

Journal

JOURNAL OF RESEARCH AND DEVELOPMENT ON INFORMATION AND COMMUNICATION TECHNOLOGY

ISSN 1859-3534

THÔNG TIN VÀ TRUYỀN THÔNG - MINISTRY OF INFORMATION AND COMMUNICATIONS

**Volume
2021**

**Number 1
March**

Chuyên san
Các công trình nghiên cứu
phát triển công nghệ
thông tin và truyền thông

Director & Editor-in-chief

Nguyen Van Ba

Technical Editor-in-Chief

Le Hoai Bac

University of Science,

Vietnam National University Ho Chi Minh City

Associate Technical Editor-in-Chief

Tran Xuan Nam

Le Quy Don University of Technology, Vietnam

Nguyen Khanh Van

Hanoi University of Science and Technology

Editor

Nguyen Thanh Thuy,

University of Engineering and Technology,

Vietnam National University Hanoi

Nguyen Xuan Huy

Institute of Information Technology,

Vietnam Academy of Science and Technology

Tu Minh Phuong

Posts and Telecommunications Institute of Technology, Vietnam

Tzung-Pei Hong

National University of Kaohsiung, Taiwan

Nguyen Xuan Huan

Middlesex University, London, England

Giuseppe D'Aniello

University of Salerno, Italia

Secretary

Nguyen Quang Hung

E-mail: nguyenquanghung@mic.gov.vn

Journal Address

115 Tran Duy Hung Rd., Cau Giay District, Hanoi

Tel: (84)-24-37737136, Fax: (84)-24-37737130,

Website <https://ictmag.vn>

Email: chuyensanbcvt@mic.gov.vn

Journal Office Branch

27 Nguyen Binh Khiem Street, District 1, Ho Chi Minh City,

Vietnam

Tel/Fax: (84)-28-38992898

Journal marketing and distribution

Pham Thi Lam

Email: ptlam@mic.gov.vn, Tel: 84-948503999

Tổng biên tập

Nguyễn Văn Bá

Thường trực Hội đồng biên tập**Chủ tịch**

GS. TS. Lê Hoài Bắc,

Trường Đại học Khoa học Tự nhiên,

Đại học Quốc gia TP. Hồ Chí Minh

Phó Chủ tịch

GS. TS. Trần Xuân Nam,

Trường Đại học Kỹ thuật Lê Quý Đôn

PGS. TS. Nguyễn Khanh Văn,

Trường Đại học Bách khoa Hà Nội

Thành viên

GS. TS. Nguyễn Thanh Thủy,

Trường Đại học Công nghệ, Đại học Quốc gia Hà Nội

PGS. TSKH. Nguyễn Xuân Huy,

Viện Công nghệ thông tin, Viện Hàn Lâm KHCN Việt Nam

GS. TS. Từ Minh Phương,

Học viện Công nghệ Bưu chính Viễn thông

GS. TS. Tzung-Pei Hong,

National University of Kaohsiung, Đài Loan

GS. TS. Nguyễn Xuân Huân,

Middlesex University, Anh

TS. Giuseppe D'Aniello,

University of Salerno, Italia

Thư ký

TS. Nguyễn Quang Hưng

E-mail: nguyenquanghung@mic.gov.vn

License No 365/GP-BTTTT, dated 19th December 2014;

Modified license No 233/GP-BTTTT dated 23rd May 2017.

Printed by Công ty TNHH MTV In Tàp chí Công san.

Printed and delivered in May 2021.

Price: 40,000 VND

TABLE OF CONTENTS

A Guidance Method for Robustness Surrogate Assisted Multi-objective Evolutionary Algorithms	<i>Dinh Nguyen Duc, Long Nguyen Hoai Nguyen Xuan</i>	1
A Comprehensive Survey of Frequent Itemsets Mining on Transactional Database with Weighted Items	<i>Huan Phan, Bac Le</i>	18
Robust Watermarking Method using QIM with VSS for Image Copyright Protection.....	<i>Thanh Ta Minh, Dan Giang Ngoc</i>	28
An Integrated Approach for Table Detection and Structure Recognition	<i>Hai-Hong Phan, Ngo Dai Duong</i>	41

A Guidance Method for Robustness Surrogate Assisted Multi-objective Evolutionary Algorithms

Dinh Nguyen Duc¹, Long Nguyen², Hoai Nguyen Xuan³

¹ Academy of Military Science and Technology, Hanoi, Vietnam

² National Defense Academy, Hanoi, Vietnam

³ Vietnam AI Academy, Hanoi, Vietnam

Correspondence: Dinh Nguyen Duc, nddinh76@gmail.com

Communication: received 23 February, revised 16 April, accepted 30 April

Digital Object Identifier: 10.32913/mic-ict-research.v2021.n1.948

Abstract: In the real world, multi-objective problems (MOPs) are relatively common in optimization in the areas of design, planning, decision support, etc. In fact, problems include two or many objectives, there is a class of problems called expensive problems that are problems with complex mathematical models, large computational costs, etc. They can not be solved by normal techniques, they are usually to be solved with techniques such as simulation, decomposing, problem transformation. In particular, using a surrogate model with Kriging, neural networks techniques in combination with an evolutionary algorithm is a subtle choice, with many positive results, being studied and applied in practice. However, the use of a surrogate model with Kriging, neural networks combining selection strategy, sampling... can reduce the robustness of the algorithms during the search. This paper analyzes the issues affecting the robustness of the multi-objective evolutionary algorithms (MOEAs) using surrogate models and suggests the use of a guidance technique to increase the robustness of the algorithm, through analysis, experiment and results are competitive and effective to improve the quality of MOEAs using a surrogate model to solve expensive problems.

Keywords: MOEA, robustness, surrogate, Kriging, FNN, CSEA, K-RVEA

I. INTRODUCTION

In optimization studies including multi-objective optimization, the main focus is on finding global or Pareto-optimized solutions that represent the best objective values. In particular, algorithms using surrogate models, techniques are applied to reduce the number of calculations of the original objective function, increase the number of times using a surrogate objective function based on predicting progress of evolution. At that time, the user may not always be interested in finding the so-called global best solutions,

especially when these solutions are quite sensitive to the inevitable fluctuations in applying techniques in the use of surrogate models. In such cases, the problem of research is how to find powerful solutions that are less sensitive to disturbances from the characteristics of the technique used.

In this paper, we analyzed the effects of the mechanism using surrogate models in the multi-objective evolutionary algorithms, focusing on techniques using Kriging, neural network models, hypothesize and propose the use of automatic guidance techniques to create self-adaptive mechanisms for algorithms to improve the sustainability of the solution, convergence quality and diversity of the solution set, exploration and exploitation of the search process.

The remainder of the paper is organized as follows: A brief summary of surrogate assisted multi-objective evolutionary algorithms is given in Section II and the discussion of the robustness of MOEAs in Section III. Detail of our proposed guidance techniques are shown in Section IV. The experimental results are presented in Section V to examine the effectiveness and efficiency of proposed methods. Conclusion and future work are given in Section VI.

II. SURROGATE ASSISTED MULTI-OBJECTIVE EVOLUTIONARY ALGORITHMS

1. Surrogate models

Surrogate models are used to approximate in simulation way to reduce the computational cost for expensive problems. The models are described as below:

If we call $f(\vec{x})$ is an original fitness function of a MOP, then, we have $f'(\vec{x})$ is a meta function, which is indicated as follows:

$$f'(\vec{x}) = f(\vec{x}) + e(\vec{x}) \quad (1)$$

function $e(\vec{x})$ is the approximated error. In this case, the fitness function $f(\vec{x})$ is not to be known, the values (input or output) are cared. Based on the responses of the simulator from a chosen dataset, a surrogate is constructed, then the model generates easy representations that describe the relations between preference information of input and output variables. There are some approaches for the surrogate models, which are divided into some kinds such as: the Radial Basis Function (RBF), the Polynomial Response Surface (PRS), the Support Vector Machine (SVM), neural network and the Kriging (KRG). The common concepts of using surrogate models are: instead of using known fitness functions always, to reduce the number of calculations, the evaluations are suggested to use surrogate functions, which are simpler. Therefore, without loss of generality, this paper investigates the use of the original and surrogate functions with Kriging model in the solving process of the MOEAs for expensive problems.

In [1], the authors proposed a method named "Kriging", which is a response surface method bases on spatial prediction techniques. This method minimizes the mean squared error to build the spatial and temporal correlation among the values of an attribute. In [2], the authors developed a parametric regression model which design and analysis of computer experiments, called DACE. The model is an extension of Kriging approach for at least three dimensions problems. The model is a combination of a known function $f(x)$ and a Gaussian random process $f'(x)$ that is assumed to have mean zero and co-variance, like as:

$$\begin{aligned} E(f'(\vec{x}^{(i)}), f'(\vec{x}^{(j)})) &= Cov(f'(\vec{x}^{(i)}), f'(\vec{x}^{(j)})) \\ &= \sigma^2 R(\vec{\theta}, \vec{x}^{(i)}, \vec{x}^{(j)}), \end{aligned} \quad (2)$$

here, $\vec{x}^{(j)}$ is the correlation function with parameters θ , σ is the process variance of the response and $R(\vec{\theta}, \vec{x}^{(i)}, \vec{x}^{(j)})$.

With KRG model and popular MOEAs, there are some works, recently [3–8], for details:

In [4] the authors built a procedure which uses Kriging approximations and NSGA-II algorithm to optimize the aircraft design. The Kriging based genetic algorithm procedure gives shapes which are optimized for both the initial peak and perceived noise level of the signature along with minimum drag. The design problem concerns three objectives: the aircraft drag coefficient and the strength of the ground boom signature (using both the initial pressure rise and the perceived noise level) for each design point. In the study, Kriging approximation models are constructed from a large number of members of an initial population

using the results from repeated analyses carried out with BOOM-UA (a specified analysis tool). The approximation model is combined with the NSGA-II algorithm to minimize ground boom intensity while maintaining acceptable levels of aerodynamic cruise performance.

In [6], by providing a Design of Experiments (DOE) capability into the framework, the authors hybridized the desirable characteristics of evolutionary algorithms (EAs) and surrogate models such as RSM to obtain an efficient optimization system. Within this context, the DOE samples a number of design candidates at which the analysis code, the surrogate model is then constructed for the computationally expensive problem. Different sampling and DOE strategies can be used; Latin hypercube, RSM or DACE/Kriging.

In [7], in the problems with two objective functions, thermal resistance and pumping power have been selected to assess the performance of the micro-channel heat sink. The design variables related to the width of the micro-channel at the top and bottom, depth of the micro-channel, and width of fin, which contribute to objective functions, have been identified and a three-level full factorial design was selected to exploit the design space. The numerical solutions obtained at these design points were utilized to construct surrogate models Kriging and Radial Basis Neural Network. A hybrid multi-objective evolutionary algorithm coupled with surrogate models is applied to find out global Pareto-optimal solutions.

In [8], the author deeply discussed on surrogate MOEAs and indicated that: Kriging model appears to perform well in most situations, however it is much more computationally expensive than the rest. It is obvious that a careful consideration of RSM could lead to a situation where all objectives in a multi-objective problem are modeled using different methods, in order to maintain high quality and reduce optimization costs.

In [5] the KRG is used, the term Kriging points directly to the origin of these prediction methods dating back to the sixties, when the mining engineer Krige used Gaussian Random Field Models (GRFMs) to predict the concentration of ore in gold and uranium mines.

In [9], authors proposed such a method, MOEA/D-EGO, for dealing with expensive MOPs. The algorithm decomposes an MOP into a number of single-objective optimization sub-problems. At each iteration, a predictive distribution model is built for each individual objective in the MOP by using Fuzzy clustering and Gaussian stochastic process modeling. Then, a predictive model for the objective of each sub-problem can be induced. MOEA/D is used for maximizing the expected improvement metrics of all the sub-problems and several test points are then selected for

evaluation.

III. ROBUSTNESS OF MULTI-OBJECTIVE EVOLUTIONARY ALGORITHMS

1. General issues

The concept of robust algorithm was introduced in [10], there are two types of input for robustness system: control factors and noise factors, in what, the control factors can easily be controlled by decision maker and noise factors are hard to control of are beyond the control of a decision maker and they can be came from the environment. The methods used for robust algorithm can be broadly classified into two major groups: methods relying on gradient-based optimization for single optimization, and stochastic population-based methods for multi-objective optimization. Both use iteration-based approaches to improve solutions. In mathematical, two types of robust solutions are defined as below:

Definition 1 (Type I): A solution $\vec{s}^*$ is called robust solution of type I, if it is the global feasible Pareto-optimal solution to the following multi-objective minimization problem, in which it respect to a δ -neighborhood $\beta_\delta(\vec{s})$ of a solution $\vec{s}$ in the global Pareto optimal set S :

$$\begin{aligned} & \text{minimize } (f_1^e(\vec{s}), f_2^e(\vec{s}), \dots, f_k^e(\vec{s})); \\ & \text{subject to } \vec{s} \in S \\ & \text{where } f_i^e(\vec{s}) \text{ is defined as follows:} \end{aligned} \quad (3)$$

$$f_i^e(\vec{s}) = \frac{1}{|\beta_\delta(\vec{s})|} \int_{\vec{y} \in \beta_\delta(\vec{s})} f_i(\vec{y}) dy$$

which $|\beta_\delta(\vec{s})|$ denotes the volume of the neighborhood.

Definition 2 (Type II): A solution $\vec{s}^*$ is called a multi-objective robust solution of type II, if it is the global feasible Pareto optimal solution to following multi-objective minimization problem:

$$\begin{aligned} & \text{minimize } \vec{f}(\vec{s}) = (f_1(\vec{s}), f_2(\vec{s}), \dots, f_k(\vec{s})); \\ & \text{subject to } \vec{s} \in S \\ & \text{and } \frac{\|\vec{f}^p(\vec{s}) - \vec{f}(\vec{s})\|}{\|\vec{f}(\vec{s})\|} \leq \eta \end{aligned} \quad (4)$$

where $\vec{f}^p(\vec{s})$ is the perturbed objective vector and it can be set, for example, as the worst function value in the neighborhood. η is predefined parameter, it is small real value.

In [11], in a general review, the authors analysed the depending on how robust the original global efficient front is respect to above definition, they can be the following four main scenarios:

- The complete original efficient front is robust: this case the original efficient front remains as an efficient

front with respect to the main effective objective functions. The entire set of original Pareto-optimal solutions is robust and is the goal of the optimization during the global efficient front constructed with the mean effective objectives are somewhat worse than that constructed with original ones.

- A part off the original efficient front is no more robust: this case, in such a problem, the task of a robust algorithm would be to identify only that part of the efficient front which is robust. In this case, the efficient front corresponding to the main effective objectives does not span over the entire original efficient region.
- The complete original global efficient front is non-robust; instead an original local efficient front is robust: this case, the global efficient front of the original problem is completely dominated by a local efficient front with respect to the mean effective objectives, thereby meaning that the original global Pareto-optimal solutions are not robust solutions and are sensitive to local perturbation. This case we need to find the robust efficient front for the problems.
- A part of the original global efficient front is robust together with a part of an original local efficient front.

There are many researchers worked on the analyzing robustness of multi-objective optimization. In [12], through an introduction of a measurement of the robustness of solutions in a MOEA, the authors also discussed on the effectiveness of combining the two approaches, different possibilities having been put forward and assessed by studying various multi-objective benchmark problems. In [13], the authors proposed a method to measure the multi-objective robustness of a design alternative using the sensitivity region concept and an approach using that measure to obtain robust Pareto solutions of multi-objective optimization problems. In [14], the author introduced an approach to compute solution paths for parameter-dependent multi-objective problems. The proposal combines techniques from numerical path following and from multi-objective optimization and allows to treat different kinds of problems. In [15], the authors focused on the (decision-maker's attitude to risk - DMAR) as a source of uncertainty and present two new types of robustness in MOP. In the first type, Dominance Robustness (DR), the robust Pareto solutions are those which, if perturbed, would have a high chance to move to another Pareto solution. In the second type, preference robustness (PR), the robust Pareto solutions are those that are close to each other in configuration space. The authors proposed methods to quantify these robustness concepts, modify existing evolutionary multi-objective optimization techniques to capture robustness against the DMAR, and present test problems to examine both DR and

PR.

2. Surrogate assisted multi-objective evolutionary algorithms issues

The method of using the surrogate model is a relatively effective and delicate choice in solving expensive problems. The techniques commonly used in the surrogate model are RBF, PRS, SVM, neural network and the Kriging,... combined with machine learning, deep learning, clustering, statistics, forecasting. However, determining the frequency of using the surrogate function to replace the original function, how to choose a reference solutions,... will create adaptive noises. These noises are treated as transformations of the surrogate function instead of the original function. This may cause a decrease in the optimal properties of the algorithm. Therefore, the use of a surrogate model can create a noisy environment for the search process of MOEA. Based on formula (1), the interference using the surrogate model can be defined as follows:

$$F_{noise} = f'(\vec{x}) = f(\vec{x}) + e(\vec{x}) \quad (5)$$

Here, e is also a noise, which is a noisy function. In the surrogate model, the difference between surrogate function values and the original fitness function ones is used to calculate the noise's standard deviation and noise's distribution.

The research issue we need to solve is keeping the robustness for the algorithm during the usage of the surrogate model. Focusing on the performance of the algorithm through basic measurements such as GD, IGD [16], we need to keep the minimized values all the time of searching. This means we need to feasible front to be closed and spread along the Pareto optimal front in every generation.

IV. A GUIDANCE METHOD FOR SURROGATE ASSISTED ALGORITHMS

As hypothesized above, to reduce the noise of the environment, we need to maintain the population in the best of convergence and diversity during the search. Recently, many researchers discuss on the way of using the guidance method to drive the evolutionary process for the balance of exploitation and exploration of the process. It also maintains the convergence and diversity of the main population (obtained solutions) during the search continuously.

To build up the technique, we focus on two recent surrogate assisted multi-objective algorithms: K-RVEA [17] (using Kriging model) and CSEA [18] (using Feedforward Neural Networks - FNN). They are up to date surrogate assisted multi-objective algorithms, we will work on the algorithms with more investigation.

