P(C|A)$ คือความน่าจะเป็นของข้อมูลที่มีคุณลักษณะเป็น A จะอยู่ในคลาส C

$P(A|C)$ คือความน่าจะเป็นของข้อมูลที่อยู่ในคลาส C และมีคุณลักษณะเป็น A

$P(C)$ คือความน่าจะเป็นของคลาส C

$P(A)$ คือจำนวนคุณลักษณะทั้งหมด

2.1.6 ป่าสุ่ม (Random Forest) [25]

ป่าสุ่มเป็นเทคนิคการจำแนกที่ได้รับความนิยมในปัจจุบันพัฒนาต่อยอดมาจากวิธีการจำแนกต้นไม้ตัดสินใจ (Decision Tree) หลาย ๆ ต้นเพื่อเพิ่มประสิทธิภาพการทำนายให้แม่นยำมากขึ้น โดยเอาค่าการตัดสินใจของแต่ละตัวแบบมาลงคะแนนเพื่อเลือกกลุ่ม (Class) ใดถูกเลือกมากที่สุด

2.1.7 ซัพพอร์ทเวกเตอร์แมชชีน (Support Vector Machine) [26]

ซัพพอร์ทเวกเตอร์แมชชีนเป็นเทคนิคการจำแนกข้อมูลโดยการวิเคราะห์การถดถอย (Regression) และการจัดกลุ่มข้อมูล (Classification) โดยการหาค่าสัมประสิทธิ์ของสมการเพื่อสร้างเส้นแบ่งกลุ่มข้อมูลที่ดีที่สุดใช้ได้ทั้งสมการเชิงเส้นตรง (Linear) และสมการไม่เชิงเส้นตรง (Non Linear)

2.2 เทคนิคการปรับเพิ่มข้อมูลด้วยวิธีสุ่ม

Synthetic Minority Oversampling Technique (SMOTE) คือเทคนิคที่ช่วยแก้ปัญหาข้อมูลที่มีจำนวนแตกต่างกันมากใน แต่ละคลาสส่งผลให้ผลลัพธ์ของการจำแนกอาจอยู่ใน ข้อมูลกลุ่มที่มีจำนวนมากกว่า เทคนิค SMOTE นี้ทำการ เพิ่มจำนวนของข้อมูลในกลุ่มที่น้อยกว่าให้ เท่ากับกลุ่มที่มีข้อมูลมากกว่า โดยอาศัยหลักการ เคเนียร์เซนเบอร์ (K-Nearest Neighbors) วิธีทำเริ่มจากสุ่มข้อมูลของกลุ่มที่มีจำนวนน้อย มาจำนวน 1 ข้อมูลหลังจากนั้นคำนวณหาระยะทางด้วยวิธี Euclidean distance เพื่อหาข้อมูลที่ใกล้เคียงจำนวน k ข้อมูล สร้างข้อมูลเทียม ที่อยู่ระหว่าง ข้อมูลที่สุ่มกับค่า k และสุ่มเลือกข้อมูลอีกหลายครั้งเพื่อสร้างข้อมูลเทียม จนครบตามจำนวนที่ต้องการ

2.3 ทบทวนงานวิจัยที่เกี่ยวข้อง

ปัจจุบันมีงานวิจัยหลายชิ้นพยายามสร้างตัวแบบสำหรับทำนาย เปปไทด์ประเภทต่าง ๆ เช่นเปปไทด์ต้านมะเร็ง, เปปไทด์ต้านจุลชีพ และ โปรตีนที่เกี่ยวข้องกับกลไกต่าง ๆ ของร่างกายข้อมูลที่นำมาใช้สำหรับฝึกฝนและทดสอบตัวแบบมาจากหลายแหล่ง ดังนี้

ในปี 2019 Md. Mehedi Hasan [16] ได้ นำเสนอ DiscriminationHLPpred-Fuse: improved and robust prediction of hemolytic peptide and its activity by fusing multiple feature representation Discrimination เพื่อสร้างตัวแบบสำหรับทำนาย hemolytic peptide และ กิจกรรมของเปปไทด์นี้ ซึ่งใช้วิธีจำแนกจำนวน 5 ขั้นตอนวิธีดังนี้ SVM, RF, GB, KNN และ AdaBoost ใช้คุณลักษณะจำนวน 9 คุณลักษณะหลักดังนี้องค์ประกอบกรดอะมิโน Amino Acid Composition (AAC), องค์ประกอบกรดอะมิโนคู่ Dipeptide Composition (DPC), Amino acid index (AAI), Binary profile (BPF), Composition transition-distribution (CTD), Conjoint triad (CTF), Quasi-sequence order (QSO), Grouped dipeptide composition (GDPC) และ Grouped tripeptide composition (GTPC) ข้อมูลแบ่งออกเป็น 2 ส่วนคือ Training และ Independent หลังจากนั้นแบ่งอีกครั้งเป็น 2 ชั้น และแบ่งข้อมูลฝึกฝนและทดสอบแบบ 10 fold cross validation

ประเมินประสิทธิภาพด้วยค่าความแม่นยำปรับสมดุล (balanced accuracy (BACC)), ความไว (Sensitivity), ความจำเพาะ (Specificity), สัมประสิทธิ์ของ Matthews (MCC), พื้นที่ใต้กราฟเส้นโค้ง The area under curve (AUC) และ p-value ผลที่ดีที่สุดได้จากวิธีจำแนก PTPD มีค่าความไวสูงกว่า ตัวแบบที่เคยตีพิมพ์แล้วจำนวน 1 วิธี ผลสรุปดังตารางที่ 1

ตารางที่ 1. แสดงผลการทดสอบเปรียบเทียบตัวแบบ HLPred-Fuse กับ HemoPI

Methods	MCC	BAC C	Sn	Sp	AUC	P-value
HLPred	0.58	0.792	0.82	0.76	0.86	-
--Fuse	5		1	2	9	
HemoPI	0.57	0.780	0.88	0.67	0.84	0.40
	7		7	3	2	4