1. K-RVEA algorithm

Base on RVEA [19], a reference vector based MOEA for solving many-objective evolutionary problems (MaOPs), the authors proposed Kriging-assisted reference vector guided evolutionary algorithm (K-RVEA) [17]. In RVEA, search process is guided by a set of predefined reference vectors inspired from the direction vectors with the idea is to partition the objective space into a number small subspaces using a set of reference vectors. An elitism selection strategy is employed inside each subspace and using an Angle-Penalized Distance (APD) to measure the distance of the solutions to the ideal point and the closeness of the solutions to the reference vectors. An illustration of using reference vectors is shown in the Figure 1.

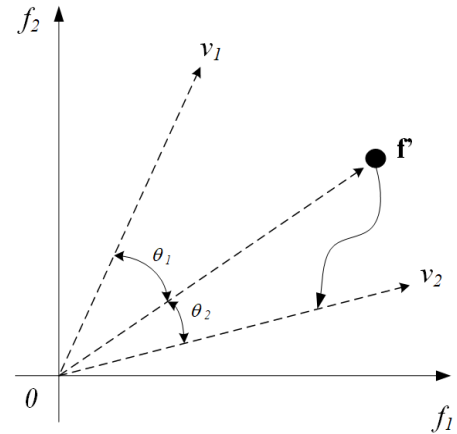


Figure 1. An example showing how to associate an individual with a reference vector. In this example, f' is a translated objective vector, v_1 and v_2 are two unit reference vectors. θ_1 and θ_2 are the angles between f' and v_1 , v_2 , respectively. Since $\theta_2 < \theta_1$, the individual denoted by f' is associated with reference vector v_2 .

In K-RVEA, based on RVEA, the authors introduced new algorithm which has three main phases: initialization, using and update the surrogate.

- 1) **Initialization:** The algorithm uses the Latin hypercube sampling to generate an initial population. Two archives A_1 and A_2 are filled with individuals which are evaluated with the original expensive functions. A Kriging model for each objective function is built by individuals which is given by the archive A_1 .
- 2) **Using the surrogate:** This phase, the K-RVEA uses Kriging models instead of the original functions to calculate objective function values. Kriging models are used for fitness evaluations for a prefixed number of generations without updating them.
- 3) **Updating the surrogate:** After an evolutionary search using the Kriging models for a fixed number of generations, the Kriging models will be updated. The selection of individuals to be re-evaluated using

the original functions, which will also be used for updating the surrogates, is essential for the performance of surrogate-assisted evolutionary algorithms (SAEAs). The authors suggested to use information from the underlying evolutionary algorithm, RVEA, and uncertainty information from the Kriging models for selecting individuals to be re-evaluated and then for re-training the surrogate. The selected individuals are then re-evaluated with the original functions and these data samples are added to both archives A_1 and A_2 . The authors also use a reference vector based selection strategy to keep a fixed number of individuals in the archive A_1 , it will eliminate extra individuals from A_1 . By the way to assign reference vectors in various clusters see Figure 2, then randomly select one data point from each cluster and eliminate the rest of the data. This technique helps the algorithm to maintain a fixed number of diverse set of training data in A_1 .

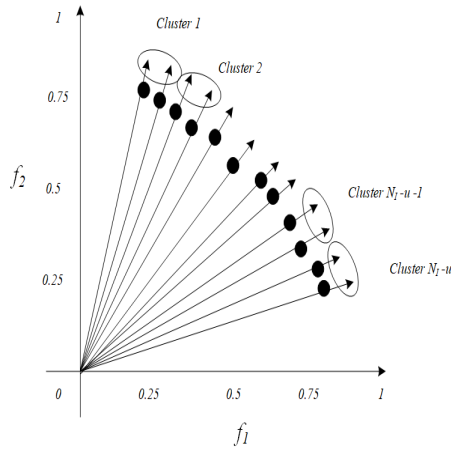


Figure 2. Clustering of active adaptive reference vectors V_a into a predefined number of clusters u .

2. CSEA algorithm

The authors in [18] proposed a classification based surrogate-assisted multi-objective algorithm for expensive many-objective optimization called CSEA. A classification criterion is proposed in CSEA to divide solutions evaluated using the expensive objective functions into two different categories. Then a surrogate is employed for learning the classification criterion to predict candidate solutions in categories and use an selection strategy to select better converged solutions for the evaluations. In CSEA, the primary features are the classification criterion and the surrogate management strategy.

- 1) **Classification criterion:** The authors introduced a classification accuracy heavily depends on the classi-

fication criterion. Based on two types of solutions in MOPs: dominated and non-dominated, a set of reference solutions are chosen to construct the Pareto dominance boundary, which will divide all solutions into two different categories.

- 2) **The selection of reference solution:** To choose a set of reference solutions to form the classification boundary, a radial space division based selection strategy. In this selection strategy, all solutions evaluated using the expensive fitness function are first projected into a 2-dimensional radial space.

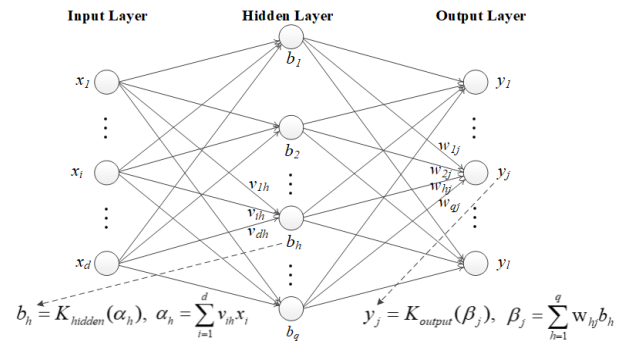


Figure 3. Illustration of a 3-layer FNN used in the CSEA algorithm.

To manage the surrogate model, the authors design the CSEA with four basic steps:

- **Initialization:** Initializing the parameters (K : number of reference solutions, $gmax$: number of solutions evaluated by surrogate model) and structure of the FNNs (which is one type of artificial neural networks in which connections between the neurons do not form a cycle). There are three main components of the FNN needed to be initialized: the network structure, the weights and the activation function.
- **Update:** Updating the weights of the FNN with the training data set. The authors used Levenberg-Marquardt back-propagation method to update the weights of the FNN.
- **Validation:** The authors used an error on test data as a measure to estimate the prediction uncertainty of the FNN. The errors are calculated by the Cross validation.
- **Surrogate assisted selection:** This step, some promising solutions are selected for the fitness evaluation. The selection is based on the predicted performance as well as the reliability. The selection used the errors on test data to estimate the reliability of the prediction of the FNN. A pair of errors (Test Error Rates of two categories) forms a point which is located in a region of the reliability configuration, the uncertainty of the FNN. The space presents the relationship between the uncertainty of the FNN and its test errors.

TABLE I
GD VALUES OF K-RVEA AND M-K-RVEA ON DTLZS BENCHMARK SET

MOEAs	E^1 w_{max}	2000	4000	6000	8000	10000	12000	14000	16000	18000	20000
DTLZ1											
K-RVEA	10	50.7441	7.4836	5.6692	4.1103	3.2282	2.7263	2.5096	2.4711	2.4463	2.3877
	20	41.9546	10.0818	10.9039	3.7888	3.7041	3.7041	4.0672	4.0672	4.0097	4.0097
	30	39.0192	9.4198	7.2260	5.4170	3.9909	3.9062	3.9777	4.0149	3.9320	3.7725
	50	42.5498	10.4695	6.7609	5.9252	5.5349	4.6894	4.7492	4.3930	4.1622	4.0916
	70	43.5096	12.6410	8.0107	6.6715	7.1815	6.2105	5.7701	5.2643	5.0035	4.7697
M-K-RVEA	ADC	46.6079	4.7267	3.3725	3.3725	3.3725	3.3725	3.3725	3.3725	3.3725	3.3725
DTLZ2											
K-RVEA	10	0.1032	0.0023	0.0014	0.0011	0.0009	0.0008	0.0008	0.0007	0.0007	0.0006
	20	0.1151	0.0014	0.0009	0.0007	0.0006	0.0005	0.0005	0.0005	0.0004	0.0004
	30	0.0999	0.0015	0.0010	0.0008	0.0006	0.0005	0.0005	0.0004	0.0004	0.0004
	50	0.1076	0.0015	0.0009	0.0006	0.0005	0.0005	0.0004	0.0004	0.0004	0.0004
	70	0.1165	0.0014	0.0010	0.0007	0.0006	0.0005	0.0005	0.0004	0.0004	0.0004
M-K-RVEA	ADC	0.1227	0.0019	0.0011	0.0008	0.0007	0.0006	0.0006	0.0005	0.0005	0.0004
DTLZ3											
K-RVEA	10	183.1624	43.1944	36.1737	33.9130	30.8676	29.8677	28.4940	30.4472	29.3890	30.4415
	20	171.6720	65.1569	57.1465	47.8764	47.3865	42.2447	40.1520	38.8519	36.7741	36.5496
	30	220.4797	74.7820	56.8731	52.7536	52.7041	48.8366	44.5964	42.4959	41.5517	40.6848
	50	194.0912	85.5257	83.2413	70.2972	59.2193	50.7384	50.3928	49.8387	48.0394	44.5380
	70	151.5813	77.2795	68.4175	59.9071	51.6176	47.8923	49.5671	51.3997	48.6738	48.7996
M-K-RVEA	ADC	183.1624	69.5647	54.4809	48.0276	47.8869	43.2640	42.0306	40.3582	39.2478	42.3035
DTLZ4											
K-RVEA	10	0.1739	0.0148	0.0093	0.0063	0.0051	0.0043	0.0038	0.0034	0.0031	0.0029
	20	0.1613	0.0113	0.0072	0.0054	0.0044	0.0039	0.0035	0.0033	0.0031	0.0029
	30	0.1634	0.0108	0.0077	0.0060	0.0051	0.0044	0.0038	0.0035	0.0033	0.0031
	50	0.1841	0.0115	0.0074	0.0057	0.0048	0.0042	0.0038	0.0034	0.0031	0.0028
	70	0.1564	0.0129	0.0085	0.0073	0.0063	0.0055	0.0050	0.0047	0.0045	0.0042
M-K-RVEA	ADC	0.1564	0.0129	0.0085	0.0073	0.0063	0.0055	0.0050	0.0047	0.0045	0.0042
DTLZ5											
K-RVEA	10	0.1407	0.0068	0.0057	0.0054	0.0050	0.0046	0.0042	0.0042	0.0042	0.0041
	20	0.1571	0.0096	0.0069	0.0063	0.0056	0.0053	0.0052	0.0050	0.0049	0.0048
	30	0.1207	0.0111	0.0082	0.0068	0.0063	0.0058	0.0056	0.0054	0.0051	0.0051
	50	0.1584	0.0052	0.0046	0.0042	0.0041	0.0041	0.0040	0.0040	0.0039	0.0038
	70	0.1368	0.0101	0.0072	0.0065	0.0062	0.0056	0.0055	0.0052	0.0050	0.0049
M-K-RVEA	ADC	0.1286	0.0071	0.0055	0.0048	0.0045	0.0042	0.0039	0.0039	0.0037	0.0037
DTLZ6											
K-RVEA	10	0.9746	0.4591	0.4409	0.3822	0.3919	0.3913	0.3306	0.3261	0.3428	0.3140
	20	0.9513	0.4647	0.4897	0.4926	0.4734	0.4403	0.4151	0.3610	0.3625	0.3456
	30	0.9965	0.4503	0.3943	0.4227	0.3993	0.3845	0.3798	0.3657	0.3867	0.3529
	50	0.9480	0.5927	0.5124	0.4898	0.4478	0.4311	0.4444	0.4308	0.4108	0.3957
	70	0.9599	0.5421	0.4830	0.4446	0.4565	0.4716	0.4370	0.4339	0.4248	0.4184
M-K-RVEA	ADC	1.0474	0.4951	0.5234	0.5070	0.4695	0.4490	0.4279	0.4099	0.3823	0.3597
DTLZ7											
K-RVEA	10	2.0241	0.0026	0.0012	0.0008	0.0006	0.0005	0.0005	0.0004	0.0004	0.0003
	20	2.6543	0.0012	0.0007	0.0005	0.0004	0.0003	0.0003	0.0003	0.0002	0.0002
	30	2.6896	0.0009	0.0005	0.0004	0.0003	0.0003	0.0003	0.0002	0.0002	0.0002
	50	2.3290	0.0010	0.0007	0.0005	0.0004	0.0004	0.0003	0.0003	0.0002	2.0384
	70	2.0943	0.0007	0.0005	0.0005	0.0004	0.0003	0.0003	0.0003	0.0002	0.0002
M-K-RVEA	ADC	2.0241	0.0009	0.0007	0.0004	0.0004	0.0003	0.0003	0.0003	0.0002	0.0002
DTLZ8											
K-RVEA	10	0.0548	0.0322	0.0322	0.0322	0.0322	0.0322	0.0322	0.0322	0.0322	0.0322
	20	0.0587	0.0327	0.0327	0.0327	0.0327	0.0327	0.0327	0.0327	0.0327	0.0327
	30	0.0706	0.0358	0.0358	0.0358	0.0358	0.0358	0.0358	0.0358	0.0358	0.0358
	50	0.0728	0.0303	0.0303	0.0303	0.0303	0.0303	0.0303	0.0303	0.0303	0.0303
	70	0.0635	0.0357	0.0357	0.0357	0.0357	0.0357	0.0357	0.0357	0.0357	0.0357
M-K-RVEA	ADC	0.0635	0.0310	0.0310	0.0310	0.0310	0.0310	0.0310	0.0310	0.0310	0.0310
DTLZ9											
K-RVEA	10	3.4231	2.2998	1.5890	1.3484	1.4657	2.2714	1.4092	1.2228	1.1880	1.1553
	20	6.3959	2.5242	2.3244	2.1598	1.9384	1.9384	2.3200	1.5506	1.6016	2.1090
	30	4.2762	2.1689	1.9599	1.7168	1.6040	1.4022	1.5312	1.4316	1.3738	1.2709
	50	3.9907	1.5387	1.1807	0.9758	2.3940	2.1892	1.8939	1.8939	1.8430	1.7425
	70	4.5920	2.5457	1.9841	2.2016	2.0792	2.0511	1.9588	1.8728	1.9472	1.7881
M-K-RVEA	ADC	3.7733	2.5289	2.2851	2.1059	2.0260	2.0260	1.8907	1.8280	1.2761	1.1746

In the CSEA algorithm, there are two loops. The main loop represents the evolution using the original objective function. In the second loop, the solutions are selected through the classification of the surrogate model.

The experiments in [18] confirmed the author's hypoth-

eses that surrogates classifying solutions into dominated and non-dominated solutions are promising for solving expensive many-objective optimization problems. CSEA is an up to date surrogate assisted multi-objective algorithm, so we build up our work on the algorithm with more investigation.