ในปี 2019 Sayamon Hongjaisee, Chanin Nantasenamat, Tanawan Samleerat Carraway และ Watshara Shoombuatong [27] ได้นำเสนอ “HIVCoR: A sequence-based tool for predicting HIV-1 CRF01_AE coreceptor usage” เพื่อสร้างตัวแบบสำหรับทำนาย HIV-1 CRF01_AE coreceptor ซึ่งใช้วิธีจำแนกป่าสุ่ม (Random Forest) และซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine) โดยใช้คุณลักษณะ องค์ประกอบกรดอะมิโน (Amino Acid Composition (AAC)), องค์ประกอบกรดอะมิโนเทียม (Pseudo Amino AcidComposition (PseAAC)) และ Relative Synonymous Codon Usage (RSCU) แบ่งข้อมูลฝึกฝนและทดสอบแบบ 10 fold cross validation ข้อมูล benchmark ชนิด R5 Internal set จำนวน 30 และ External set จำนวน 120 และ ข้อมูล benchmark ชนิด X4 Internal set จำนวน 30 และ External set จำนวน 10 ประเมินประสิทธิภาพด้วยค่าความแม่นยำ (Accuracy), ความไว (Sensitivity), ความจำเพาะ (Specificity), สัมประสิทธิ์ของ Matthews (MCC) และ พื้นที่ใต้กราฟเส้นโค้ง The area under curve (AUC) ผลที่ดีที่สุดได้จากวิธีจำแนก PTPD มีค่าความไวสูงกว่า ตัวแบบที่เคยตีพิมพ์แล้วจำนวน 5 วิธี ดังแสดงในตารางที่ 2

ตารางที่ 2. เปรียบเทียบประสิทธิภาพของ 6 ตัวแบบ

Validation	Model	AC (%)	SN (%)	Sp (%)	MCC	AUC
10-fold CV	SVM/ LMT	88.43	90.32	88.11	0.65	0.96
	HIVCoR	95.29	94.38	100.00	0.86	1.00
Internal test	SVM/ LMT	73.13	72.89	73.83	0.46	0.85
	HIVCoR	93.67	93.23	94.30	0.87	0.99
External test	SVM/ LMT	74.24	75.95	73.99	0.36	0.86
	HIVCoR	93.80	93.90	93.79	0.71	0.99

3. วิธีดำเนินการวิจัย

3.1 การเก็บรวบรวมและเตรียมข้อมูล

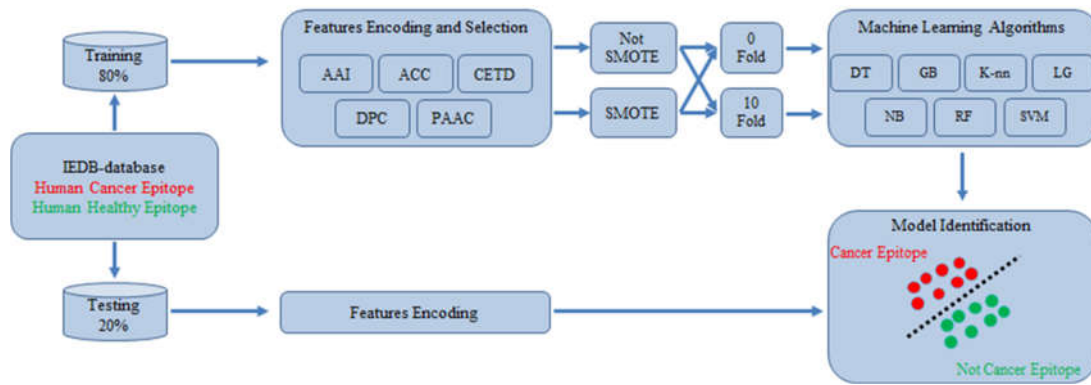
ข้อมูลที่นำมาใช้ในงานวิจัยนี้รวบรวมจาก IEDB (<https://www.iedb.org/>) ข้อมูลเอพิโทปจากเซลล์ของผู้ป่วยมะเร็งและจากเซลล์ของผู้ที่สุขภาพดีหลังจากนั้นแยกข้อมูลที่เข้าช้อนออกจะได้ข้อมูลที่มีรายละเอียดตามตารางที่ 3. แปลงข้อมูลจาก csv format เปลี่ยนเป็น fasta format เพื่อสกัดคุณลักษณะ

ตารางที่ 3 แสดงจำนวนข้อมูลที่ได้จาก IEDB คัดข้อมูลที่เข้าออก และเพิ่มข้อมูลให้เท่ากันด้วยวิธี SMOTE

แหล่งที่มา (iedb)	จำนวน (เอพิโทป)	จำนวนไม่ซ้ำ (เอพิโทป)	SMOTE
เซลล์ของผู้ป่วยมะเร็ง	947	804	2,421
เซลล์ของผู้ที่สุขภาพดี	2,549	2,421	2,421

3.2 การสกัดคุณลักษณะ

ในงานวิจัยนี้ได้ทำการสกัดคุณลักษณะ โดยใช้เครื่องมือ Pfeature [13] จำนวน 5 คุณลักษณะ แต่ละคุณลักษณะประกอบด้วยคุณลักษณะย่อยมีรายละเอียดดังตารางที่ 4 และมีขั้นตอนการดำเนินงานตามรูปที่ 2



รูปที่ 2 แสดงขั้นตอนการเลือกตัวแบบที่เหมาะสมสำหรับทำนายเอพิโทปของเซลล์มะเร็ง

ตารางที่ 4 แสดงกรดอะมิโน 20 ชนิด

Abbreviation Name			Sidechain
Alanine	Ala	A	hydrophobic
Cysteine	Cys	C	hydrophobic
Aspartate	Asp	D	negative
Glutamate	Glu	E	negative
Phenylalanine	Phe	F	hydrophobic
Glycine	Gly	G	hydrophobic
Histidine	His	H	positive
Isoleucine	Ile	I	hydrophobic
Lysine	Lys	K	positive
Leucine	Leu	L	hydrophobic
Methionine	Met	M	hydrophobic
Asparagine	Asn	N	polar
Proline	Pro	P	hydrophobic
Glutamine	Gln	Q	polar
Arginine	Arg	R	positive
Serine	Ser	S	polar
Threonine	Thr	T	polar
Valine	Val	V	hydrophobic
Tryptophan	Trp	W	hydrophobic
Tyrosine	Tyr	Y	polar

3.2.1 องค์ประกอบกรดอะมิโน (Amino Acid Composition (AAC))

องค์ประกอบกรดอะมิโน คือการคำนวณหาความถี่ของกรดอะมิโนทั้ง 20 ชนิดในสายเปปไทด์ โดยสามารถคำนวณได้จากสมการที่ 6

$$f(x) = \frac{B(x)}{L} \quad x \in \{A, C, D, \dots, Y\} \quad (6)$$

3.2.2 องค์ประกอบกรดอะมิโนคู่ (Dipeptide Composition (DPC))

องค์ประกอบกรดอะมิโนคู่ คือการเรียงสับเปลี่ยนกรดอะมิโน 20 ชนิดทีละคู่สามารถเข้ากันได้ ทำให้เกิดคุณลักษณะใหม่จำนวน 400 คุณลักษณะ สามารถคำนวณได้จากสมการที่ 7

$$f(x, y) = \frac{B(x, y)}{L - 1} \quad x \in \{A, C, D \dots Y\} \quad (7)$$

3.2.3 ดัชนีกรดอะมิโน (Amino acid index (AAI))

ดัชนีกรดอะมิโน คือชุดค่าตัวเลข 20 ค่าของคุณสมบัติทางเคมีกายภาพและชีวภาพที่แตกต่างกันของกรดอะมิโน AAindex1 ทำการรวบรวมดัชนีที่เผยแพร่พร้อมกับผลลัพธ์ของคลัสเตอร์มาทำการวิเคราะห์โดยใช้สัมประสิทธิ์สหสัมพันธ์เป็นระยะห่างระหว่าง 2 ดัชนีไว้ทั้งหมด 556 ดัชนี เช่น Hydrophobicity index (Argos et al., 1982), alpha-CH chemical shifts (Andersen et al., 1992), Signal sequence helical potential (Argos et al., 1982), Residue volume (Bigelow, 1967)

3.2.4 Composition enhanced Transition and Distribution (CETD)

Composition enhanced Transition and Distribution (CETD) แสดงคุณสมบัติของเปปไทด์ ไม่ชอบน้ำ (Hydrophobicity), ปริมาตรของ Waals (Norm Waals Volume), ขั้ว (Polarity),

ความสามารถในการโพลาไรซ์ (Polarizability), ประจุ (Charge), โครงสร้างรอง (Secondary Structure) และการเข้าถึงตัวทำละลาย (Solvent Access) ของการเข้ารหัสสามประเภท ดังนี้

3.2.4.1 องค์ประกอบ (C) คือจำนวนกรดอะมิโนของคุณสมบัติเฉพาะเช่น ความไม่ชอบน้ำ หากรด้วยจำนวนกรดอะมิโนทั้งหมด

3.2.4.2 การเปลี่ยนภาพ (T) คือลักษณะความถี่ร้อยละที่กรดอะมิโนของคุณสมบัติเฉพาะตามด้วยกรดอะมิโนที่มีคุณสมบัติต่างกัน

3.2.4.3 การกระจาย (D) คือวัดความยาวโซ่ที่กรดอะมิโนตำแหน่งแรก 25, 50, 75 และ 100 ที่คุณสมบัติเฉพาะตั้งอยู่ตามลำดับ

3.2.5 องค์ประกอบกรดอะมิโนเทียม (Pseudo Amino Acid Composition (PAAC))

องค์ประกอบกรดอะมิโนเทียมประกอบด้วยคุณลักษณะย่อยจำนวน 20 คุณลักษณะในปี 2544 Kuo-Chen Chou ได้เสนอตัวอย่างโปรตีนสำหรับปรับปรุงการทำนายการไล่กลไลโซชันของโปรตีนและทำนายประเภทโปรตีนเมมเบรนเช่นเดียวกับวิธีองค์ประกอบกรดอะมิโน (Amino acid composition (AAC)) วิธีนี้จะใช้เมทริกซ์ความถี่ของกรดอะมิโนเป็นหลักเพื่อกำหนดลักษณะของโปรตีนซึ่งช่วยให้โปรตีนไม่มีความคล้ายคลึงกันของโปรตีนอื่น ๆ เมื่อเปรียบเทียบกับ วิธีองค์ประกอบกรดอะมิโน (ACC) แล้วข้อมูลเพิ่มเติมจะรวมอยู่ในเมทริกซ์เพื่อแสดงคุณลักษณะในท้องถิ่นบางอย่าง เช่น ความสัมพันธ์ระหว่างสิ่งตกค้างในระยะทางที่กำหนด เมื่อจัดการกับกรณีขององค์ประกอบกรดอะมิโนเทียม มักใช้ทฤษฎีบทค่าคงที่ของ Chou

จำนวนคุณลักษณะหลักและคุณลักษณะย่อยที่นำมาใช้ในงานวิจัยนี้แสดงในตารางที่ 5

ตารางที่ 5 แสดงคุณลักษณะหลักและจำนวนคุณลักษณะย่อย

คุณลักษณะ	คุณลักษณะย่อย
Amino Acid Composition (AAC)	20
Amino Acid Index (AAI)	553
Composition enhanced Transition and Distribution (CETD)	189
Dipeptide Composition (DPC)	400
Pseudo Amino Acid Composition (PAAC)	20

3.3 การแบ่งข้อมูล

หลังจากนั้นเป็นการทำ cross validation โดยแบ่งข้อมูลออกเป็น 10 ส่วน ซึ่งข้อมูลจำนวน 9 ส่วนจะนำมาใช้สอนเครื่องให้เรียนรู้ (train) ส่วนข้อมูลที่เหลืออีก 1 ส่วนจะนำมาใช้สำหรับเป็นข้อมูลทดสอบ (test) โดยสลับข้อมูลทดสอบ จำนวน 10 ครั้ง วิธีนี้สามารถป้องกันความลำเอียง (bias) ที่เกิดจากข้อมูลที่มีการกระจายตัวน้อย ซึ่งส่งผลให้ตัวแบบทำนายข้อมูลที่ใช้ทำนายจริงมีความแม่นยำน้อย (overfitting)

3.4 การประเมินประสิทธิภาพของ ตัวแบบ

ในงานวิจัยนี้ วิธีจำแนกทั้ง 7 วิธี ใช้ค่าเริ่มต้นของแต่ละวิธี โดยไม่มีการปรับค่าใด ๆ การประเมินประสิทธิภาพของตัวแบบแต่ละตัวที่เป็นแบบ สองกลุ่ม (Binary Classes) ใช้ค่าสำหรับการประเมินดังนี้ [16-20]

ความแม่นยำ (Accuracy) คืออัตราของผลรวมถูกต้องทั้งหมดกับผลรวมทั้งหมด ดังสมการที่ 8