¹ Evaluations

TABLE II
IGD VALUES OF K-RVEA AND M-K-RVEA ON DTLZS BENCHMARK SET

MOEAs	w_{max}	E^1	2000	4000	6000	8000	10000	12000	14000	16000	18000	20000
DTLZ1												
K-RVEA	10		37.5017	15.9551	15.9551	15.9551	8.6306	8.6306	8.6306	6.4437	6.4437	6.4437
	20		69.2926	12.2018	9.0805	6.0770	6.0770	6.0770	6.0770	6.0770	6.0770	6.0770
	30		73.5400	18.0276	18.0276	18.0276	18.0276	18.0276	18.0276	18.0276	18.0276	18.0276
	50		75.3314	21.8503	21.8503	21.8503	18.8011	18.8011	18.8011	18.8011	18.8011	18.8011
	70		56.6293	17.9351	17.9351	17.9351	15.7866	15.7866	11.9698	11.9698	11.9698	11.9698
M-K-RVEA	ADC		30.8156	9.9032	9.9032	9.9032	9.9032	9.9032	9.9032	9.9032	9.9032	9.9032
DTLZ2												
K-RVEA	10		0.4863	0.0523	0.0417	0.0372	0.0343	0.0326	0.0314	0.0302	0.0293	0.0284
	20		0.4862	0.0461	0.0368	0.0325	0.0297	0.0287	0.0281	0.0273	0.0265	0.0259
	30		0.4720	0.0467	0.0382	0.0326	0.0293	0.0283	0.0268	0.0260	0.0249	0.0244
	50		0.5233	0.0476	0.0387	0.0340	0.0306	0.0295	0.0285	0.0276	0.0272	0.0267
	70		0.4848	0.0502	0.0427	0.0372	0.0338	0.0324	0.0302	0.0292	0.0283	0.0272
M-K-RVEA	ADC		0.5617	0.0528	0.0398	0.0354	0.0335	0.0320	0.0305	0.0295	0.0287	0.0278
DTLZ3												
K-RVEA	10		500.9795	154.0806	154.0806	154.0806	154.0806	154.0806	154.0806	154.0806	154.0806	154.0806
	20		608.2876	195.9133	185.9003	185.9003	185.9003	185.9003	185.9003	177.8553	177.8553	177.8553
	30		421.7119	282.9219	229.2645	229.2645	229.2645	222.9678	132.3652	132.3652	132.3652	132.3652
	50		533.7593	229.4352	225.6308	225.6308	196.0629	196.0629	196.0629	196.0629	196.0629	195.8602
	70		672.9475	281.1927	274.5212	246.6838	246.6838	183.5229	183.5229	183.5229	183.5229	183.5229
M-K-RVEA	ADC		500.9795	235.2239	204.4000	204.4000	204.4000	182.6566	182.6566	182.6566	182.6566	182.6566
DTLZ4												
K-RVEA	10		0.9201	0.1337	0.0794	0.0595	0.0549	0.0495	0.0477	0.0458	0.0444	0.0434
	20		0.9733	0.0753	0.0606	0.0543	0.0503	0.0465	0.0447	0.0418	0.0408	0.0396
	30		0.8311	0.0745	0.0626	0.0568	0.0520	0.0466	0.0441	0.0421	0.0408	0.0399
	50		0.8441	0.0786	0.0637	0.0577	0.0541	0.0512	0.0474	0.0442	0.0420	0.0413
	70		0.8792	0.0833	0.0662	0.0629	0.0582	0.0552	0.0542	0.0524	0.0512	0.0504
M-K-RVEA	ADC		0.8792	0.0833	0.0662	0.0629	0.0582	0.0552	0.0542	0.0524	0.0512	0.0504
DTLZ5												
K-RVEA	10		0.3699	0.0395	0.0397	0.0384	0.0343	0.0330	0.0329	0.0327	0.0301	0.0301
	20		0.3739	0.0583	0.0525	0.0391	0.0371	0.0357	0.0355	0.0347	0.0345	0.0344
	30		0.4081	0.0722	0.0597	0.0496	0.0492	0.0478	0.0457	0.0449	0.0421	0.0421
	50		0.3483	0.0381	0.0337	0.0317	0.0307	0.0367	0.0364	0.0364	0.0364	0.0363
	70		0.3860	0.0494	0.0426	0.0412	0.0411	0.0394	0.0393	0.0388	0.0384	0.0376
M-K-RVEA	ADC		0.4247	0.0421	0.0371	0.0380	0.0373	0.0345	0.0338	0.0336	0.0333	0.0333
DTLZ6												
K-RVEA	10		8.2380	1.8545	0.8540	0.8110	0.8110	0.8110	0.8110	0.8110	0.7906	0.7724
	20		8.4386	2.2229	1.3579	1.3532	1.3532	1.2986	1.2986	1.2557	1.2557	1.2557
	30		8.2757	2.4312	2.4304	1.8914	1.7946	1.7637	1.6556	1.6556	1.6556	1.6437
	50		8.2283	2.7022	2.5015	2.4640	2.4640	2.4640	2.2781	2.2781	2.0575	2.0575
	70		8.4906	2.4451	2.4451	2.4451	2.1171	2.0384	2.0384	2.0384	2.0384	2.0384
M-K-RVEA	ADC		8.4051	2.9249	1.8047	1.6488	1.6488	1.6488	1.6488	1.6334	1.5731	1.5731
DTLZ7												
K-RVEA	10		8.1067	0.0611	0.0397	0.0332	0.0295	0.0276	0.0256	0.0243	0.0231	0.0222
	20		7.4668	0.0492	0.0406	0.0338	0.0317	0.0296	0.0280	0.0268	0.0256	0.0246
	30		8.6059	0.0514	0.0424	0.0380	0.0358	0.0341	0.0312	0.0303	0.0293	0.0282
	50		8.8847	0.0707	0.0599	0.0543	0.0469	0.0440	0.0420	0.0407	0.0398	0.0477
	70		8.2937	0.0725	0.0598	0.0559	0.0506	0.0487	0.0459	0.0442	0.0433	0.0424
M-K-RVEA	ADC		8.1067	0.0618	0.0546	0.0456	0.0433	0.0411	0.0359	0.0343	0.0320	0.0312
DTLZ8												
K-RVEA	10		0.2309	0.2050	0.2050	0.2050	0.2050	0.2050	0.2050	0.2050	0.2050	0.2050
	20		0.2590	0.1599	0.1599	0.1599	0.1599	0.1599	0.1599	0.1599	0.1599	0.1599
	30		0.2190	0.1827	0.1827	0.1827	0.1827	0.1827	0.1827	0.1827	0.1827	0.1827
	50		0.2516	0.2357	0.2357	0.2357	0.2357	0.2357	0.2357	0.2357	0.2357	0.2357
	70		0.2397	0.1935	0.1935	0.1935	0.1935	0.1935	0.1935	0.1935	0.1935	0.1935
M-K-RVEA	ADC		0.2397	0.1790	0.1790	0.1790	0.1790	0.1790	0.1790	0.1790	0.1790	0.1790
DTLZ9												
K-RVEA	10		11.3847	6.3491	6.1299	6.1299	5.6413	5.3711	4.9843	4.9843	4.9843	4.9843
	20		11.1618	6.3076	6.2910	6.2910	6.2910	6.2910	5.9713	5.9713	5.9713	5.9066
	30		11.1510	7.7752	7.7752	7.7563	7.7378	7.7212	7.7103	7.7103	7.7103	7.7099
	50		11.2579	7.9947	7.9363	7.9294	7.5453	7.5453	7.5453	7.5453	7.5453	7.5453
	70		11.1964	6.6274	6.3928	6.3837	6.3837	6.3837	6.3837	6.3837	6.3837	6.3830
M-K-RVEA	ADC		11.2529	6.5146	6.4943	6.4943	6.4943	6.4943	6.4943	6.4943	6.4697	6.4697

3. Noisy issues

In our work, we analysed the evolutionary process's behavior during the search with surrogate-assisted multi-objective evolutionary algorithms.

K-RVEA: For MOEA in general, according to the searching process, the quality of the achieved solution pop-

ulation varies according to the general rule, that is: in the first generations, the convergence quality of the algorithm is poor, the calculation quite diverse due to the fact that random generated solutions have not really evolved over generations. After that, the solution quality gradually got better, the diversity also gradually decreased and towards

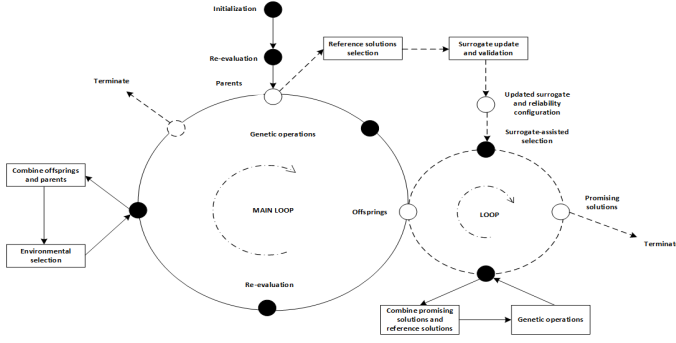


Figure 4. Demonstrate the main loop of evaluation using the original objective function and sub-loop of the selected solutions from the surrogate model.

the Pareto optimal front spread. The last generations of the process with a specified number of generations, the quality of the population converged was good because the process over many generations. An algorithm are assessed well if the diversity of the population is maintained when the solutions spread evenly over the Pareto optimization front. In many analyzes of the evolutionary trend of the process, in order to ensure a quality balance between convergence and diversity, it is necessary to guide the algorithm to keep a balance between exploration and exploitation cascade of evolution. Therefore, in the early stages, the ability to converge and the end stage need to enhance the diversity of the population. In K-RVEA, the authors use parameters to determine whether the model update is fixed, which can lack the algorithm's adaptability, possibly making noises for environment which is against the robustness of the algorithm. We will test the K-RVEA algorithm on the class of expensive problems, with different parameters to evaluate the above comments.

We used w_{max} , the number of generations before updating the Kriging model with values: 10, 20, 30, 50 and 70 with DTLZs problems [20] (DTLZs from 1 to 9) in 20000 evaluations. We reported the results on GD and IGD, the popular metrics to measure the performance of the algorithm. The results are shown in Tables I, II. Some of the behaviors of the evolution process are reported in Figures 5 to 10.

CSEA: The strategy to select reference solutions is an important factor for the performance of the algorithm. The space division based on a boundary that is formed from the selected solutions. There is a key parameter for CSEA which is K : number of reference solutions that are selected. The role of K is discussed on the proposal. When value of K is small, the number of reference solutions of category II solutions is small, the algorithm has better convergence performance but is poor in terms of diversity. By contrast, that has better diversity but is poorly converged with larger

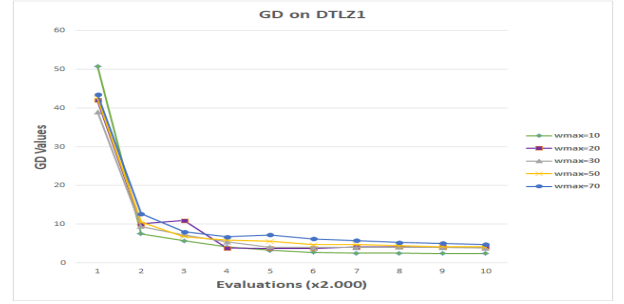


Figure 5. The behavior of K-RVEA on DTLZ1 (GD)

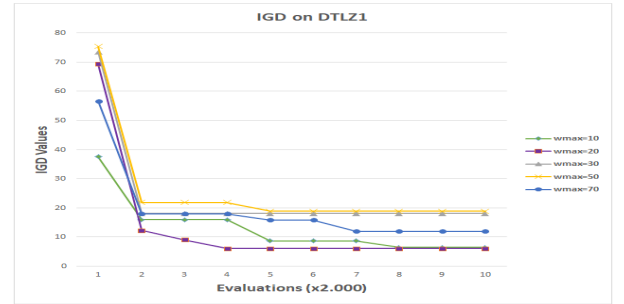


Figure 6. The behavior of K-RVEA on DTLZ1 (IGD)

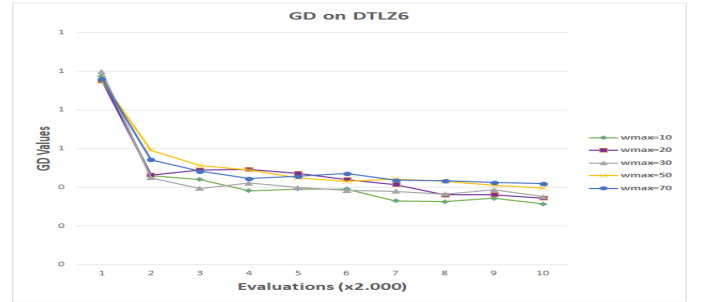


Figure 7. The behavior of K-RVEA on DTLZ6 (GD)

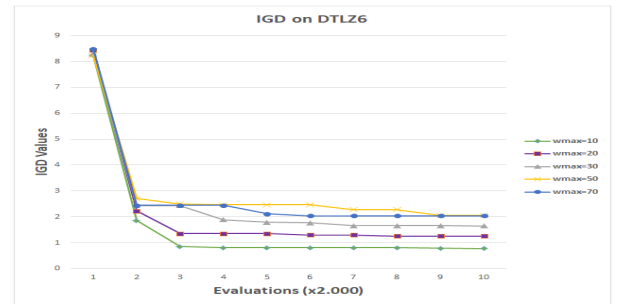


Figure 8. The behavior of K-RVEA on DTLZ6 (IGD)

value of K . The choice of the number of reference solutions provides a way of balancing the convergence and diversity of the selected category II solutions.

The authors also empirically investigate the impact of the number of reference solutions K on the performance

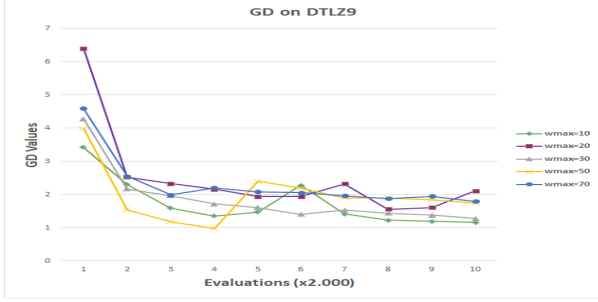


Figure 9. The behavior of K-RVEA on DTLZ9 (GD)

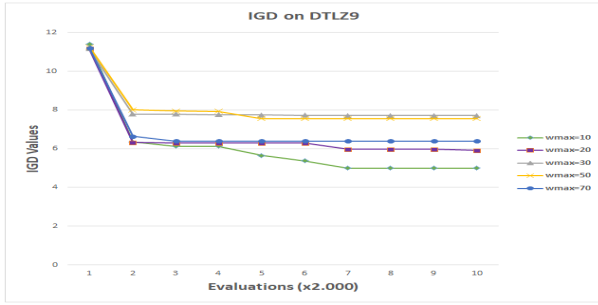


Figure 10. The behavior of K-RVEA on DTLZ9 (IGD)

of CSEA for solving problems with different numbers of objectives. CSEA with different settings of parameter K is tested on different benchmark sets. Through the results, the best choice of K is recorded to set for the algorithm. Then, the parameter is fixed to six for all the benchmark sets, it is known as a practical parameter value.

Similar to analysis of the evolutionary process's behavior in K-RVEA, the convergence quality in the early stages and the diversity of the population in the end stage need to be enhanced. In CSEA, the authors use K parameter to select reference solutions to form the classification boundary, a radial space division based on this selection strategy. Because the value of K is not adaptive, so it might lack the algorithm's adaptability, possibly making the noises for environment which is against the robustness of the algorithm.

4. Guidance techniques

a) An adaptive control

Based on the analysis from the above experimental results with K-RVEA, for reducing the noise, the assumptions about appropriate parameter adjustment during evolution create an effective balance between convergence quality and diversity of the population. Specifically, adjusting the parameter on the number of generations before updating the model according to the assessed factors and assuming that the direct influence in SAEAs is: the processing time

of the evolution, the quality of the current population, and the computational characteristics of the problem.

Building the rate of using surrogate function and frequency of updating Kriging model flexibly, depending on the quality of the population in the current generation, to determine the quality of the solution, we use two defined values at a time t , with the test problem j as follows:

$$Q_t = \frac{n_{ND} * n_{gen}}{q_{size} * t_{size}} C_j \quad (6)$$

In which, n_{ND} is the number of non-dominated solutions in the main population, q_{size} is the size of the population, current generation n_{gen} , t_{size} is the number of generations. C_j is the complexity of the problem, this is a parameter dependent on the complexity of the test function, it depends on the size of decision spaces, number of objectives and number of calculations for objective functions, we suggest to choose practical values $C_j \in [0.8, 1]$, the default value is 1. The constant 60 is used, which is the upper range of w_{max} , it helps the control to get the maximized value, 50. Based on the Q_t value, we will calculate the value of $w_{max} \in R$ (the rate of using surrogate function and the frequency of updating the Kriging model) in the range [10, 50], a normalize function is defined to adjust valid w_{max} values.

$$w_{max} = \text{normalize}(Q_t * 60) \quad (7)$$

The modified version M-K-RVEA is described in Algorithm 1.

b) An adaptive number of reference solutions

Analyzing from the experimental results with CSEA, we hypothesize that, in order to reduce the noise by keeping balance the evolution between exploration and mining as well as to maintain the convergence and diversity of the population, the selection of a reference solution requires adaptability. It depends on factors that we can identify with the following influences:

- The process of evolution: evolution takes place in the process from the initial state of creation, the central stage and the final stage, the quality of the main population (number of non-dominated solutions) at each stage are different. Therefore, the use of reference solutions should be adjusted according to the quality of the current population in terms of convergence, diversity. In other words, at each time, it is necessary to select the number of different reference solutions to create a balance between exploration and exploitation, convergence and diversity.
- The sample problems are different, the quality of the population in each stage has different development, so at each specific time, based on the quality of the population to choose the appropriate number of reference

TABLE III
 DTLZS PROBLEMS FOR EXPERIMENTS

Name	Specific characteristics
DTLZ1	The Pareto front is linear (Hyper-plane). There are $(11^k - 1)$ local optimal fronts where k is a user-specified parameter
DTLZ2	For $M > 3$, the Pareto optimal solutions lie inside the first quadrant of the unit sphere in a three-objective plot with f_M as one of the axes
DTLZ3	There are $(3^K - 1)$ local fronts that are parallel to the global Pareto front where k is a user-specified parameter
DTLZ4	The Pareto optimal solutions are non-uniformly distributed along the Pareto front
DTLZ5	The front is a curve and the Pareto optimal solutions are non-uniformly distributed along the Pareto front
DTLZ6	A more difficult version of the DTLZ5 problem: the non-linear distance function g makes it harder to convergence against the Pareto optimal curve
DTLZ7	The Pareto front is formed by 2^{M-1} disjoint regions in the objective space
DTLZ8	The Pareto front is a combination of a straight line and a hyper-plane. The straight line is the intersection of the first $(M - 1)$ constraints with $f_1 = f_2 = \dots = f_{M-1}$ and the hyper-plane is represented by another constraint g_M
DTLZ9	The Pareto front is a curve with $f_1 = f_2 = \dots = f_{M-1}$. The solution density gets thinner towards the Pareto front

solutions. It also makes the algorithm highly adaptive in maintaining the balanced convergence quality and diversity of the population.

- The difficulty of the problems is different, making predictions about the difficulty of the problem, deciding on the reference solution from users is also an important factor to guide the evolution flexibly in maintaining a balance between exploration and mining. This can both help the problem create leaps that can overcome local points, and the combination of decision makers and evolution.

Based on the analysis from the above experimental results, the assumptions about appropriate parameter adjustment during evolution is to create an effective balance between convergence quality and diversity of the population. Specifically, adjusting the parameter on the number of generations before updating the model according to the assessed factors and assuming that the direct influence in algorithms using a surrogate function is: the process time of the process evolution, the quality of the current population and the computational characteristics of the problem.

Designing a dynamic selection strategy to select reference solutions to form the boundary for the classification, we take into account the influence of the quality of the solution population in the current generation. Similar to the adaptive control technique for K-RVEA, to determine the quality of the solution, we calculate Q_t value by using equation 6. Based on the Q_t value, we will calculate the value of $K \in R$ (the number of reference solutions will be selected to form the boundary) in the range [3, 9], a normalize function is defined to adjust valid K values. We use constant 6 as a generated point to control the K in the above range.

$$K = \text{normalize}(Q_t * 6) \quad (8)$$

The modified version M-CSEA is described in Algorithm 2.

V. EXPERIMENTAL STUDIES

1. Testing problems and parameter settings

In the experiments, we used the origin K-RVEA and the modified version M-K-RVEA with the usage of the above adaptive w_{max} , and the origin CSEA and the modified version called M-CSEA with the usage of the above dynamic K . We also used some selected expensive problems in DTLZs [20] benchmark set. In the benchmark set, we have multiple kinds of expensive problems, which can be described in Table III.

In our experiments, the parameters are: the distribution index of crossover $n_c = 20$ and the distribution index of mutation $n_m = 20$, the crossover probability $pc = 1.0$ and the mutation probability $p_m = 1/d$, where d is the number of decision variables. Each case study, we run 30 independent times for analysing experimental results.

2. Performance metrics

Performance metrics are usually used to compare algorithms in order to form an understanding of which algorithm is better and in what aspects. However, it is hard to define a concise definition of algorithmic performance. In general, when doing comparisons, a number of criteria are employed [16]. Since each metric has both advantage and disadvantage, we will look at all two popular criteria: the generational distance (GD) and the inverse generational distance (IGD).

The GD measure is defined as the average distance from a set of solutions, denoted P , found by evolution to the global Pareto optimal set (POS) [21]. The first-norm equation is defined as:

$$GD = \frac{\sum_{i=1}^n d_i}{n} \quad (9)$$

where d_i is the Euclidean distance (in objective space) from solution i to the nearest solution in the POS, and n is the size of P . This measure is considered for convergence

Algorithm 1: M-K-RVEA

```

1 Input:
2 The maximum number of expensive function
  evaluations:  $FE_{max}$ ; the number of individuals to
  be re-evaluated and to be used for updating
  Kriging models is  $u$ ; the number of generations
  before updating Kriging models:  $w_{max}$ .
3 Output:
4 Set of non-dominated solutions which are
  evaluated from  $A_2$  set.
5 /* Initialization */
  1) Initialize a new population of size  $N_I$  with Latin
    hypercube sampling, set the number of function
    evaluations  $FE = 0$ , initialize the counters: the
    generation counter for using Kriging models  $w = 1$ ,
    the number of updates  $t_u = 0$ . Set archives  $A_1$  and
     $A_2$  to be empty,  $|A_1| = |A_2| = \emptyset$ 
  2) Evaluate the initial population with the original
    objective functions and add them to  $A_1$  and  $A_2$ ,
    update  $FE = FE + N_I$ ,  $|A_1| = |A_1| + N_I$  and
     $|A_2| = |A_2| + N_I$ 
  3) Using archive  $A_1$  to train a Kriging model for each
    objective function.
  4) while  $FE < FE_{max}$ 
    /* Using Kriging models */
  5) Calculate  $Q_t$  by equation 6 then re-calculate  $w_{max}$ 
    by equation 7
  6) while  $w \leq w_{max}$ 
  7) Run RVEA steps
    • Generate offspring population
    • Combine parent and offspring populations
    • Select parents from the population combined for
      the next generation.
    • Update reference vectors
    • Update  $w = w + 1$ 
  8) end while
    /* Updating Kriging models */
  9) Select  $u$  individuals according to either the amount
    of uncertainty or the value of angle penalized
    distance (APD) [17] and re-evaluate them with the
    original objective functions and update
     $FE = FE + u$ . Add these individuals to two archives
     $A_1$ ,  $A_2$  and update  $|A_1| = |A_1| + u$  and
     $|A_2| = |A_2| + u$ 
  10) Remove  $|A_1| - N_I$  individuals from  $A_1$  with
    specified concept of managing training data archive
    [17], update  $w = 1$  and  $t_u = t_u + 1$  and go to step 3.
  11) end while

```

Algorithm 2: M-CSEA

```

1 Input:
2 Population size:  $N$ ; the number of reference
  solutions:  $K$ ; the number of hidden neurons:  $H$ ;
  the maximum number of fitness evaluations (FEs):
   $t_{max}$ ; the maximum number of surrogate
  predictions before being updated:  $g_{max}$ 
3 Output:
4 Set of non-dominated solutions from population  $P$ .
5 /* Initialization */
  1) Initialize  $P$  with  $11d - 1$  solutions using Latin
    hypercube sampling method.
  2) Set  $t = 11d - 1$  with number of decision variables  $d$ .
  3) Initialize the neural network with  $H$  hidden neurons
    with  $m$  is the number of objectives
  4) Copy individuals from  $P$  to the archive  $Arc$ .
  5) while  $t \leq t_{max}$ 
  6) Calculate  $Q_t$  by equation 6 then re-calculate  $K$  by
    equation 8
  7) Get  $K$  reference individuals from  $P$  to population
     $P_R$ 
    /* Updating, validating the surrogate model */
  8) Classifying the combination of  $P_R$  and  $Arc$  into
    two categories  $C$  ( $C_1$  and  $C_2$ ) by Classify()
    procedure is presented in [18]
  9) Rate each solutions in  $C_2$  and store in  $rr$ 
  10) Calculate  $tr = \min\{rr, 1 - rr\}$ 
  11) Partitioning the combination of  $Arc$  and  $C$  into
     $D_{train}$  and  $D_{test}$  with DataPartition() procedure
    in [18]
  12) Training neural network net by Train() procedure
    [18] with  $D_{train}$  and  $T$ 
  13) Validation net with  $D_{test}$  into  $[p_1, p_2]$ 
    /* Using the surrogate model */
  14) Using a surrogate-assisted selection strategy (with
    SASselection( $P, P_R, p_1, p_2, g_{max}, tr$ ) procedure in
    [18]) to  $Q$  population
  15) Combine the archive  $Arc$  and  $Q$  to  $Arc$ 
  16) Using a radial projection based selection (with
    RadialSelection( $P \cup Q, N$ ) procedure in [18]) to
    select  $N$  individuals to  $P$ .
  17) Set  $t = t + |Q|$ 
  18) end while

```

aspect of performance. Therefore, it could happen that the set of solutions is very close to the Pareto optimal front (POF), but it does not cover the entire of the POF.

The measure IGD takes into account both convergence and spread to all parts of the POS. The first-norm equation

for IGD is as follows:

$$IGD = \frac{\sum_{i=1}^{\bar{N}} \bar{d}_i}{\bar{N}} \quad (10)$$

where $\bar{d}_i$ is the Euclidean distance (in objective space) from solution i in the POS to the nearest solution in P , and $\bar{N}$ is the size of the POS. In order to get a good value for IGD (ideally zero), P needs to cover all parts of the POS. However, this method only focuses on the solution that is closest to the solution in the POS indicating that a solution in P might not take part in this calculation.

3. Results and Analysis

a) M-K-RVEA

In order to analyze the performance of M-K-RVEA, with maximum of 20000 evaluations in 30 independent runs, we reported the results of GD, IGD measurement in rows ADC on Tables I, II. We also reported and compared the behaviors of K-RVEA (with the default w_{max} is 30) and M-K-RVEA.

Analyzing and comparing the results, we can found that M-K-RVEA with adaptive parameter control can get more balance of convergence and diversity for obtained population. In details, in the early stages, it gets better in 6 problems on GD and 6 problems on IGD. During the generations (the later stages), M-K-RVEA gets equal or better on 5 problems, one is the same. It is also better in 6 problems on IGD. In the last stages, M-K-RVEA is gets better in 7 problems on GD and 5 problems in IGD. In other problems, K-RVEA gets the same results with M-K-RVEA or a bit better than M-K-RVEA.

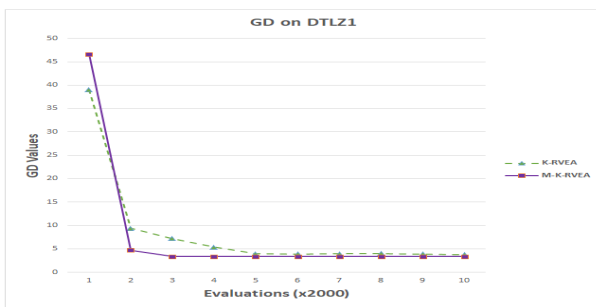


Figure 11. The behavior of M-K-RVEA and K-RVEA on DTLZ1 (GD)

In general, through the experimental results, it is proposed to use adaptive parameter control that has an effective effect to implement the optimal target optimization algorithm using general surrogate functions, algorithms based on Kriging model in particular. The obvious meaning of using the response parameter is to assess the quality of the population at the time of determining whether

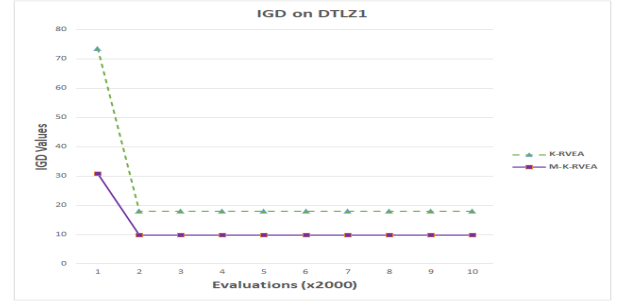


Figure 12. The behavior of M-K-RVEA and K-RVEA on DTLZ1 (IGD)

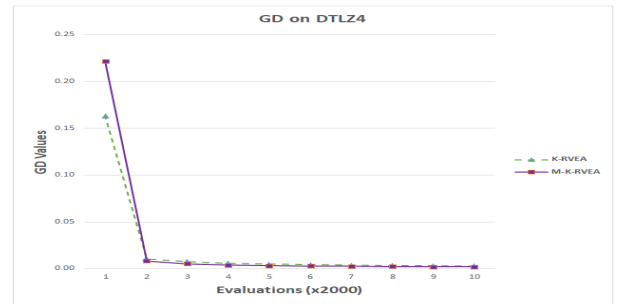


Figure 13. The behavior of M-K-RVEA and K-RVEA on DTLZ4 (GD)

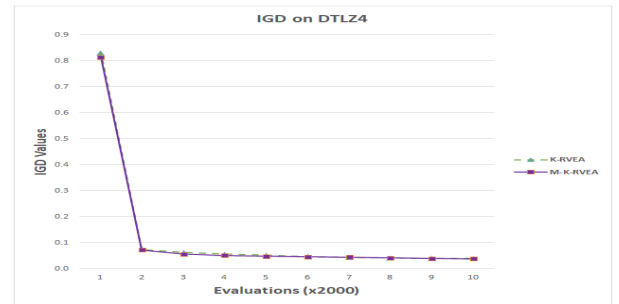


Figure 14. The behavior of M-K-RVEA and K-RVEA on DTLZ4 (IGD)

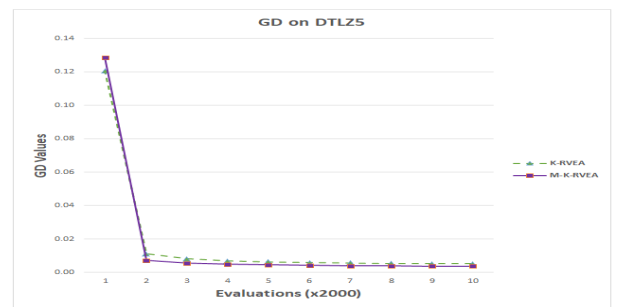


Figure 15. The behavior of M-K-RVEA and K-RVEA on DTLZ5 (GD)

to update the Kriging model or not? The results using adaptive parameter control demonstrate the hypothesis of the impact of the time process, the current convergence quality of the population, and the characteristics of the problem on maintaining the equilibrium of convergence

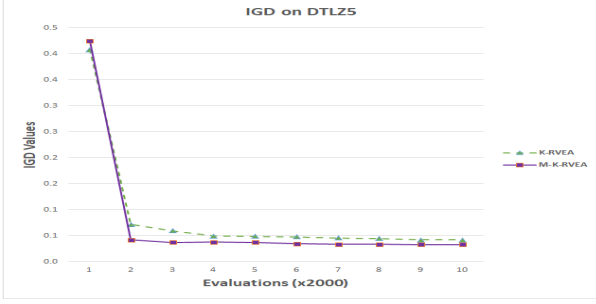


Figure 16. The behavior of M-K-RVEA and K-RVEA on DTLZ5 (IGD)

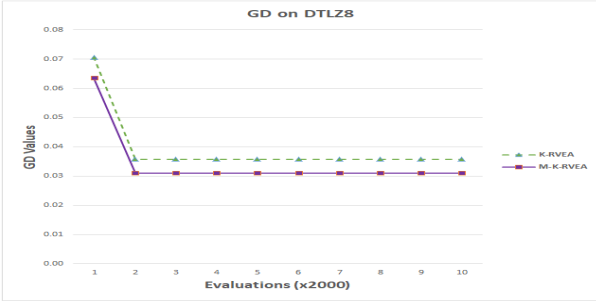


Figure 17. The behavior of M-K-RVEA and K-RVEA on DTLZ8 (GD)

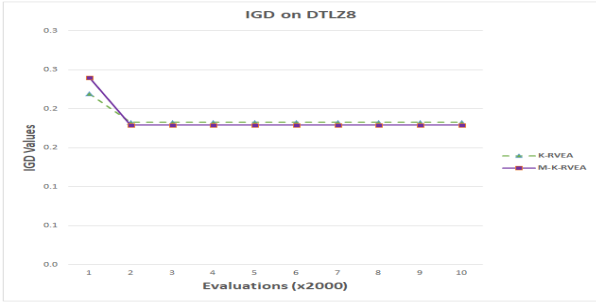


Figure 18. The behavior of M-K-RVEA and K-RVEA on DTLZ8 (IGD)

quality and diversity of the population, through ensuring a balance between the exploration and exploitation process of evolution. The effectiveness of parameter selection in the current stages also depends on the quality and effectiveness of the previous stage, proving that a timely adjustment to maintain a balance between exploration and exploitation is needed to set. However, from the results of some problems, K-RVEA has a better quality, although not large, it also gives us issues that need to be adequately posed such as: stability, sustainability, random factors course, etc. These factors influence the choice of parameters appropriately, or guide evolution to better efficiency, it will reduce noises to improve the robustness of the algorithm.

b) M-CSEA

The results for the experiments are shown in Tables IV, V. We also reported and compared the behaviors

of CSEA (with the default K is 6) and M-CSEA, some of screenshots are shown in Figures: 19 to 26.

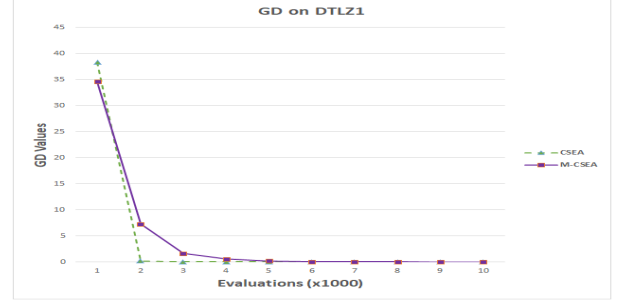


Figure 19. The behavior of M-CSEA and CSEA on DTLZ1 (GD)

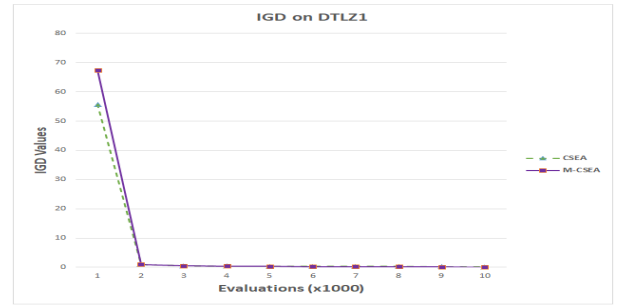


Figure 20. The behavior of M-CSEA and CSEA on DTLZ1 (IGD)

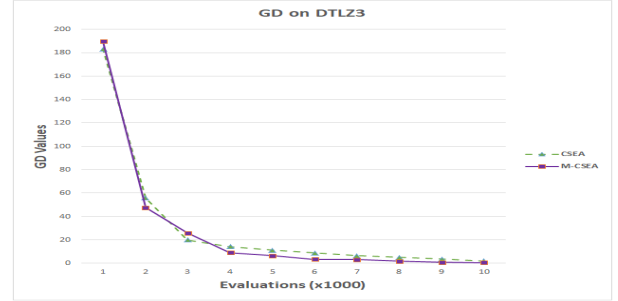


Figure 21. The behavior of M-CSEA and CSEA on DTLZ3(GD)

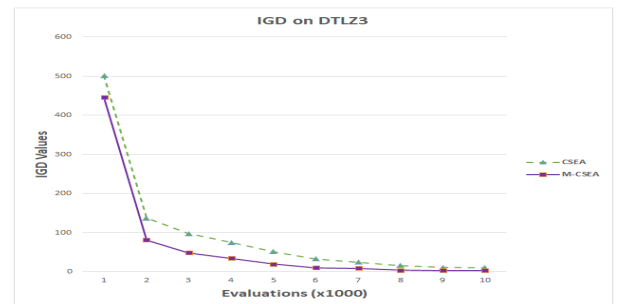


Figure 22. The behavior of M-CSEA and CSEA on DTLZ3 (IGD)

Analysing the results, with maximum of 10000 evaluations in 30 independent runs, in comparing we can found

TABLE IV
GD VALUES OF CSEA AND M-CSEA ON DTLZS BENCHMARK SET

MOEAs \ Evaluations	1000	2000	3000	4000	5000	6000	7000	8000	9000	10000
DTLZ1										
CSEA	38.304222	0.249517037	0.093476884	0.067478989	0.049841399	0.039018872	0.037277541	0.036718244	0.038506013	0.023741603
M-CSEA	34.55195154	7.293731419	1.659906579	0.593031306	0.253315597	0.093995118	0.049742977	0.045665123	0.007362405	0.004428484
DTLZ2										
CSEA	0.111937965	0.003129364	0.001690321	0.001120471	0.00073779	0.000533606	0.000432797	0.000362262	0.000307509	0.000255133
M-CSEA	0.10952606	0.003268781	0.001166987	0.000817451	0.000578712	0.000431733	0.00034105	0.000305797	0.00026639	0.000231731
DTLZ3										
CSEA	183.1624496	55.90315143	19.85598041	14.28060864	11.19028972	8.80644816	6.527776933	5.115575763	3.579570234	2.017174022
M-CSEA	189.4951725	47.64296364	25.87443973	9.138374728	6.40377985	3.355073622	3.030661077	1.628800101	0.748851809	0.575976429
DTLZ4										
CSEA	0.16005859	0.009324148	0.00158179	0.000911676	0.00062676	0.00050403	0.00043422	0.000366466	0.000306355	0.000280448
M-CSEA	0.14623881	0.00786399	0.002191418	0.001288786	0.000863088	0.000676222	0.00057595	0.000471852	0.000425611	0.00039424
DTLZ5										
CSEA	0.143045187	0.009094896	0.002914526	0.001666095	0.000700365	0.00038713	0.000316341	0.000294842	0.000244925	0.000211265
M-CSEA	0.129459625	0.005469629	0.002643803	0.001694055	0.000970129	0.000418626	0.000343096	0.000282539	0.000252634	0.000134312
DTLZ6										
CSEA	0.937372419	0.490257754	3.61595E-06	1.61995E-06	1.18903E-06	9.76567E-07	8.60556E-07	7.69527E-07	7.12257E-07	6.60703E-07
M-CSEA	0.982408347	0.392787218	0.103280342	0.001726882	0.000270277	0.000137908	1.72537E-05	1.31121E-05	7.19235E-06	5.97902E-06
DTLZ7										
CSEA	2.115699646	0.037620821	0.002642258	0.001528894	0.000967458	0.000684218	0.000527802	0.000431631	0.000384352	0.000343302
M-CSEA	2.631468079	0.051293186	0.008160607	0.00309386	0.001154022	0.000798814	0.000458733	0.000350707	0.000289746	0.000250762
DTLZ8										
CSEA	0.058641647	0.017278221	0.017278221	0.017278221	0.017278221	0.017278221	0.017278221	0.017278221	0.017278221	0.017278221
M-CSEA	0.052658985	0.015295372	0.015295372	0.015295372	0.015295372	0.015295372	0.015295372	0.015295372	0.015295372	0.015295372