ความไว (Sensitivity) คืออัตราผลบวกจริงของการทดสอบหรือความน่าจะเป็นของผลบวกจริง ดังสมการที่ 9

ความจำเพาะ (Specificity) คืออัตราผลลบจริงของการทดสอบหรือ ความน่าจะเป็นของผลลบจริง ดังสมการที่ 10

สัมประสิทธิ์ของ Matthews (MCC) เหมาะสำหรับการประเมินข้อมูลที่ไม่สมดุลโดยจะให้ความสำคัญกับค่าความสับสนทั้ง 4 ค่าคือผลบวกจริง (TP), ผลลบจริง (TN), ผลบวกเท็จ (FP) และผลลบเท็จ (FN) ดังสมการที่ 11 ซึ่งจะต่างจากค่า F1 ที่ไม่สนใจค่าผลลบจริง

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (8)$$

$$Sensitivity = \frac{TP}{TP + FN} \quad (9)$$

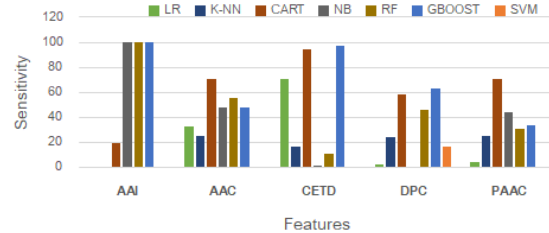
$$Specificity = \frac{TN}{TN + FP} \quad (10)$$

$$MCC = \frac{TP \times TN - FP \times FN}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}} \quad (11)$$

หลังจากนั้นนำค่าอัตราผลบวกจริงและอัตราผลลบจริงแสดงในกราฟเส้นโค้ง Receiver Operating Characteristic (ROC) [23] เพื่อหาจุดตัดการแบ่งกลุ่มที่ดีที่สุดและคำนวณหา

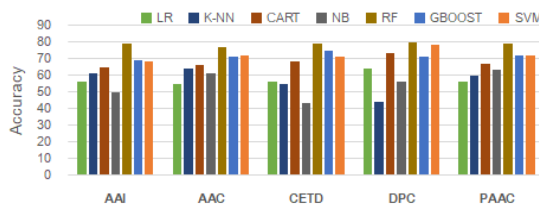
พื้นที่ใต้กราฟ The area under curve (AUC) เพื่อประเมินภาพรวม
ตัวแบบโดยพื้นที่ใต้กราฟ ยิ่งเข้าใกล้ 1 แสดงว่าตัวแบบนั้น
ทำนายได้ดีมาก

4. ผลการทดลองและการอภิปราย



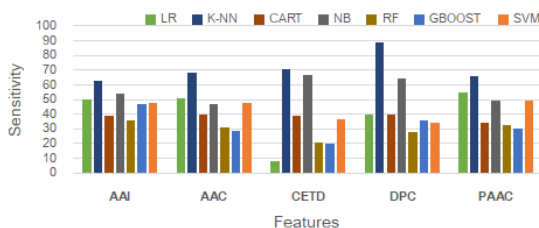
รูปที่ 3. แสดงกราฟเปรียบเทียบค่าความแม่นยำของคุณลักษณะ
และวิธีจำแนกของข้อมูลที่ไม่ได้ปรับสมดุล

จากรูปที่ 3 แสดงข้อมูลที่ยังไม่ได้ปรับสมดุลวิธีจำแนก SVM มี
ค่าความแม่นยำสูงที่สุดเมื่อใช้คุณลักษณะ AAI, AAC, CETD
และ DPC ในขณะที่คุณลักษณะ PAAC มีค่าความแม่นยำสูงที่สุด
เมื่อถูกใช้โดยวิธีจำแนก KNN



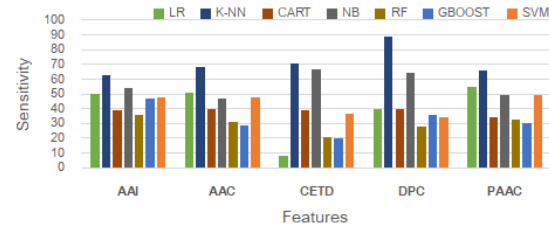
รูปที่ 4. แสดงกราฟเปรียบเทียบค่าความแม่นยำของ
คุณลักษณะ และวิธีจำแนกของข้อมูลที่ได้ปรับสมดุลด้วย SMOTE

จากรูปที่ 4 แสดงให้เห็นว่าในข้อมูลที่ได้ปรับสมดุลด้วย SMOTE
วิธีจำแนก RF มีค่าความแม่นยำสูงที่สุดในทุกคุณลักษณะ



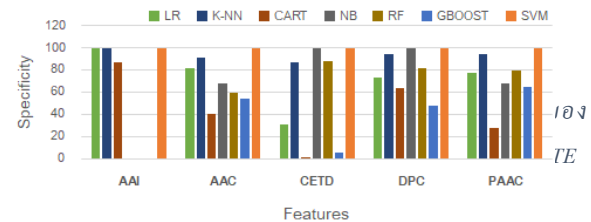
รูปที่ 5. แสดงกราฟเปรียบเทียบค่าความไวของคุณลักษณะ
และวิธีจำแนกของข้อมูลที่ไม่ได้ปรับสมดุล

จากรูปที่ 5 แสดงให้เห็นว่าในข้อมูลที่ยังไม่ได้ปรับสมดุลวิธี
จำแนก GBOOST มีค่าความไวสูงที่สุดเมื่อใช้คุณลักษณะ AAI,
CETD และ DPC ในขณะที่คุณลักษณะ PAAC มีค่าความไว สูงสุด
เมื่อถูกใช้โดยวิธีจำแนก NB, RF และ GBOOST



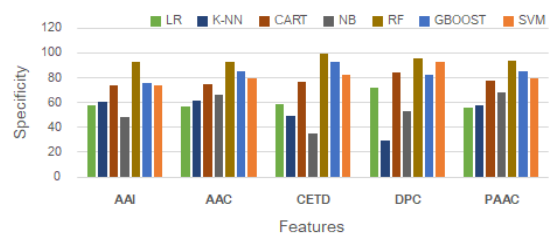
รูปที่ 6. แสดงกราฟเปรียบเทียบค่าความไวของคุณลักษณะ
และวิธีจำแนกของข้อมูลที่ได้ปรับสมดุลด้วย SMOTE