TABLE V
IGD VALUES OF CSEA AND M-CSEA ON DTLZS BENCHMARK SET

MOEAs \ Evaluations	1000	2000	3000	4000	5000	6000	7000	8000	9000	10000
DTLZ1										
CSEA	55.70363999	0.984059887	0.535546843	0.422683497	0.319925917	0.307731215	0.304742746	0.303695285	0.257744343	0.079166362
M-CSEA	67.37843306	0.963184028	0.503859492	0.390551114	0.29077677	0.249641266	0.238247688	0.230690889	0.083235634	0.055616082
DTLZ2										
CSEA	0.477951349	0.110380642	0.074753971	0.065342481	0.057251695	0.052646793	0.050160575	0.047509835	0.045705387	0.043639601
M-CSEA	0.445656061	0.105494378	0.077526559	0.069578473	0.064227686	0.059277889	0.056366055	0.055084377	0.054108945	0.051323836
DTLZ3										
CSEA	500.9795407	137.0799636	97.40436478	74.77410613	51.0949479	32.67368451	24.24581574	15.80437507	10.63840422	9.87021718
M-CSEA	445.3088212	81.03428807	48.27255007	34.01990624	19.12160089	9.809608997	8.919420738	3.96288272	3.06367452	2.820373235
DTLZ4										
CSEA	0.920569575	0.177499921	0.083470208	0.065663179	0.061877807	0.058606665	0.053952637	0.051886509	0.048899923	0.047121413
M-CSEA	0.867610961	0.167420331	0.074209237	0.05987287	0.052081465	0.048150374	0.044544918	0.041799879	0.039348027	0.037201222
DTLZ5										
CSEA	0.479076784	0.049762862	0.019695687	0.012999939	0.008494638	0.006550633	0.005835436	0.005027044	0.004493137	0.004167442
M-CSEA	0.424670314	0.03700648	0.01817932	0.013216065	0.009003176	0.007227901	0.006568781	0.006156097	0.005924713	0.004950872
DTLZ6										
CSEA	8.52857946	1.487209415	0.158712128	0.018593501	0.009388897	0.007934719	0.006343882	0.005608391	0.005024294	0.004639098
M-CSEA	8.547355148	1.926878075	0.117556984	0.015819266	0.00962453	0.006568246	0.004810618	0.003727698	0.003479865	0.003282058
DTLZ7										
CSEA	7.993415905	0.391054762	0.108179002	0.076912183	0.059149384	0.047771425	0.041651639	0.037564026	0.033012235	0.031701786
M-CSEA	8.948192414	1.211191478	0.068162702	0.050790774	0.042785891	0.038499199	0.035339474	0.032621292	0.031239453	0.030098333
DTLZ8										
CSEA	0.24116755	0.136231997	0.136231997	0.136231997	0.136231997	0.136231997	0.136231997	0.136231997	0.136231997	0.136231997
M-CSEA	0.245041877	0.111925885	0.111925885	0.111925885	0.111925885	0.111925885	0.111925885	0.111925885	0.111925885	0.111925885

TABLE VI
PERFORMANCE TIME FOR K-RVEA AND M-K-RVEA

Problems	K-RVEA (seconds)	M-K-RVEA (seconds)	Faster (%)
DTLZ1	6,864	6,087	11.32
DTLZ2	9,845	9,694	1.53
DTLZ3	15,621	13,851	11.33
DTLZ4	17,432	14,360	17.62
DTLZ5	11,230	10,375	7.61
DTLZ6	19,901	19,866	0.18
DTLZ7	28,823	26,667	7.48
DTLZ8	118,080	105,529	10.63
DTLZ9	34,750	33,969	2.25

TABLE VII
PERFORMANCE TIME FOR CSEA AND M-CSEA

Problems	K-RVEA (seconds)	M-K-RVEA (seconds)	Faster (%)
DTLZ1	3,952	2,824	28.54
DTLZ2	5,067	5,009	1.14
DTLZ3	12,340	11,337	8.13
DTLZ4	12,53	9,781	21.96
DTLZ5	7,864	6,339	19.39
DTLZ6	13,468	13,166	2.24
DTLZ7	16,446	15,323	6.83
DTLZ8	80,781	77,554	3.99

that, M-CSEA with dynamic selection strategy can get more balance of convergence and diversity for obtained population. In details, in the early stages (since 4000 evaluations), it gets better in DTLZ1, DTLZ2, DTLZ3, DTLZ6, DTLZ8 problems on GD and DTLZ1, DTLZ3, DTLZ4, DTLZ6, DTLZ7, DTLZ8 problems on IGD. During the generations (since 7000 evaluations), M-CSEA gets better in DTLZ2,

DTLZ3, DTLZ6, DTLZ7, DTLZ8 on GD. It is also better in DTLZ1, DTLZ3, DTLZ4, DTLZ6, DTLZ7, DTLZ8 on IGD. In the last stages (completed 10000 evaluations), M-CSEA is gets better in DTLZ1, DTLZ2, DTLZ3, DTLZ5, DTLZ6, DTLZ7, DTLZ8 on GD and DTLZ1, DTLZ2, DTLZ3, DTLZ6, DTLZ7, DTLZ8 on IGD. In other problems, CSEA gets a bit better than M-CSEA.

Besides, in terms of performance time, on the test

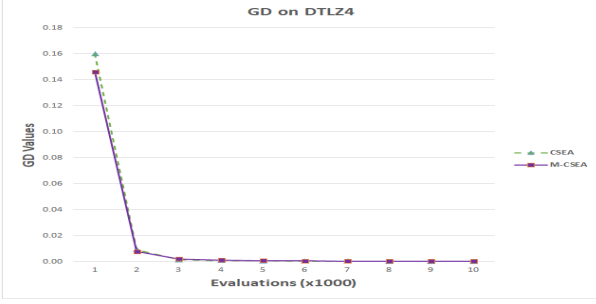


Figure 23. The behavior of M-CSEA and CSEA on DTLZ4 (GD)

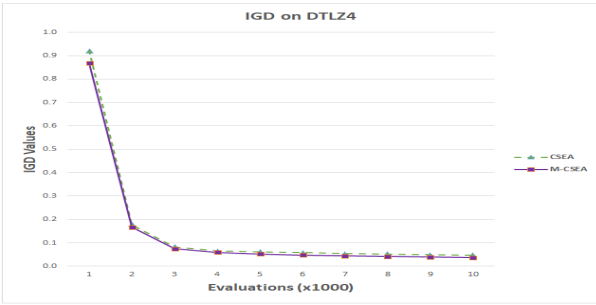


Figure 24. The behavior of M-CSEA and CSEA on DTLZ4 (IGD)

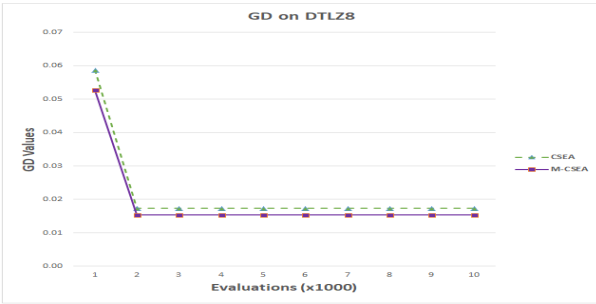


Figure 25. The behavior of M-CSEA and CSEA on DTLZ8 (GD)

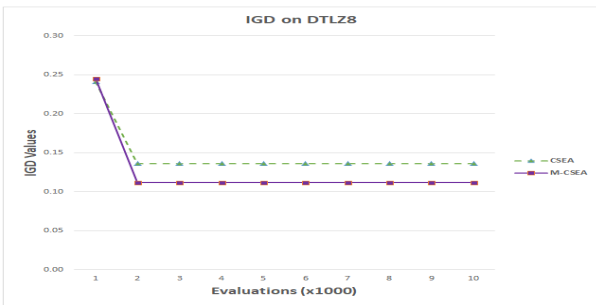


Figure 26. The behavior of M-CSEA and CSEA on DTLZ8 (IGD)

environment, HP ProLiant DL380 with CPU: 2xIntel Xeon Silver 4110 Processor 11M cache 2.10GHz, RAM: 4x31GB PC4-21300, HDD: 4x1TB 7200 rpm SATA 6Gb/s, run-time recognition result between K-RVEA and M-K-RVEA; CSEA and M-CSEA are shown in Table VI and VII.

From the mechanism of applying the guidance method in M-K-RVEA and M-CSEA we can found that, M-K-RVEA with adaptive guidance mechanism, in the early stages, the number of using the origin function increases because of the need to increase convergence, however, when the evolution time is long enough, the training dataset or the external archive is good quality, the frequency of which is mainly used by the surrogate function. On the other hand, the original function is complex, the cost is much larger than the surrogate function, so with a reasonable range of variation in the proposed paper, the overall calculation time will be significantly reduced. This depends on the complexity of the original problem. The M-K-RVEA example greatly reduces time in DTLZ1, DTLZ3, DTLZ4 and DTLZ8 problems.

With the M-CSEA, with an adaptive guidance mechanism in determining the number of samples for model training, in the first phase, it is necessary to use a lot of data from the population to train the model, however with the evolution time is long enough, the number of samples needed is reduced, the search process will be faster, due to the reduced number of samples used. M-CSEA is faster on DTLZ1, DTLZ4, DTLZ5 problems. Here, the different performance efficiency between M-K-RVEA and M-CSEA shows that in addition to the impact assessment raised factors, there are other influencing factors from the techniques used in search, selection, crossover and mutation of each algorithm.

Overall assessment, in addition to a significant improvement in the convergence quality and diversity of the obtained solutions as well as maintaining a balance between exploitation and exploration of the evolutionary process, the proposed guidance method also helps to reduce the computation time of surrogate assisted multi-objective algorithms, especially suitable for solving expensive problems.

In general, through the experimental results, it is proposed to use new selection strategy with adaptive number of reference solutions, which has an effect to implement the optimal target optimization algorithm using general surrogate functions, selection of reference solutions for sampling in machine learning in the FNNs. The obvious meaning of using the response parameter is to assess the quality of the population to determine how many reference solutions are needed for forming the boundary of the classification. The results using dynamic selection strategy by the K parameter demonstrate the hypothesis of the impact of the time process, the current convergence quality of the population, and the characteristics of the problem on maintaining the equilibrium of convergence and diversity quality of the population, through ensuring a balance between the exploration and exploitation process

of evolution. The effectiveness of adaptive selection in the current stages also depends on the quality and effectiveness of the previous stage, proving that a timely adjustment to maintain a balance between exploration and exploitation is needed. However, from the results of some problems, CSEA has a better quality, although not large, it also gives us issues that need to be adequately posed such as: stability, sustainability, random factors course, etc. Those factors influence the choice of parameters appropriately, or guide evolution to better efficiency, the robustness of the algorithm will be improved by reducing the noises with the dynamic selection strategy.

VI. CONCLUSIONS

Using surrogate model is an advanced and effective technique in solving expensive problems with MOEAs. However, on the principle of using the surrogate function to replace the original objective function with its own techniques, although it makes the algorithm flexible and diverse, it also makes the evolutionary environment more turbulent, impact the power of the algorithm. This is a research problem set out in order to improve the quality of algorithms using surrogate models, creating effective algorithms that are comprehensive, capable of applying to problems in practice.

The focus of study is on the effects of noisy environments on the performance aspects of MOEAs using surrogate model, the paper has analyzed and evaluated the effects of techniques used in MOEAs using surrogate models that are Kriging, FNNs and guidance techniques. The paper has also analyzed and evaluated factors in parameter selection about frequency of using surrogate model, sample selection methods for training models, the binding factors between sustainability and convergence and diversity through the values of the GD, IGD measures over time of evolution. The paper also proposes to use instructional techniques to automatically adjust algorithms to maintain the convergence and diversity properties as well as the exploration and exploitation capabilities of the evolutionary process. Through experiments on K-RVEA, CSEA and the modified versions, the results showed that with proposed guidance techniques, the algorithms have been improved on quality and thereby reducing noise for the evolutionary environment, contributing to improve the robustness, creating more effective and comprehensive innovative algorithms.

In the next studies, we will study the impacts of practical problems in design, decision making... combined with the above research results to adjust and supplement flexible guidance techniques, responding more effectively to solving expensive problems in practice.

REFERENCES

- [1] G. Matheron, "Principles of geostatistics," *Economic geology*, vol. 58, no. 8, pp. 1246–1266, 1963.
- [2] S. S. Garud, I. A. Karimi, and M. Kraft, "Design of computer experiments: A review," *Computers & Chemical Engineering*, vol. 106, pp. 71–95, 2017.
- [3] F. Bittner and I. Hahn, "Kriging-assisted multi-objective particle swarm optimization of permanent magnet synchronous machine for hybrid and electric cars," in *Electric Machines & Drives Conference (IEMDC), 2013 IEEE International*. IEEE, 2013, pp. 15–22.
- [4] S. Choi, J. J. Alonso, and H. S. Chung, "Design of a low-boom supersonic business jet using evolutionary algorithms and an adaptive unstructured mesh method," *AIAA paper*, vol. 1758, p. 2004, 2004.
- [5] M. T. Emmerich, K. C. Giannakoglou, and B. Naujoks, "Single- and multiobjective evolutionary optimization assisted by gaussian random field metamodelling," *IEEE Transactions on Evolutionary Computation*, vol. 10, no. 4, pp. 421–439, 2006.
- [6] L. Gonzalez, J. Periaux, K. Srinivas, and E. Whitney, "A generic framework for the design optimisation of multidisciplinary uav intelligent systems using evolutionary computing," in *44th AIAA Aerospace Sciences Meeting and Exhibit*, 2006, p. 1475.
- [7] A. Husain and K.-Y. Kim, "Enhanced multi-objective optimization of a microchannel heat sink through evolutionary algorithm coupled with multiple surrogate models," *Applied Thermal Engineering*, vol. 30, no. 13, pp. 1683–1691, 2010.
- [8] I. Voutchkov and A. Keane, "Multi-objective optimization using surrogates," in *Computational Intelligence in Optimization*. Springer, 2010, pp. 155–175.
- [9] Q. Zhang, W. Liu, E. Tsang, and B. Virginas, "Expensive multiobjective optimization by moea/d with gaussian process model," *IEEE Transactions on Evolutionary Computation*, vol. 14, no. 3, pp. 456–474, 2010.
- [10] S. Maroufpoor, O. Bozorg-Haddad, and X. Chu, "Geostatistics: principles and methods," in *Handbook of Probabilistic Models*. Elsevier, 2020, pp. 229–242.
- [11] K. Deb and H. Gupta, "Introducing robustness in multi-objective optimization," *Evolutionary computation*, vol. 14, no. 4, pp. 463–494, 2006.
- [12] A. Gaspar-Cunha and J. A. Covas, "Robustness in multi-objective optimization using evolutionary algorithms," *Computational Optimization and Applications*, vol. 39, no. 1, pp. 75–96, 2008.
- [13] S. Gunawan and S. Azarm, "Multi-objective robust optimization using a sensitivity region concept," *Structural and Multidisciplinary Optimization*, vol. 29, no. 1, pp. 50–60, 2005.
- [14] M. Dellnitz and K. Witting, "Computation of robust pareto points," *International Journal of Computing Science and Mathematics*, vol. 2, no. 3, pp. 243–266, 2009.
- [15] L. T. Bui, D. Essam, H. A. Abbass, and D. Green, "Performance analysis of evolutionary multi-objective optimization methods in noisy environments," in *Proceedings of the 8th Asia Pacific symposium on intelligent and evolutionary systems*. Citeseer, 2004, pp. 29–39.
- [16] E. Zitzler, L. Thiele, and K. Deb, "Comparison of multi-objective evolutionary algorithms: Empirical results," *Evolutionary Computation*, vol. 8, no. 1, pp. 173–195, 2000.
- [17] T. Chugh, Y. Jin, K. Miettinen, J. Hakanen, and K. Sindhy, "K-rvea: a kriging-assisted evolutionary algorithm for many-objective optimization," *Reports of the Department of Mathematical Information Technology, Series B, Scientific Computing no. B*, vol. 2, 2016.
- [18] L. Pan, C. He, Y. Tian, H. Wang, X. Zhang, and Y. Jin,

“A classification-based surrogate-assisted evolutionary algorithm for expensive many-objective optimization,” *IEEE Transactions on Evolutionary Computation*, vol. 23, no. 1, pp. 74–88, 2018.

- [19] R. Cheng, Y. Jin, M. Olhofer, and B. Sendhoff, “A reference vector guided evolutionary algorithm for many-objective optimization,” *IEEE Transactions on Evolutionary Computation*, vol. 20, no. 5, pp. 773–791, 2016.
- [20] K. Deb, L. Thiele, M. Laumanns, and E. Zitzler, “Scalable test problems for evolutionary multi-objective optimization, TIK-Report no. 112,” *Computer Engineering and Networks Laboratory (TIK), Swiss Federal Institute of Technology (ETH), Zurich*, Tech. Rep., 2001.
- [21] V. L. Vachhani, V. K. Dabhi, and H. B. Prajapati, “Survey of multi objective evolutionary algorithms,” in *2015 International Conference on Circuits, Power and Computing Technologies [ICCPCT-2015]*. IEEE, 2015, pp. 1–9.



Hoai Xuan NGUYEN received the Ph.D. degree in computer science from the University of New South Wales, Kensington, Australia. He has been involved with academics including teaching and research since 1998. Currently, he is an Associate Professor and Director of Vietnam Institute of Artificial Intelligence, Hanoi, Vietnam.

He is researching in the field of Intelligence Computation, Natural language processing; Operations research.

Email: nxhoai@gmail.com.



Dinh Duc NGUYEN has his B.Sc (2000) and M.Tech (2010) degree from Le Quy Don Technical University. Currently he is a PhD Candidate at the Military Information Technology Institute, Academy of Military Science and Technology. His research interests are in the areas of Evolutionary Multi-objective Optimizations, Computational Intelligence, Multi-agent systems, Simulation systems, UAV and their applications to defence and security domain.

Email: nddinh76@gmail.com.



Long NGUYEN received his PhD in 2015 in Fundamentals of Mathematics for Information Technology from Le Quy Don Technical University. Currently he is a Associate Professor at the Department of Information Technology, National Defense Academy. His research interests are in the areas of Evolutionary Multi-objective Optimizations, Computational Intelligence, Multi-agent systems, Cyber warfare, Artificial Intelligence and their applications to defence and security domain.

Email: longit76@gmail.com.

A Comprehensive Survey of Frequent Itemsets Mining on Transactional Database with Weighted Items

Huan Phan^{1,3}, Bac Le^{2,3}

¹Faculty of Mathematics and Computer Science, Ho Chi Minh City University of Science, Vietnam

²Faculty of Information Technology, Ho Chi Minh City University of Science, Vietnam

³Vietnam National University, Ho Chi Minh City, Vietnam

Correspondence: Huan Phan, huanphan@hcmussh.edu.vn

Communication: received 5 April, revised 19 May, accepted 31 May

Digital Object Identifier: 10.32913/mic-ict-research.v2021.n2.967

Abstract: In 1993, Agrawal et al. proposed the first algorithm for mining traditional frequent itemset on binary transactional database with unweighted items - This algorithm is essential in finding hidden relationships among items in your data. Until 1998, with the development of various types of transactional database - some researchers have proposed a frequent itemsets mining algorithms on transactional database with weighted items (the importance/meaning/value of items is different) - It provides more pieces of knowledge than traditional frequent itemsets mining. In this article, the authors present a survey of frequent itemsets mining algorithms on transactional database with weighted items over the past twenty years. This research helps researchers to choose the right technical solution when it comes to scale up in big data mining. Finally, the authors give their recommendations and directions for their future research.

Keywords: Data mining, frequent itemsets, weighted items

I. INTRODUCTION

Association rule mining is an important technique in data mining. The goal of mining is to discover relationships between data values in a dataset. The first model of association rule mining is the binary model, which is also known as the fundamental model, was proposed in 1993 by Agrawal et al. [1] to analyse data transaction, detect relationships between items sold in supermarkets. From there, they can have a reasonable investment and business plan, and organize the sales counter to have revenue in the most profitable trading sessions. In addition, this knowledge can be applied to predict the number of upcoming best-selling items and customer shopping trends. Synthesizing this knowledge to plan operations, production and business in a more convenient way will reduce the time for statistics, market research, etc.

The proposed algorithms for association rule mining are divided into two phases:

Phase 1: Find all itemsets that satisfy the minimum support threshold *minsup*. This phase takes a lot of time to process.

Phase 2: Generate association rules in turn from itemsets that satisfy *minsup* in phase 1 and these association rules must satisfy minimum confidence threshold *minconf*.

Agrawal has proposed the Apriori algorithm [2] - Frequent Itemsets (FI) mining, an algorithm that scans the data many times and has an exponential complexity. To improve the time in frequent itemsets mining, many researchers have proposed an efficient frequent itemsets mining algorithm based on the storage structures that reduce the search space such as SE-Tree [8], Prefix-Tree [4], IT-Tree [6], FP-Tree [3], etc. However, in practice, generating the frequent itemsets is time consuming and has a very large number of itemsets. Therefore, some researchers have proposed mining the Closed Frequent Itemsets (CFI) with fewer frequent itemsets such as A-Close [5], Charm [6], etc. Meanwhile, some other scientists also proposed mining the Maximal Frequent Itemsets (MFI) such as Pincer-Search [7], Max-Miner [8], etc. In some practical applications, to apply FI, CFI and MFI costs a lot of computation as well as a large number of frequent itemsets. Bayardo proposed mining the Maximum Length Frequent Itemsets (LFI) which is a subset of the frequent itemsets FI and contains only frequent itemsets of maximum length like the DepthProject [9], MAXLEN-FI [10], etc.

In the last years of the 20th century, along with the development of diverse transaction data and inheritance from Agrawal's traditional frequent itemsets mining; Ramkumart et al. have proposed the WIS algorithm [11] that mines

the frequent itemsets with weighted items (the importance / significance level of the items is different) containing more knowledge than the traditional frequent itemsets mining (without weighted of items). Besides, Cai et al. also proposed MINWAL algorithm [12] to solve the above problem. After that, many authors researched and proposed algorithms [13-21] to solve this problem. However, all algorithms approach and solve in the direction of satisfying the “downward closure property”. In 2011, Huai et al. proposed the WHIUA algorithm [22] which, following the Apriori approach and “not satisfy the downward closure property”, significantly increases the search space - this is a big challenge for researchers of data mining. In the next decade, a number of algorithms were proposed such as IWFPP [23], PWA [27], WAC [30], etc. but most of them still solve the problem in the direction of satisfying the “downward closure property” and the algorithm are similar to traditional frequent itemsets mining. Therefore, it has motivated the authors to research and summarize a number of algorithms for mining frequent itemsets on weighted transaction data proposed in the period 1998 to present.

In Part II, the article presents the basic concepts of association rule mining, frequent itemsets, Apriori algorithm and common data structures. In Part III, we summarize some algorithms for mining frequent itemsets on transaction database with weighted items over the years. Discussion and recommendations are presented in Part IV; Conclusions and development directions are presented in Part V.

II. THE BASIC CONCEPTS

1. Mining association rules

Let $I = \{i_1, i_2, \dots, i_m\}$ be a set of m properties, each property is called an item. The set $W = \{w_1, w_2, \dots, w_m\}, \forall w_j \geq 0$ is the weight (significance/importance) of each item. The set of items $X = \{i_1, i_2, \dots, i_k\} \forall i_j \in I (1 \leq j \leq k)$ is called itemset, itemset with k items is called k -itemset. $\mathcal{D}$ is transaction database, including n distinct records called set of transaction $T = \{t_1, t_2, \dots, t_n\}$, each transaction $t_k = \{q_1 i_{k_1}, q_2 i_{k_2}, \dots, q_m i_{k_m}\}, i_{k_j} \in I (1 \leq k_j \leq m)$ and $q_j \geq 0$ is the quantity/probability of item i_{k_j} in transaction t_k .

Definition 1: Let $X, Y \subseteq I$ with $X \cap Y = \emptyset$, the association rule is a implication of the form $X \rightarrow Y$, satisfying two given thresholds (minsup - minimum support; minconf - minimum confidence), where X is called the antecedent and Y is the consequent.

Definition 2: The support of itemset $X \subseteq I$, denoted $\text{sup}(X)$ - the ratio between the number of transactions in $\mathcal{D}$ containing X and n transactions.

$$\text{sup}(X) = \frac{|\{t \in \mathcal{D} \mid X \subseteq t\}|}{n} \quad (1)$$

TABLE I
TRANSACTION DATABASES CLASSIFICATION

Data type	Feature description
Binary (traditional)	$\forall w_j = 1, \forall q_j = 1 (1 \leq j \leq m)$
Uncertain (Fuzzy)	$\forall w_j, q_j \in [0, 1] (1 \leq j \leq m)$
Unweighted quantity and item	$\forall w_j = 1, \forall q_j \geq 0 (1 \leq j \leq m)$
Weighted quantity and item	$\forall w_j \in [0, 1], \forall q_j \geq 0 (1 \leq j \leq m)$
Weighted item*	$\forall w_j \in [0, 1], \forall q_j = 1 (1 \leq j \leq m)$

(*) In this article, the authors only focus on surveying the research related to the frequent itemsets on transaction database with weighted items.

Definition 3: The support of the association rule $X \rightarrow Y$, denoted $\text{sup}(X \rightarrow Y)$ - the ratio between the number of transactions in $\mathcal{D}$ containing $\{X \cup Y\}$ and n transactions.

$$\text{sup}(X \rightarrow Y) = \text{sup}(X \cup Y) \quad (2)$$

Definition 4: Confidence of association rule $X \rightarrow Y$, denoted $\text{conf}(X \rightarrow Y)$ - the ratio between the number of transactions containing $\{X \cup Y\}$ and the number of transactions containing X in $\mathcal{D}$.

$$\text{conf}(X \rightarrow Y) = \frac{\text{sup}(X \cup Y)}{\text{sup}(X)} \quad (3)$$

When we say that the confidence of a rule is 95%, it means that 95% of the transactions that contain the antecedent X also contain the consequent Y - “the conditional probability that the event Y occurs is 95%”, with the condition “event X occurs”. The confidence of the association rule represents the correlation between the events X and Y . The confidence is used to measure the significance of the rule. Usually in tasks, users are only interested in high confidence association rules.

Pseudocode 1: Association Rules Mining

Input: Transaction database $\mathcal{D}$, minsup, minconf

Output: Association Rules

-
1. $FI = \{\emptyset\}$ //frequent itemset – Phase 1
 2. For each $Z \in P_{\geq 1}(I)$ do
 3. If $\text{sup}(Z) \geq \text{minsup}$ then
 4. $FI = FI \cup Z$
 5. $AR = \{\emptyset\}$ //Association Rules – Phase 2
 6. For each $fi \in FI$ do
 7. If $\text{conf}(X \rightarrow \{fi \setminus X\}) \geq \text{minconf}$ then
 8. $AR = AR \cup (X \rightarrow \{fi \setminus X\})$
 9. Return AR
-

$(P_{\geq 1}(I))$: all powerset of I have at least one item)

Algorithm 1 is divided into 2 phases

Phase 1: (Lines 1 to 4) this is the phase to find all itemset that satisfy the minsup threshold;

Phase 2: (Lines 5 to 8) generate association rules from itemset that satisfy minsup in phase 1 and these association rules must satisfy minconf (association rule $X \rightarrow \{fi \setminus X\}$ has the antecedent X and the consequent $Y = \{fi \setminus X\}$ satisfying $X \cap Y = \emptyset$).

Analyze space search algorithm 1:

Phase 1: transaction database $\mathcal{D}$ has $|I| = m$, all itemset generated from I have at least one item - $P_{\geq 1}(I)$, we have a space of combinations between items generated from m items of $2^m - 1$, this is search space in phase 1 for identifying itemsets that satisfy the minsup threshold.

Phase 2: after determining the set of combinations or itemsets satisfying minsup, this phase generates association rules from the itemsets in turn. Suppose, for itemset X with m items, $l(1 \leq l \leq m)$ and $r(1 \leq r \leq m-l)$ are the number of items in the precondition and the conclusion. Then, we have C_l^m number of ways to choose the antecedent with l items over m items from itemsets X and C_r^{m-l} number of ways to choose the consequent with r items on $(m-l)$ items left from itemset X :

$$\sum_l^m \sum_r^{m-l} C_l^m C_r^{m-l} \quad (4)$$

Applying Newton's triangle, equation (4) is rewritten as follows:

$$\sum_l^m \sum_r^{m-l} C_l^m C_r^{m-l} = 3^m - 2^{m+1} + 1 \quad (5)$$

Most of the proposed association rule mining algorithms focus on the improvement technique in phase 1 (determination of itemsets satisfying minsup). All these algorithms assume that phase 1 takes the most time in the whole association rule mining process. However, according to the above analysis, phase 2 is more complex and has a large generation space; or in the other words, in algorithm 1, every phase is necessary to improve and enhance calculating performance. For example, we have itemset $X = \{A, B, C\}$, $m = 3$, the total number of association rules generated from itemset X as calculated by equation (5) is $C_1^3 \times (C_2^2 + C_1^2) + C_2^3 \times (C_1^1) = 3 \times (1 + 2) + 3 \times (1) = 12$, that means for itemset X have 3 items, it is necessary to consider whether the 12 association rules satisfy the minconf threshold or not. Similarly, if itemset X has 4 items ($m = 4$), then we need to consider 50 association rules that satisfy minconf. In addition, the above analysis also clearly shows that the search space in both phases does not depend on the number of transactions of the dataset D but only on the number of items of the dataset mining.

P_1 : the space for generating combinations $(2^m - 1)$ in Phase 1;

TABLE II
ESTIMATING THE GENERATION SPACE OF COMBINATIONS AND THE NUMBER OF ASSOCIATION RULES GENERATED FROM M-ITEMSET

m	10	20	30	40	50
P_1	$2^{10} - 1$	$2^{20} - 1$	$2^{30} - 1$	$2^{40} - 1$	$2^{50} - 1$
P_2	$3^{10} - 2^{11} + 1$	$3^{20} - 2^{21} + 1$	$3^{30} - 2^{31} + 1$	$3^{40} - 2^{41} + 1$	$3^{50} - 2^{51} + 1$
P_2/P_1	57	3.325	191×10^3	11×10^6	637×10^6

P_2 : the number of association rules generated from m - itemset $(3^m - 2^{m+1} + 1)$ in Phase 2.

Table II, shows that the ratio between the space generated in Phase 2 and Phase 1 is very large when the number of items are increased, that is Phase 2 has a very large space compared to Phase 1.

Definition 5: Let $X \subseteq I$, X is called the frequent Denoted FI is the set of the frequent itemsets.

Definition 6: Let $X \subseteq I$, X is called the Closed frequent itemset - if X is a frequent itemset and there is no parent set of the same support. Denoted CFI is the set of closed frequent item sets.

Definition 7: Let $X \subseteq I$, X is called the maximal frequent itemset - if X is a frequent itemset and there is no parent set of frequent itemsets. Denoted MFI is the set of the maximal frequent itemsets.

In addition, when it is necessary to mine the association rules with the largest number of items, the researchers also propose mining the frequent itemset of maximum length - which is a subset of the maximal frequent itemset and has the maximum length ($\mathbf{LFI} \subseteq \mathbf{MFI} \subseteq \mathbf{CFI} \subseteq \mathbf{FI}$).

Definition 8: Let $X \in \mathbf{FI}$, X is called the frequent itemset of maximum length - if $\forall Y \in \mathbf{FI}$ then $|X| \geq |Y|$, that means, the number items of itemset X greater than or equal the number items of any frequent itemsets contained in FI. Denoted LFI is the set of the maximum length frequent itemsets.

Some properties of frequent itemsets: these are the fundamental properties used for reducing the search space - these properties are called the Downward Closure Property (DCP), which is also known as Apriori properties.

Property 1: $\forall X \subseteq Y : \sup(X) \geq \sup(Y)$;

Property 2: $\forall X \subseteq Y, \sup(Y) \geq \text{minsup} : \sup(X) \geq \text{minsup}$;

Property 3: $\forall X \subset Y, \sup(X) < \text{minsup} : \sup(Y) < \text{minsup}$.

2. Apriori algorithm

The algorithm proposed by Agrawal in 1994 [2], is considered historic in the association rule mining, because it is far beyond the reach of familiar algorithms. Apriori is the foundational algorithm for finding frequent itemsets using candidate generation methods. The algorithm is characterized by Breadth-First Search using the Apriori property:

any infrequent $(k - 1)$ -itemset cannot be a sub-itemset of the frequent k -itemsets.

Pseudocode 2: Frequent Itemset Mining

Input: Transaction database $\mathcal{D}$, minsup

Output: Frequent Itemsets

-
1. Generate candidate C_{k+1} from frequent k -itemset;
 2. Scan data - calculate support of candidates;
 3. Add the candidates to the $(k + 1)$ -itemset list.
-

Advantage: The algorithm is based on a fairly simple comment that any sub-itemset of frequent itemsets are also frequent itemsets (property 2). Therefore, in the process of finding the candidate itemset, the algorithm only needs to use the candidate itemset that appeared in the previous step, not all of the candidate itemsets (up to that point). So, the memory is significantly released.

Disadvantage: The algorithm has to scan the data $(\max + 1)$ times with $\max$ being the length of the longest frequent itemset. The Apriori algorithm reduces the space based on the Apriori property. However, when the number of frequent itemsets generated is large, the fact that the $\max$ is large or the minsup is small will result in generating a lot of candidate itemsets and having to traverse the data many times, the algorithm has a high cost. In practice, 2^{100} candidates (in the worst case) need to be generated to find the frequent itemset of size 100.

From the above disadvantages, many researchers have proposed algorithms to increase the performance of generating frequent itemsets based on the organization of data storage and corresponding effective search strategies.

3. The common data structure on mining the traditional frequent itemsets

In addition to some algorithms with Apriori approach, we present a survey of common data structures used by many authors for storing the search space in the mining of frequent itemset and search strategies along with storage space. Search strategies on tree structure: Depth First Search - DFS, Breadth First Search - BFS or a combination of both.

a) Lattice

This is a data structure that is used a lot in the proposed algorithms. Transaction database $\mathcal{D}$ have $|I| = m$, all itemset (powerset) are generated from $I - P(I)$, we have the space of combinations between items generated from m items is 2^m , this is the full search space. In 1999, Pasquier proposed the A-Close algorithm [5] based on the Lattice structure to mining the closed frequent itemset (nearly 2,100 citations). **Search strategy:** combining both dimensions (Top-Down and Bottom-Up).