จากรูปที่ 6 แสดงให้เห็นว่าในข้อมูลที่ได้ปรับสมดุลด้วย SMOTE
วิธีจำแนก KNN มีค่าความไวสูงที่สุดในทุกคุณลักษณะ



รูปที่ 7. แสดงกราฟเปรียบเทียบค่าความจำเพาะของ
คุณลักษณะและวิธีจำแนกของข้อมูลที่ไม่ได้ปรับสมดุล

จากรูปที่ 7 แสดงให้เห็นว่าในข้อมูลที่ยังไม่ได้ปรับสมดุลวิธี
จำแนก SVM มีค่าความจำเพาะสูงที่สุดในทุกคุณลักษณะ



รูปที่ 8. แสดงกราฟเปรียบเทียบค่าความจำเพาะของ
คุณลักษณะและวิธีจำแนกของข้อมูลที่ได้ปรับสมดุลด้วย SMOTE

จากรูปที่ 8 แสดงให้เห็นว่าในข้อมูลที่ได้ปรับสมดุลด้วย SMOTE
วิธีจำแนก RF มีค่าความจำเพาะสูงที่สุดในทุกคุณลักษณะ

คุณลักษณะ AAI ข้อมูลไม่ได้ปรับสมดุลวิธีจำแนกการถดถอยโลจิสติก (LR), เคนเนยเรสเนเบอร์ (K-NN) และซัพพอร์ตเวกเตอร์แมชชีน (SVM) มีค่า ความแม่นยำสูงสุดที่ 75% ค่าความจำเพาะที่ 100% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.56, 0.51 และ 0.50 ตามลำดับส่วนข้อมูลปรับสมดุลวิธีจำแนกป่าสุ่ม (RF) ค่าความแม่นยำสูงสุดที่ 79% ค่าความจำเพาะที่ 93% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.71

คุณลักษณะ AAC ข้อมูลไม่ได้ปรับสมดุลวิธีจำแนกเคนเนยเรสเนเบอร์ (K-NN) และซัพพอร์ตเวกเตอร์แมชชีน (SVM) มีค่าความแม่นยำสูงสุดที่ 75% ค่าความจำเพาะที่ 91% และ 100% ตามลำดับและมีพื้นที่ใต้กราฟเส้นโค้ง 0.66 และ 0.50 ตามลำดับส่วนข้อมูลปรับสมดุลวิธีจำแนกป่าสุ่ม (RF) ค่าความแม่นยำสูงสุดที่ 77% ค่าความจำเพาะที่ 93% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.73

ตารางที่ 6. ผลประเมินประสิทธิภาพคุณลักษณะและตัวแบบ ที่ไม่ได้ปรับสมดุลข้อมูล

Feature	Classifier	Ac (%)			Sn (%)			Sp (%)			MCC			AUC		
		0 fold	10 fold	test	0 fold	10 fold	test	0 fold	10 fold	test	0 fold	10 fold	test	0 fold	10 fold	test
AAI	LR	75	75	75	3	3	0	99	99	100	0.07	0.07	N/A	0.59	0.59	0.56
	K-NN	73	73	75	24	24	0	90	90	100	0.17	0.17	N/A	0.65	0.65	0.51
	CART	64	64	70	32	32	19	75	75	87	0.07	0.07	0.06	0.55	0.55	0.53
	NB	55	55	25	44	44	100	59	59	0	0.03	0.03	N/A	0.55	0.55	0.50
	RF	77	77	25	12	12	100	98	98	0	0.21	0.21	N/A	0.70	0.70	0.47
	GBOOST	71	71	25	19	19	100	88	88	0	0.10	0.10	N/A	0.66	0.66	0.52
	SVM	77	77	75	11	11	0	98	98	100	0.20	0.20	N/A	0.69	0.69	0.50
AAC	LR	75	75	69	3	3	33	99	99	81	0.08	0.08	0.14	0.60	0.60	0.59
	K-NN	75	75	75	29	29	25	91	91	91	0.25	0.25	0.20	0.67	0.67	0.66
	CART	69	69	48	42	42	71	79	79	40	0.21	0.21	0.10	0.59	0.59	0.56
	NB	73	73	63	22	22	48	90	90	68	0.16	0.16	0.14	0.59	0.59	0.59
	RF	78	78	58	21	21	55	98	98	59	0.32	0.32	0.12	0.72	0.72	0.61
	GBOOST	75	75	52	17	17	48	95	95	54	0.19	0.19	0.01	0.63	0.63	0.51
	SVM	77	77	75	15	15	0	98	98	100	0.25	0.25	N/A	0.65	0.65	0.50
CETD	LR	74	74	41	9	9	71	96	96	31	0.09	0.09	0.01	0.59	0.59	0.49
	K-NN	74	74	69	25	25	16	90	90	87	0.19	0.19	0.04	0.62	0.62	0.55
	CART	66	66	24	37	37	94	76	76	1	0.12	0.12	-0.14	0.57	0.57	0.47
	NB	46	46	75	68	68	1	38	38	100	0.05	0.05	0.03	0.53	0.53	0.54
	RF	78	78	69	17	17	11	98	98	88	0.28	0.28	-0.02	0.66	0.66	0.45
	GBOOST	73	73	28	11	11	97	95	95	5	0.09	0.09	0.04	0.59	0.59	0.53
	SVM	76	76	75	7	7	0	99	99	100	0.15	0.15	N/A	0.71	0.71	0.50
DPC	LR	73	73	65	37	37	42	84	84	73	0.22	0.22	0.13	0.61	0.61	0.60
	K-NN	75	75	77	21	21	24	93	93	94	0.20	0.20	0.26	0.66	0.66	0.68
	CART	73	73	61	37	37	58	84	84	63	0.22	0.22	0.18	0.60	0.60	0.60
	NB	40	40	75	83	83	0	26	26	100	0.09	0.09	N/A	0.58	0.58	0.50
	RF	80	80	72	25	25	46	98	98	81	0.38	0.38	0.26	0.72	0.72	0.71
	GBOOST	74	74	52	27	27	63	90	90	48	0.21	0.21	0.10	0.63	0.63	0.57
	SVM	79	79	79	22	22	16	98	98	100	0.34	0.34	0.34	0.69	0.69	0.73
PAAC	LR	75	75	66	2	2	34	99	99	77	0.04	0.04	0.11	0.59	0.59	0.59
	K-NN	76	76	77	32	32	25	90	90	94	0.28	0.28	0.27	0.61	0.61	0.67
	CART	68	68	39	41	41	71	77	77	28	0.17	0.17	0.00	0.59	0.59	0.50
	NB	73	73	62	23	23	44	90	90	68	0.17	0.17	0.11	0.60	0.60	0.58
	RF	79	79	67	21	21	31	98	98	79	0.32	0.32	0.10	0.71	0.71	0.61
	GBOOST	76	76	57	17	17	34	95	95	65	0.20	0.20	-0.01	0.64	0.64	0.49
	SVM	77	77	75	17	17	0	97	97	100	0.26	0.26	N/A	0.62	0.62	0.50