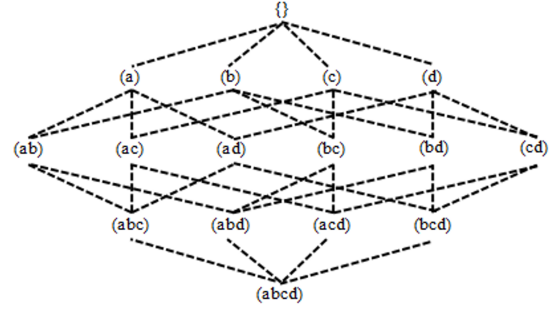


Figure 1. Lattice

Advantages: simple, ensure mining in full the frequent itemsets.

Disadvantages: takes a lot of memory during mining.

b) Set Enumeration Tree (SE-Tree)

In 1998, Bayardo proposed the Max-Miner algorithm [8] based on an enumeration tree (over 2,000 citations) to mine the maximal frequent itemsets.

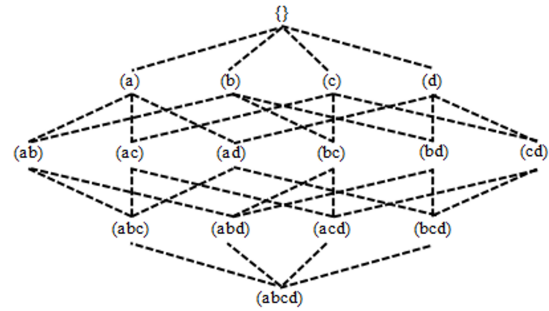


Figure 2. Set Enumeration Tree (SE-Tree)

Search strategy: DFS.

Advantages: simple, ensure mining in full the frequent itemsets.

Disadvantages: unbalance in the mining process (depending on the frequency of the item).

c) Prefix-Tree

In 2002, Liu proposed the OpportuneProject algorithm [4] based on the Prefix-Tree structure (nearly 300 citations) to quickly mining the frequent itemset. Similar to an SE-Tree, the contents of the archive are shortened through prefixes.

d) IT-Tree (Itemset Tidset Tree)

In 1999, Zaki proposed the Charm algorithm [6] along with an IT-Tree (Itemset Tidset Tree) structure like a list tree structure and a combination of storing transaction identifier of itemset (nearly 1,600 citations).

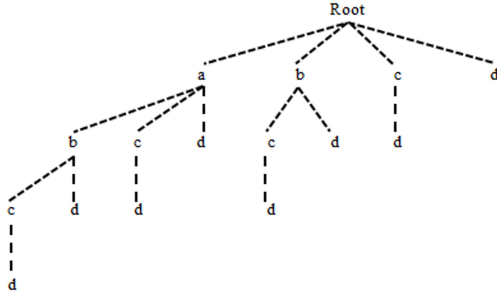


Figure 3. Prefix-Tree

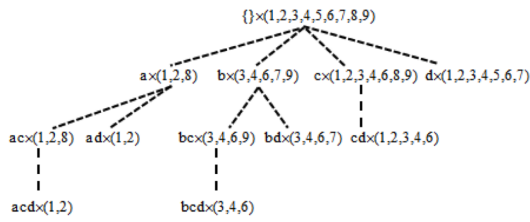


Figure 4. IT-Tree

Search strategy: DFS.

Advantages: simple, ensure mining in full the frequent itemsets.

Disadvantages: unbalance during mining and memory intensive storage.

e) *FP-Tree (Frequent Pattern Tree)*

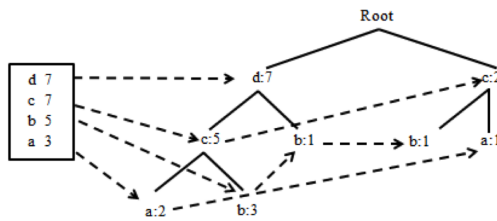


Figure 5. FP-Tree

In 2000, Jiawei Han et al. proposed the FP-Tree structure [3] (nearly 9,000 citations) and the FP-Growth algorithm for mining frequent itemsets. This is an extended tree from Prefix-Tree, with an additional header table that stores items, frequencies and links to nodes in the tree. From that, the FP-Tree structures have been used and improved by many researchers.

Although the FP-Tree structure is more reduced than the above ones', FP-Tree is an extension from Prefix-Tree,

so there are still many limitations similar to Prefix-Tree - unbalanced tree.

III. SURVEY OF FREQUENT ITEMSETS MINING ON TRANSACTIONAL DATABASE WITH WEIGHTED ITEMS

1. Algorithm grouping

In practice, the properties in the data are not equal. There are some properties that are emphasized, we say those properties have a higher importance level than others. Over the past 20 years, this has always been a very interesting research and many proposals have been made (by researchers) to solve this problem. In this survey, the authors categorize as well as group the research works of domestic/foreign researchers based on the method of calculating the measures related to the weights of the itemset.

a) *Some foreign research work*

Group 1: In 1998, G. D. Ramkumar et al. proposed the WIS algorithm [11].

Description WIS algorithm (1998)	
Apriori property	Satisfy the downward closure property
Weighted items	Weight of each item $w_j \in [0, 1]$
Approach	Apriori
Association rule	Weighted binary
Transaction weight	$tw(t_k) = \frac{\sum_{i_j \in t_k} w_j}{ t_k }$
Weighted support	$ws(X) = \frac{\sum_{t_k \in t(X)} tw(t_k)}{\sum_{t_k \in T} tw(t_k)}$

In 2003, Feng Tao et al. proposed the WARM algorithm [13]. The proposed algorithm inherits the calculation method of G. D. Ramkumar and extends the unbound weighted itemset belong to $[0, 1]$. To speed up the computation, Feng Tao uses the Lattice structure, which is a structure widely used in mining data without weights.

Description WARM algorithm (2003)	
Apriori property	Satisfy the downward closure property
Weighted items	Weight of each item w_j (unbound $\in [0, 1]$)
Approach	Lattice
Association rule	Weighted binary
Transaction weight	$tw(t_k) = \frac{\sum_{i_j \in t_k} w_j}{ t_k }$
Weighted support	$ws(X) = \frac{\sum_{t_k \in t(X)} tw(t_k)}{\sum_{t_k \in T} tw(t_k)}$

Group 2: In 1998, Chun Hing Cai et al. proposed the MINWAL(O) and MINWAL(W) algorithm [12].

Description MINWAL algorithm (1998)	
Apriori property	Satisfy the downward closure property
Weighted items	Weight of each item $w_j \in [0, 1]$
Approach	Apriori-like
Association rule	Weighted binary
Weighted itemset	$w(X) = \frac{1}{ X } \left(\sum_{i_j \in X} w_j \right)$
Weighted support	$ws(X) = w(X) \times \sup(X)$

In 2005, Unil Yun et al. proposed the WFIM algorithm to inherit the calculation method from Chun Hing Cai. From 2005 to now, Unil Yun's team has more than 20 research works related to weighted association rule mining. However, these works are all based on the WFIM algorithm [14].

Description WFIM algorithm (2005)	
Apriori property	Satisfy the downward closure property
Weighted items	Weight of each item w_j (unbound $\in [0, 1]$)
Approach	FP-Tree (2005), Prefix-Tree (2012)
Association rule	Weighted binary
Weighted itemset	$w(X) = \frac{1}{ X } \left(\sum_{i_j \in X} w_j \right)$
Weighted support	$ws(X) = w(X) \times \sup(X)$

Group 3: In 2011, Zi-guo Huai et al. proposed the WHIUA algorithm [22]. Although nearly a decade has passed, Zi-guo Huai's work has only two citations. According to the author's survey, this is a challenging approach in mining association rule with weighted - the WHIUA algorithm approaches in the direction of NOT satisfying the downward closure property, leading to a very large search space and is not suitable for common pruning strategies.

Description WHIUA algorithm (2011)	
Apriori property	NOT Satisfying the downward closure property
Weighted items	Weight of each item $w_j \in [0, 1]$
Approach	Apriori combines hash table
Association rule	Weighted binary
Weighted itemset	$w(X) = \text{Max}_{i_j \in X} w(i_j)$
Weighted support	$w \sup(X) = w(X) \times \sup(X)$

Group 4: In 2013, Guo-Cheng Lan et al. proposed the PWA algorithm [27]. In order to implement effective pruning strategies and satisfy the downward closure property, the author proposed the upper bound weights of the transactions as well as the itemset.

Description WFIM algorithm PWA (2013)	
Apriori property	NOT Satisfying the downward closure property
Weighted items	Weight of each item $w_j \in [0, 1]$
Approach	Apriori-like
Association rule	Weighted binary
Weighted itemset	$w(X) = \frac{1}{ X } \left(\sum_{i_j \in X} w_j \right)$
Weighted support	$ws(X) = \sum_{X \in t_k, t_k \in TDB} w(X)$
Transaction-weight upper-bounds	$twub(t_k) = \text{Max}_{i_j \in t_k} w(i_j)$
Weighted frequent upper-bound itemset	$wsup(X) = \sum_{X \in t_k, t_k \in TDB} twub(t_k)$

Group 5: In 2015, Xuyang Wei et al. proposed the IWFP algorithm [34]. In order to implement effective pruning strategies and satisfy the downward closure property, the author proposed the upper bound weights of the transactions as well as the itemset.

Description IWFP algorithm (2015)	
Apriori property	NOT Satisfying the downward closure property
Weighted items	Weight of each item $w_j \in [0, 1]$
Approach	Apriori
Association rule	Weighted binary
Weighted itemset	$w(X) = \frac{\prod_{i_j \in X} w(i_j)}{\sum_{i_j \in X} w(i_j)}$
Weighted support	$w \sup(X) = \sup(X) \times \frac{\prod_{i_j \in X} w(i_j)}{\sum_{i_j \in X} w(i_j)}$

In addition, there has been a lot of research related to the mining of weighted association rules in recent years - these algorithms, mostly using the above calculation methods (5 groups) belong with data structures and appropriate pruning strategy or adding constraints.

2. Some domestic research work

Through a survey of the domestic research on mining the weighted frequent itemset of items, there is a prominent research group of Le Hoai Bac and Vo Dinh Bay. From 2009 to now, the above research group has had related works which are summarized as the following table:

Description WIT-FWIs algorithm (2010)	
Apriori property	Satisfy the downward closure property
Weighted items	Weight of each item $w_j \in [0, 1]$
Approach	IT-Tree (2010, 2013, 2017), DBV (2016), Prefix-Tree (2016, 2018, 2020, 2021)
Association rule	Weighted binary
Transaction weight	$tw(t_k) = \frac{\sum_{i_j \in t_k} w_j}{ t_k }$
Weighted support	$ws(X) = \frac{\sum_{t_k \in T(X)} tw(t_k)}{\sum_{t_k \in T} tw(t_k)}$

The research team uses the method to calculate weighted based on group 1 (G.D.Ramkumar, 1998) and use data structures with the appropriate pruning strategies to increase mining efficiency.

In addition, the domestic research groups are interested in research in the mining data number or high-interested samples and there are not many works related to the research problem. Therefore, the authors only focus on surveying the prominent group above.

3. Some algorithms for mining the frequent itemsets with weighted items

Synthesizing some of the frequent itemsets mining works in transaction database with weighted items by time and classified in 5 groups.

Table III, presents the timeline of some frequent itemsets mining algorithm on the transaction database with weighted items from 1998 to present, including information fields: top author's name, algorithm's name, satisfy Apriori property, with the data structure approach, the number of citations of the work [11-50] (Cai [11] to Bui [50]), the year of publication and the group of algorithms.

TABLE III
SYNTHESIZING 40 FREQUENT ITEMSETS MINING WORKS WITH
WEIGHTED ITEMS [11-50] (MARCH 2021)

First author	Algorithms	DCP	Approach	Cite	Yr	Group
Ramkumar	MINWAL[11]	✓	Apriori	645	1998	2
Cai	WIS[12]	✓	Apriori	84	1998	1
Tao	WARM[13]	✓	Lattice	487	2003	1
Yun	WFIM[14]	✓	FP-Tree	72	2005	2
Geng	GTCWFP[15]	✓	FP-Tree	4	2008	2
Ahmed	IWFP[16]	✓	FP-Tree	21	2008	2
Le	WTFWI[17]	✓	IT-Tree	10	2010	1
Jeong	DWFPIM[18]	✓	FP-Tree	6	2010	2
Pears	EWGen[19]	✓	Apriori-like	6	2010	2
Wang	DWC[20]	✓	FP-Tree	1	2011	1
Yun	WAF[21]	✓	FP-Tree	53	2011	2
Huai	WHIUA[22]	×	Apriori + Hash	5	2011	3
Ahmed	IWFP[23]	✓	FP-Tree	64	2012	2
Yun	MWFM[24]	✓	lattice	46	2012	2
Vo	WTFWIs[25]	✓	IT-Tree	123	2013	1
Vo	FWCIs[26]	✓	IT-Tree	13	2013	1
Lan	PWA[27]	×	Apriori-like	9	2013	4
Yun	MCWP[28]	✓	FP-Tree	32	2013	2
Mohan	FWIDIFF[29]	✓	IT-Tree	1	2014	1
Yun	WAC[30]	✓	FP-Tree	20	2014	2
Nguyen	SWITD[31]	✓	IT-Tree	4	2015	1
Lan	PWAI[32]	×	Apriori-like	5	2015	4
Wei	IWFPIM[33]	×	FP-Tree	5	2015	5
Lee	AWMFPs[34]	✓	FP-Tree	15	2016	2
Nguyen	IWSFWIs[35]	✓	IT-Tree + BDV	9	2016	1
Yun	IM_WMFI [36]	✓	FP-Tree	58	2016	2
Qin	WidTMWFIM[37]	✓	IT-Tree		2016	1
Yun	WFPmax[38]	✓	FP-Tree	13	2016	2
Lee	FWI[39]	✓	Prefix-Tree	17	2017	1
Lin	RWFIMMine[40]	✓	SE-Tree	9	2017	2
Vo	WTFWCIDiff[41]	✓	IT-Tree	12	2017	1
Bui	NFWI[42]	✓	Prefix-Tree	4	2017	1
Kiran	WFRIM[43]	✓	FP-Tree		2017	2
Zhao	SWFP[44]	✓	FP-Tree	5	2018	2
Klangwisan	WFRIMIWS[45]	✓	FP-Tree		2018	2
Cengiz	Pre/PostWAR M[46]	✓	Lattice		2019	1 + 2
Dewan	CPTDW[47]	✓	Prefix-Tree		2019	2
Yue	ENFSWI[48]	✓	Prefix-Tree		2019	1
Vo	TFWIN+[49]	✓	Prefix-Tree	4	2020	1
Bui	NFWCI[50]	✓	Prefix-Tree		2021	1

(✓: satisfy Apriori property; ×: not satisfy Apriori property)

IV. DISCUSSION AND RECOMMENDATION

1. Data structure

The authors have investigated 40 works on the frequent itemset mining on transaction database with weighted items [11-50]. Some algorithms follow Apriori approach, while others use the common data structures. The below table is a statistics of algorithms using data structures by each group. Table IV, shows the most used FP-Tree data structure (40.00%), followed by IT-Tree data structure (20.00%) and approach Apriori (15.00%) - this is the basic approach to mining frequent itemsets with weighted items. In addition, the above table also shows that the researches also focus on groups 1 and 2 (90.00%).

TABLE IV
RATIO OF ALGORITHMS IN THE SYNTHESIZE USING THE
COMMON DATA STRUCTURES FOR MINING FI WITH WEIGHTED
ITEMS

Approach	Ratio of algorithm numbers in the synthesize by group (%)					Total (%)
	1	2	3	4	5	
Apriori	2.50	5.00	2.50	5.50		15.00
Lattice	7.50					7.50
SE-Tree		2.50				2.50
Prefix-Tree	12.50	2.50				15.00
IT-Tree	20.00					20.00
FP-Tree	2.50	35.00			2.50	40.00
Total (%)	45.00	45.00	2.50	5.50	2.50	100.00

a) Apriori property

In frequent itemsets mining on unweighted transaction data of items, this is the fundamental property for reducing the search space for the itemsets satisfying the user-specified frequency threshold minsup. However, the observations in Table III show that 90.00% (corresponding to 36 works in groups 1 and 2) algorithms use satisfying Apriori property in the process of pruning and reducing the space generated frequent itemset on transaction databases with weighted items; only 10.00% (corresponding to 4 works in groups 3, 4 and 5) of algorithms solve the problem of mining frequent itemsets in transaction databases with weighted items that does NOT satisfy the Apriori property (the downward closure property) - this is a big challenge in mining frequent itemsets because the search space is very large. This requires researchers to come up with the storage techniques as well as the strategies to reduce the search space.

In addition, the algorithms in groups 3, 4 and 5 approach in the direction of NOT satisfying the Apriori property (only 04/40 works) together with the proposed algorithm base the Apriori algorithm.

2. The method to calculate the weighted itemset

In this section, the authors discuss the measures to calculate the 5 groups of algorithms above: transaction weight, weighted itemset, weighted support, transaction-weight upper-bound, weighted frequent upper-bound itemset and support of interest in algorithms.

TABLE V
DESCRIPTION OF THE MEASURES IN THE CALCULATION
METHOD OF THE ALGORITHM GROUP

The measures including the calculation method of each group of algorithms	Algorithm group				
	1	2	3	4	5
Transaction weight	✓				
Weighted itemset		✓	✓	✓	✓
Weighted support	✓	✓	✓	✓	✓
Support		✓	✓		✓
Transaction-weight upper-bound				✓	
Weighted frequent upper-bound itemset				✓	

Table V, shows the weights frequently calculated in all 5 groups of algorithms - this is the basic measure for

evaluating whether itemsets are frequent or not. Only group 1 considers transaction weight; group 4 proposes adding the transaction-weight upper-bound and the weighted frequent upper-bound itemset; The itemset weights are calculated in groups 2-4. In addition, the frequent itemsets are the support of groups 2,3 and 5 - this is an important measure to show the frequency of frequent itemsets appearance in transaction database. In groups 1 and 4, the support itemsets are not used - the frequent itemsets mined from these 2 groups are difficult to assess as frequent or not.