ตารางที่ 7. ผลประเมินประสิทธิภาพคุณลักษณะและตัวแบบ ที่ปรับสมดุลข้อมูลด้วยวิธี SMOTE

Feature	Classifier	Ac (%)			Sn (%)			Sp (%)			MCC			AUC		
		0 fold	10 fold	test	0 fold	10 fold	test	0 fold	10 fold	test	0 fold	10 fold	test	0 fold	10 fold	test
AAI	LR	60	58	56	50	56	50	63	60	58	0.11	0.16	0.07	0.60	0.60	0.60
	K-NN	60	74	61	63	93	63	59	55	61	0.19	0.53	0.21	0.65	0.65	0.67
	CART	66	74	65	40	79	39	74	69	74	0.13	0.48	0.12	0.57	0.57	0.56
	NB	52	55	50	59	62	54	50	47	48	0.08	0.09	0.02	0.56	0.56	0.55
	RF	77	90	79	30	90	36	91	91	93	0.25	0.80	0.36	0.69	0.69	0.71
	GBOOST	67	72	69	34	73	47	77	71	76	0.10	0.44	0.22	0.59	0.59	0.63
AAC	SVM	66	72	68	42	72	48	73	73	74	0.14	0.45	0.20	0.64	0.64	0.68
	LR	56	58	55	54	59	51	56	58	57	0.09	0.17	0.07	0.59	0.59	0.58
	K-NN	63	76	64	68	93	68	61	60	62	0.25	0.56	0.26	0.70	0.70	0.69
	CART	68	76	66	41	77	40	76	74	75	0.17	0.51	0.15	0.59	0.59	0.58
	NB	62	61	61	52	58	47	65	64	66	0.16	0.22	0.11	0.61	0.61	0.60
	RF	77	87	77	30	80	31	92	93	93	0.28	0.74	0.30	0.70	0.70	0.73
CETD	GBOOST	71	79	71	31	73	29	84	84	85	0.16	0.57	0.16	0.57	0.57	0.64
	SVM	69	78	72	38	77	48	79	79	80	0.17	0.56	0.27	0.68	0.68	0.70
	LR	56	62	56	51	63	48	58	60	59	0.07	0.24	0.06	0.55	0.55	0.56
	K-NN	56	71	55	73	96	71	51	47	49	0.20	0.49	0.17	0.67	0.67	0.63
	CART	68	77	68	43	78	39	75	77	77	0.17	0.55	0.16	0.59	0.59	0.58
	NB	44	55	43	68	75	67	37	35	35	0.04	0.10	0.02	0.52	0.52	0.53
DPC	RF	78	89	79	14	79	21	99	98	99	0.28	0.78	0.35	0.73	0.73	0.72
	GBOOST	70	81	75	20	73	20	87	88	93	0.08	0.62	0.18	0.57	0.57	0.58
	SVM	70	76	71	40	74	37	80	78	82	0.19	0.53	0.20	0.65	0.65	0.64
	LR	68	74	64	48	76	40	74	72	72	0.20	0.48	0.11	0.68	0.68	0.62
	K-NN	44	62	44	88	99	89	30	26	29	0.17	0.36	0.19	0.67	0.67	0.66
	CART	72	80	73	44	78	40	80	82	84	0.24	0.60	0.25	0.62	0.62	0.62
PAAC	NB	55	67	56	64	85	64	52	50	53	0.14	0.37	0.15	0.61	0.61	0.60
	RF	82	89	80	31	81	28	98	97	96	0.42	0.79	0.36	0.77	0.77	0.74
	GBOOST	72	79	71	39	76	36	82	83	82	0.21	0.58	0.19	0.65	0.65	0.60
	SVM	78	88	78	25	82	34	95	93	93	0.29	0.75	0.34	0.70	0.70	0.67
	LR	61	56	56	62	55	55	60	57	56	0.20	0.12	0.10	0.63	0.63	0.60
	K-NN	60	78	60	63	95	66	59	60	58	0.20	0.59	0.21	0.64	0.64	0.68
PAAC	CART	65	77	67	35	78	34	77	77	78	0.11	0.54	0.12	0.56	0.56	0.56
	NB	63	59	63	56	55	49	66	64	68	0.20	0.19	0.15	0.63	0.63	0.63
	RF	75	88	79	26	82	33	94	93	94	0.28	0.76	0.35	0.67	0.67	0.73
	GBOOST	70	78	72	34	73	30	84	83	85	0.20	0.56	0.17	0.61	0.61	0.64
	SVM	69	79	72	47	78	49	78	80	80	0.25	0.57	0.28	0.68	0.68	0.70

AAI: Amino Acid Index, AAC: amino acid composition, CETD: Composition enhanced Transition and Distribution, DPC: dipeptide composition, PAAC: pseudo amino acid composition

LR: การถดถอยโลจิสติก, K-NN: เคเนย์เรสเนมเบอร์, CART: ต้นไม้ตัดสินใจ, NB: เนอ์ฟเบย์, RF: ป่าสุ่ม, GBOOST: เกรเดียนท์บูตติ้ง, SVM: ซัพพอร์ตเวกเตอร์แมชชีน

Ac: ความแม่นยำ, Sn: ความไว, Sp: ความจำเพาะ, MCC: สัมประสิทธิ์ของ Matthews, AUC: พื้นที่ใต้กราฟเส้นโค้ง

คุณลักษณะ CETD ข้อมูลไม่ได้ปรับสมดุลวิธีจำแนกเคอเนอ (NB) และซัพพอร์ตเวกเตอร์แมชชีน (SVM) มีความแม่นยำสูงสุดที่ 75% ค่าความจำเพาะที่ 100% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.54 และ 0.50 ตามลำดับ ส่วนข้อมูลปรับสมดุลวิธีจำแนกป่าสุ่ม (RF) ค่าความแม่นยำสูงสุดที่ 79% ค่าความจำเพาะที่ 99% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.72

คุณลักษณะ DPC ข้อมูลไม่ได้ปรับสมดุลวิธีจำแนกซัพพอร์ตเวกเตอร์แมชชีน (SVM) มีความแม่นยำสูงสุดที่ 79% ค่าความจำเพาะที่ 91% และ 100% ตามลำดับ และมีพื้นที่ใต้กราฟเส้นโค้ง 0.66 และ 0.50 ตามลำดับ ส่วนข้อมูลปรับสมดุลวิธีจำแนกป่าสุ่ม (RF) ค่าความแม่นยำสูงสุดที่ 77% ค่าความจำเพาะที่ 93% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.73

คุณลักษณะ PACC ข้อมูลไม่ได้ปรับสมดุลวิธีจำแนกเคอเนอ เรสเนเบอร์ (K-NN) มีความแม่นยำสูงสุดที่ 77% ค่าความจำเพาะที่ 94% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.67 ส่วนข้อมูลปรับสมดุลวิธีจำแนกป่าสุ่ม (RF) ค่าความแม่นยำสูงสุดที่ 79% ค่าความจำเพาะที่ 94% และมีพื้นที่ใต้กราฟเส้นโค้ง 0.73

5. ข้อเสนอแนะ

เนื่องจากงานวิจัยชิ้นนี้มุ่งเน้นเพื่อหาตัวแบบสำหรับการทำนายเอพิโทป จากเซลล์ผู้ที่มีสุขภาพดีและเซลล์ผู้ป่วยมะเร็ง ผ่านคุณลักษณะและวิธีการจำแนกข้างต้น ซึ่งยังไม่ครอบคลุมทุกคุณลักษณะ และวิธีจำแนกในอนาคตควรมีการศึกษาคุณลักษณะอื่นรวมถึงคุณลักษณะผสมเพิ่มเติม และศึกษาวิธีการจำแนกการเรียนรู้เชิงลึก (Deep Learning) เพื่อนำมาประยุกต์ใช้สำหรับเพิ่มความแม่นยำในการทำงาน

เอกสารอ้างอิง

- [1] World Health Organization, "Cancer," 21 September 2021. [Online]. Available: <https://www.who.int/news-room/fact-sheets/detail/cancer>. [Accessed Sep. 20, 2021].
- [2] Bray F, Ferlay J, Soerjomataram I, Siegel RL, Torre LA, Jemal A., "Global cancer statistics 2018: GLOBOCAN estimates of incidence and mortality worldwide for 36 cancers in 185 countries," CA Cancer J Clin, 68(6), pp. 394–424, 2018.
- [3] Shahid Akbar, Ateeq Ur Rahman, Maqsood Hayat, Mohammad Sohail, "cACP: Classifying anticancer peptides using discriminative intelligent model via Chou's 5-step rules and general pseudo components," Chemometrics and Intelligent Laboratory Systems, Volume 196, 103912, ISSN 0169-7439, 2020
- [4] L. Breiman, J. H. Friedman, R. Olshen and C. J. Stone, "Classification and Regression Trees," Wadsworth International Group, Belmont, California, 1984.
- [5] Friedman, J. H., "Greedy Function Approximation: A Gradient Boosting Machine," Annals of Statistics, 29, pp. 1189-1232, 2000.
- [6] Quinlan, J. R. (1986). Induction of decision trees. Machinelearning, 1(1), 81-106, 1986.
- [7] Quinlan, J. R. (1993). C4. 5: "programs for machine learning," (Vol. 1). Morgan kaufmann, 1993.
- [8] Vita R, Mahajan S, Overton JA, Dhanda SK, Martini S, Cantrell JR, Wheeler DK, Sette A, Peters B., "The Immune Epitope Database (IEDB): 2018 update," Nucleic Acids Res. 2018 Oct 24. doi: 10.1093/nar/gky1006. [Epub ahead of print] PubMed PMID: 30357391, 2018.
- [9] UniProt, "The universal protein knowledgebase," Nucleic Acids Res. 45, D158–D169, 2016.
- [10] Peri S, Navarro JD, Kristiansen TZ, Amanchy R, Surendranath V, Muthusamy B, Gandhi TK, Chandrika KN, Deshpande N, Suresh S, et al: "Human protein reference database as a discovery resource for proteomics." Nucleic Acids Res, 32 Database: D497-501, 2004.
- [11] Yu Wan1, Zhuo Wang and Tzong-Yi Lee1, "Incorporating support vector machine with sequential minimal optimization to identify anticancer peptides: Wan et al. BMC Bioinformatics," 22:286, 2021.
- [12] Wang, L.; Niu, D.; Wang, X.; Khan, J.; Shen, Q.; Xue, Y., "A Novel Machine Learning Strategy for the

- Prediction of Antihypertensive Peptides Derived from Food with High Efficiency.” *Foods*, 10, 550, 2021.
- [13] Akshara Pande, Sumeet Patiyal, Anjali Lathwal, Chakit Arora, Dilraj Kaur, Anjali Dhall, Gaurav Mishra, Harpreet Kaur, Neelam Sharma, Shipra Jain, Salman Sadullah Usmani, Piyush Agrawal, Rajesh Kumar, Vinod Kumar, Gajendra P.S.Raghava: “Computing wide range of protein/peptide features from their sequence and structure : biorxiv,” April 04 2019.
- [14] Onkar Singh, Wen-Lian Hsu and Emily Chia-Yu Su., “Co-AMPPred for in silico-aided predictions of antimicrobial peptides by integrating composition-based features,” Singh et al. *BMC Bioinformatics*, 22:389, 2021.
- [15] Lei J, Sun L, Huang S, Zhu C, Li P, He J, Mackey V, Coy DH, He Q., “The antimicrobial peptides and their potential clinical applications,” *Am J Transl Res*. 11(7):3919–31, 2019.
- [16] Muthurulan Pushpanathan, Paramasamy Gunasekaran and Jeyaprakash Rajendhran, }Antimicrobial Peptides: Versatile Biological Properties,” Hindawi Publishing Corporation International Journal of Peptides, Volume 2013, Article ID 675391, 15 pages, <http://dx.doi.org/10.1155/2013/675391>, 2013.
- [17] Usmani SS, Bhalla S, Raghava GPS., “Prediction of antitubercular peptides from sequence information using ensemble classifier and hybrid features,” *Front Pharmacol*, 9:954, 2018.
- [18] Kao HJ, Nguyen VN, Huang KY, Chang WC, Lee TY., “SuccSite: incorporating amino acid composition and informative k-spaced amino acid pairs to identify protein succinylation sites,” *Genomics Proteomics Bioinform*, 18(2):208–19, 2020.
- [19] Huang CH, Su MG, Kao HJ, Jhong JH, Weng SL, Lee TY., “UbiSite: incorporating two-layered machine learning method with substrate motifs to predict ubiquitin-conjugation site on lysines,” *BMC Syst Biol*, 10(Suppl 1):6, 2016.
- [20] Chen SA, Lee TY, Ou YY., “Incorporating significant amino acid pairs to identify O-linked glycosylation sites on trans-membrane proteins and non-transmembrane proteins,” *BMC Bioinform*, 11:536, 2010.
- [21] Chou KC., “Prediction of protein cellular attributes using pseudo-amino acid composition,” *Proteins Struct Funct Bioinform*. 43(3):246–55, 2001.
- [22] Chou K-C., “Using amphiphilic pseudo amino acid composition to predict enzyme subfamily classes,” *Bioinformatics*, 21(1):10–9, 2005.
- [23] Hanley JA, McNeil BJ., “The meaning and use of the area under a receiver operating characteristic (ROC) curve,” *Radiology*, 143(1), pp. 29–36, 1982.
- [24] Epitope human publications, 2021. [Online]. Available: https://github.com/manon089/Epitope_human_papers.git. [Accessed Sep. 20, 2021].
- [25] Breiman, L., “Random Forests,” *Machine Learning* 45, 5–32, 2001.
- [26] Cortes, Corinna; Vapnik, Vladimir N., “Support-vector networks,” *Machine Learning*. 20 (3), pp. 273–297, 1995.
- [27] Sayamon Hongjaisee, Chanin Nantasenamat, Tanawan Samleerat Carraway, Watshara Shoombuatong, “HIVCoR: A sequence-based tool for predicting HIV-1 CRF01_AE coreceptor usage,” *Computational Biology and Chemistry*, Volume 80, Pages 419-432, ISSN 1476-9271, 2019.

CONTRIBUTORS' FORM

(to be modified as applicable and one signed copy attached with the manuscript)

Manuscript Title:

Author(s):

I/we certify that I/we have participated sufficiently in the intellectual content, conception and design of this work or the analysis and interpretation of the data (when applicable), as well as the writing of the manuscript, to take public responsibility for it and have agreed to have my/our name listed as a contributor. I/we believe the manuscript represents valid work. Neither this manuscript nor one with substantially similar content under my/our authorship has been published or is being considered for publication elsewhere, except as described in the covering letter. I/we certify that all the data collected during the study is presented in this manuscript and no data from the study has been or will be published separately. I/we attest that, if requested by the editors, I/we will provide the data/information or will cooperate fully in obtaining and providing the data/information on which the manuscript is based, for examination by the editors or their assignees. Financial interests, direct or indirect, that exist or may be perceived to exist for individual contributors in connection with the content of this paper have been disclosed in the cover letter. Sources of outside support of the project are named in the cover letter.

I/We hereby transfer(s), assign(s), or otherwise convey(s) all copyright ownership, including any and all rights incidental thereto, exclusively to the Journal, in the event that such work is published by the Journal. The Journal shall own the work, including 1) copyright; 2) the right to grant permission to republish the article in whole or in part, with or without fee; 3) the right to produce preprints or reprints and translate into languages other than English for sale or free distribution; and 4) the right to republish the work in a collection of articles in any other mechanical or electronic format.

We give the rights to the corresponding author to make necessary changes as per the request of the journal, do the rest of the correspondence on our behalf and he/she will act as the guarantor for the manuscript on our behalf.

All persons who have made substantial contributions to the work reported in the manuscript, but who are not contributors, are named in the Acknowledgment and have given me/us their written permission to be named. If I/we do not include an Acknowledgment that means I/we have not received substantial contributions from non-contributors and no contributor has been omitted.

	Name	Signature	Date signed	
1	_____	_____	_____	
2	_____	_____	_____	
3	_____	_____	_____	
4	_____	_____	_____	(up to 4 contributors for case report/images/review)
5	_____	_____	_____	
6	_____	_____	_____	(up to 6 contributors for original studies)

INFORMATION FOR AUTHORS

The Journal of Information Science and Technology (JIST) aims to be the forum through which researchers, faculties, graduate students and experts of the computer and information technology and other technological related fields share and discuss their high quality research work as well as innovation. Original research articles, practical applications and innovations in the broad area of computer and information technology are suitable for publication in JIST. Periodicity (Publication) is 2 issues per year (JAN - JUN and JULY - DECEMBER).

Accepted papers will be Double-blind peer reviewed by minimum 3 reviewers with "Accept" result.

1. Soft Computing
2. Human-Computer Interaction
3. Information Assurance and Security
4. Information Systems
5. Networking
6. Programming
7. Platform Technologies
8. System Integration and Architecture
9. Social and Professional Issues
10. Web Systems and Technologies
11. Multidisciplinary
12. e-Business and m-Business
13. Business and Information System
14. Social and Business Aspects of Convergence IT and Ubiquitous Computing
15. Internet of Things (IoT)
16. Cloud Computing
17. Fintech and blockchain

Manuscript Preparation Guide

The template for writing the journal is available for downloading (<https://ph02.tci-thaijo.org/index.php/JIST/>). Length of the paper should be in between 8-12 pages of A4 size. Text should be typewritten or printed with double-spaced in 11-point of A4 white paper with margins of 1.5" for top, 1.25" for left, and 1" for bottom and right sides. All pages must be numbered sequentially. Here are some guidelines.

- Abstract which length 100-200 words.
- Introduction which talks about the problem and related research.
- Materials and methods use for experimentation.
- Results of the experiment.
- Discussion and Conclusion, References

Submissions (Publication Charge: Free) (Indexed by TCI tier 2)

Please submit paper in MS Word via <https://ph02.tci-thaijo.org/index.php/JIST/> (Online Journal Submission)

Contact Address

The Association of Council of IT Deans (CITT)
140 Moo 1, Cheumsampan Road, Nongchok
Bangkok, Thailand 10530
Tel: 02-988-3655 ext 4115 Fax: 02-988-4027
E-mail: jist@mut.ac.th