The weight itemset is calculated in groups 2-4: in groups 2 and 4, this measure is averaged according to the weight of items in itemsets - losing the significance of the weights of items in data transactions; In group 3, this measure is calculated according to dominance - that is, the weighted of itemset is largest weight of the items in the itemset (representing the significance of weight or the important of items in the transaction database).

In 2013, Lan et al. proposed PWA algorithm [27] - the first algorithm in group 4: adding a measure of the transaction-weight upper-bound and the weighted frequent upper-bound itemsets to help prune the search space faster.

In group 5, Wei et al. proposed IWFP algorithm [34]. The algorithm for calculating the itemset weight is equal to the ratio between the product of the weighted items and sum of the weighted items in itemset - with this measure, the authors have not yet explained the meaning/role of the weights in the work.

In addition, basing on the method of calculating the measures of the 5 groups of algorithms, the authors have the following comments: from group 1 to 4, when changing to the traditional form of mining frequent itemset (the weighted items are the same) then the measures are corresponding - the weighted frequent becomes frequent; especially for group 5, the uncorrelated weighted frequent is the frequent when applied to the unweighted data transactions of items.

3. Recommendations

In the above section, the authors presented and discussed the frequent itemsets mining works on the weighted data transactions of items, which is surveyed from data structure, the search space pruning/reducing strategy and the method to calculate the relevant metrics in the mining process, the authors have the following recommendations:

– First, the algorithms for mining frequent itemsets on transaction database with weighted items, the experiment evaluation needed to be further compared with the traditional efficient frequent itemsets mining algorithms (that means assign the weight of items to equal 1) to show that the proposed algorithm is really efficient;

– Second, it is necessary to choose a method to calculate the appropriate measures - the importance is that the measures represent the weighted role/meaning of each items, as well as the frequency of occurrence of the itemsets and especially the metrics to use that must be suitable when changing to unweighted one (or the weight of items is equal to 1);

– Third, according to the recommendations above - the authors propose to correct the calculation formula for group 5, specifically adding $|X|$ in weighting itemset;

Description IWFP algorithm (2015)	
Weighted itemset	$w(X) = X \times \frac{\prod_{i_j \in X} w(i_j)}{\sum_{i_j \in X} w(i_j)}$
Weighted support	$w \sup(X) = \sup(X) \times w(X)$

– Fourth, researchers need to focus on proposing a frequent itemsets mining algorithm on transaction database with weighted items following the approach that does NOT satisfy the Apriori property - this is really a big challenge in this type of problem. Currently, the proposed research works only approach the Apriori algorithm and have no effective algorithm to solve this problem;

In addition, the authors also found that the approach in mining frequent itemsets on transaction database with weighted items is also based on the aggregated data structures used in traditional frequent itemsets mining. Then, the efficient algorithms in traditional mining frequent itemsets can be used for group 1 and 2 (satisfy Apriori property) and only need to calculate more measures related to the weights of items.

V. CONCLUSIONS AND FUTHER DEVELOPMENTS

In this article, the authors present an overview of common data structures used in mining frequent itemsets on the binary transaction database and a summary of some algorithms for mining frequent itemsets on transaction database weighted items in the direction of data structure analysis, search strategy on generate space, and method of calculating the related measures in researches over the past twenty years. The authors also give recommendations, and help researchers in data mining have enough knowledge when choosing an appropriate technical solution for the problem of mining frequent itemset on transaction database with weighted items.

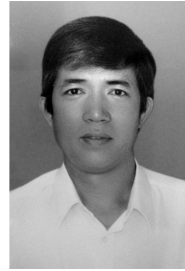
From the summary and recommendations on some common frequent itemsets mining algorithms on transaction data with weighted items presented in Parts II and IV: In future, the authors will expand the research on some data structures and algorithms in order to mine frequent itemsets on large transaction database with weighted items, as well

as parallelize the proposed algorithm in order to quickly mine the frequent itemsets on multi-core processors, the distributed computing systems such as Hadoop, Spark, etc.

REFERENCES

- [1] R. Agrawal, T. Imilienski, A. Swami, "Mining association rules between sets of large databases", *Proc of the ACM SIGMOD Int Conf on Management of Data*, Washington, DC, 207-216, 1993.
- [2] R. Agrawal, Srikant R., "Fast Algorithms for Mining Association Rules in Large Databases", *VLDB*, Chile, 487-499, 1994.
- [3] J. Han, J. Pei, Y. Yin, "Mining Frequent Patterns without Candidate Generation", *SIGMOD Conf. 2000*, 1-12, 2000.
- [4] J. Liu, Y. Pan, K. Wang, J. Han, "Mining frequent item sets by opportunistic projection", *KDD 2002*, 229-238, 2002.
- [5] N. Pasquier, Y. Bastide, R. Taouil, L. Lakhal, "Discovering frequent closed itemsets for association rules", *Proceedings of the 5th Int Conf on Database Theory*, LNCS, Springer-Verlag, Jerusalem, Israel, 398-416, 1999.
- [6] M.J. Zaki, C. J. Hsiao, "CHARM: An Efficient Algorithm for Closed Association Rule Mining", *Computer Science*, Rensselaer Polytechnic Institute 1-20, 1999.
- [7] D. Lin, Z. M. Kedem, "Pincer-Search: A New Algorithm for Discovering the Maximum Frequent Itemset", *EDBT Conference Proceedings*, 105-110, 1998.
- [8] R. J. Bayardo, "Efficiently mining long patterns from databases", *In Proc. The ACM SIGMOD Int. Conf. on Management of data*, 85-93, 1998.
- [9] R. Agarwal, C. Aggarwal, V. V. V. Prasad, "Depth First Generation of Long Patterns", *In Proc. ACM SIGMOD*, 108-118, 2000.
- [10] H. Phan, B. Le, "MAXLEN-FI: A fast algorithm for mining maximum length frequent itemsets", *DLU Journal of Science*, 8(2), 2018, 108-123.
- [11] G. D. Ramkumar, S. Ranka, S. Tsur, "Weighted Association Rules: Model and Algorithm", *Proc. ACM SIGKDD*, 1998, 1-13.
- [12] C. H. Cai, A. W. C. Fu, C. H. Cheng, W. W. Kwong, "Mining association rules with weighted items", *Proceedings. IDEAS'98. International Database Engineering and Applications Symposium (Cat. No.98EX156)*, Cardiff, UK, 68-77, 1998.
- [13] F. Tao, F. Murtagh, M. M. Farid, "Weighted association rule mining using weighted support and significance framework", *SIGKDD'03*, 661-666, 2003.
- [14] U. Yun, J. J. Leggett, "WFIM: Weighted Frequent Itemset Mining with a weight range and a minimum weight", *SDM 2005*, 636-640, 2005.
- [15] R. Geng, X. Dong, X. Zhang, W. Xu, "Efficiently Mining Closed Frequent Patterns with Weight Constraint from Directed Graph Traversals Using Weighted FP-Tree Approach", *Computing Communication Control and Management 2008. CCCM'08. ISECS International Colloquium on*, 3, 399-403, 2008.
- [16] C. F. Ahmed, S. K. Tanbeer, B. S. Jeong, Y. K. Lee, "Mining Weighted Frequent Patterns in Incremental Databases", *In: Ho TB., Zhou ZH. (eds) PRICAI 2008: Trends in Artificial Intelligence. PRICAI 2008. Lecture Notes in Computer Science*, vol 5351. Springer, Berlin, Heidelberg, 2008.
- [17] B. Le, H. Nguyen, B. Vo, "Efficient Algorithms for Mining Frequent Weighted Itemsets from Weighted Items Databases", *RIVF 2010*, 1-6, 2010.
- [18] B. S. Jeong, A. Farhan, "Efficient Dynamic Weighted Frequent Pattern Mining by using a Prefix-Tree", *The KIPS Transactions:PartD*, 17D, 253, 2010.
- [19] R. Pears, Y. S. Koh, G. Dobbie, "EWGen: Automatic Generation of Item Weights for Weighted Association Rule Mining", *ADMA (1)*, 36-47, 2010.
- [20] B. Wang, Z. Yuanpan, G. Feng Guo, "Mining weighted closed itemsets directly for association rules generation under weighted support framework", *IEEE 3rd Inter Conf on Communication Software and Networks*, 145-149, 2011.
- [21] U. Yun, K. H. Ryu, "Approximate weighted frequent pattern mining with/without noisy environments", *Knowl. Based Syst*, 24(1), 2011, 73-82.
- [22] Z. G. Huai, M. H. Huang, "A weighted frequent itemsets Incremental Updating Algorithm base on hash table", *2011 IEEE 3rd International Conference on Communication Software and Networks*, Xi'an, 201-204, 2011.
- [23] C. F. Ahmed, S. K. Tanbeer, B. S. Jeong, Y. K. Lee, H. J. Choi, "Single-pass incremental and interactive mining for weighted frequent patterns", *Expert Systems with Applications*, 39, 2012, 7976.
- [24] U. Yun, H. Shin, K. H. Ryu, E. Yoon, "An efficient mining algorithm for maximal weighted frequent patterns in transactional databases", *Knowl. Based Syst*, 33, 2012, 53-64.
- [25] B. Vo, F. Coenen, B. Le, "A new method for mining Frequent Weighted Itemsets based on WIT-trees", *Expert Systems with Applications*, 40(4), 2013, 1256-1264.
- [26] B. Vo, N. Y. Tran, D. H. Ngo, "Mining Frequent Weighted Closed Itemsets", *Advanced Computational Methods for Knowledge Engineering 2013*, 379-390, 2013.
- [27] G. C. Lan, T. P. Hong, H. Y. Lee, and C. W. Lin, "Mining Weighted Frequent Itemsets", *In Proceedings of the 30th workshop on Combinatorial Mathematics and Computation Theory (Alg'30)*, 85-89, 2013.
- [28] U. Yun, K. H. Ryu, "Efficient mining of maximal correlated weight frequent patterns", *Intell. Data Anal*, 17(5), 2013, 917-939.
- [29] A. Mohan, R. Visakh, "Weighted Frequent Pattern Mining using RDD, the Basic Spark Abstraction", *In Proceedings of the 2014 International Conference on Information and Communication Technology for Competitive Strategies (ICTCS '14)*. Association for Computing Machinery, New York, NY, USA, Article 81, 1-5, 2014.
- [30] U. Yun, E. Yoon, "An Efficient Approach for Mining Weighted Approximate Closed Frequent Patterns Considering Noise Constraints", *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 22(6), 2014, 879-912.
- [31] D. H. Nguyen, B. Vo, T. H. M. Nguyen, T.-P. Hong, "An Improved Algorithm for Mining Frequent Weighted Itemsets", *IEEE International Conference on Systems, Man, and Cybernetics*, 2579-2584, 2015.
- [32] G. C. Lan, T. P. Hong, H. Y. Lee, C. W. Lin, "Tightening upper bounds for mining weighted frequent itemsets", *Intell. Data Anal*, 19(2), 2015, 413-429.
- [33] X. Wei, Z. Li, T. Zhou, H. Zhang, G. Yang, "IWFPM: Interested Weighted Frequent Pattern Mining with Multiple Supports", *Journal of Software*, 10, 2015, 9-19.
- [34] G. Lee, U. Yun, H. Ryang, D. Kim, "Approximate Maximal Frequent Pattern Mining with Weight Conditions and Error Tolerance", *Int. J. Pattern Recognit. Artif. Intell*, 30(6), 2016, 1-42.
- [35] H. Nguyen, B. Vo, M. Nguyen, W. Pedrycz, "An efficient algorithm for mining frequent weighted itemsets using interval word segments", *Appl. Intell*, 45(4), 2016, 1008-1020.
- [36] U. Yun, G. Lee, "Incremental mining of weighted maximal frequent itemsets from dynamic databases", *Expert Syst. Appl*, 54, 2016, 304-327.
- [37] Q. Qin, L. Tan, "An Efficient Mining Algorithm for Maximal Weighted Frequent Patterns Based on WIdT-Trees", *IDEAL*

- 2016, 596-605, 2016.
- [38] U. Yun, G. Lee, K. M. Lee, "Efficient representative pattern mining based on weight and maximality conditions", *Expert Systems*, 33(5), 2016, 439-462.
 - [39] G. Lee, U. Yun, K. H. Ryu, "Mining Frequent Weighted Itemsets without Storing Transaction IDs and Generating Candidates", *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 25(1), 2017, 111-144.
 - [40] J. C. W. Lin, W. Gan, P. F. Viger, H. C. Chao, T. P. Hong, "Efficiently mining frequent itemsets with weight and recency constraints", *Applied Intelligence*, 47, 2017, 769-792.
 - [41] B. Vo, "An Efficient Method for Mining Frequent Weighted Closed Itemsets from Weighted Item Transaction Databases", *J. Inf. Sci. Eng.*, 33(1), 2017, 199-216.
 - [42] H. Bui, B. Vo, H. Nguyen, T. A. N. Hoang, T. P. Hong, "A weighted N-list-based method for mining frequent weighted itemsets", *Expert Syst. Appl.*, 96, 2018, 388-405.
 - [43] R. U. Kiran, A. Kotni, P. K. Reddy, M. Toyoda, S. Bhall, M. Kitsuregawa, "Efficient discovery of weighted frequent itemsets in very large transactional databases: A re-visit", *Proc. IEEE Int. Conf. Big Data (Big Data)*, 723-732, 2018.
 - [44] X. Zhao, X. Zhang, P. Wang, S. Chen, Z. Sun, "A weighted frequent itemset mining algorithm for intelligent decision in smart systems", in *IEEE Access*, 6, 2018, 29271-29282.
 - [45] K. Klangwisan, K. Amphawan, "Efficient weighted-frequent-regular itemsets mining using interval word segments structure", *Knowledge and Smart Technology (KST) 2018 10th International Conference on*, 59-67, 2018.
 - [46] A. B. Cengiz, K. U. Birant, D. Birant, "Analysis of Pre-Weighted and Post-Weighted Association Rule Mining", *Innovations in Intelligent Systems and Applications Conference*, Izmir, Turkey, 1-5, 2019.
 - [47] U. Dewan, C. F. Ahmed, C. K. Leung, R. A. Rizvee, D. Deng, J. Souza, "An efficient approach for mining weighted frequent patterns with dynamic weights", *ICDM 2019*, 13-27, 2019.
 - [48] Z. Yue, J. Sun, R. Liu, "Mining Frequent Weighted Itemsets Using Extended N-List and Subsume", *International Conference on Robots & Intelligent System (ICRIS)*, 513-516, 2019.
 - [49] B. Vo, H. Bui, T. Vo, T. Le, "Mining top-rank-k frequent weighted itemsets using WN-list structures and an early pruning strategy", *Knowledge-Based Systems*, 201-202, 2020, 106064-106075.
 - [50] H. Bui, B. Vo, T. A. N. Hoang, U. Yun, "Mining frequent weighted closed itemsets using the WN-list structure and an early pruning strategy", *Appl Intell*, 51(3), 2021, 1439-1459.



Huan Phan is a Ph.D. Student at the Faculty of Mathematics and Computer Science, Vietnam National University Ho Chi Minh City - University of Science, Vietnam. His research interests: artificial intelligence, data mining and high performance computing.

Email: huanphan@hcmussh.edu.vn



Bac Le is a Professor at the Faculty of Information Technology, Vietnam National University Ho Chi Minh City - University of Science, Vietnam. His research interests: artificial intelligence, soft computing, machine learning, data mining and data science.

Email: lhbac@fit.hcmuns.edu.vn

Robust Watermarking Method using QIM with VSS for Image Copyright Protection

Thanh Ta Minh, Dan Giang Ngoc
Le Quy Don Technical University

Correspondence: Ta Minh Thanh, thanhmt@lqdtu.edu.vn

Communication: received 30 October 2020, revised 4 April 2021, accepted 31 May

Digital Object Identifier: 10.32913/mic-ict-research.v2021.n1.971

Abstract: This paper proposes a new watermarking method for digital image by composing the DWT-QIM based embedding with visual secret sharing (VSS) method. Firstly, the watermark image is separated into n shares by using the k -out-of- n method, also called (k,n) visual secret sharing. One of the s is embedded into the original image for copyright protection. The other $(n-1)$ shares are registered with Vietnam Copyright Department. When a dispute happens, the verifier can extract the watermark information from the watermarked image, then, decode it with $(k-1)$ shares chosen from $(n-1)$ shares to retrieve the copyright information. Our experimental results show that our proposed method works efficiently on protecting the copyright of digital images.

Keywords: Image watermarking, copyright protection, digital content, visual secret sharing - VSS, (k,n) sharing method.

I. INTRODUCTION

1. Overview

The industrial revolution 4.0 has changed people's perception, as well as the way people approach and communicate with each other on the Internet. The advent of software and new devices such as digital cameras, high-quality scanners, printers, digital audio recorders, etc. has allowed the vast consumer world to create, process and enjoy multimedia data. Along with the development of social networks, the Internet has become a "virtual society" where information exchange takes place in all fields such as politics, military, national defense, economy, and commercial. It is in such an open and comfortable environment that issues such as piracy, information misrepresentation, and unauthorized access to digital multimedia arise. Therefore, a solution to protect digital products, preserve information, and protect copyright for digital products is urgently required. Among the current solutions, digital watermarking is one that is being researched because of its effectiveness and suitability for copyright protection of multimedia products.

In our recent proposal [23], we mentioned a watermarking method based on a combination of Discrete Wavelet

Transform (DWT) and cryptographic techniques VSS. The experimental results with that proposal have demonstrated that the above combination has enhanced the imperceptibility after embedding the watermark, as well as the stability of the watermark against some common digital transformations. However, that solution still has some limitations such as needing the original image in the process of extracting copyright information, which is known as non-blind watermarking.

Inheriting the results from previous studies, we continue to propose an improved watermarking method, with the combination of DWT-QIM and VSS. The difference in this algorithm is that the embedding and extraction of watermarks will be conducted in a smarter way, closer to the actual application through the use of quantization index modulation (QIM) to adjust image quality and persistence of copyright information.

Normally, for a digital watermarking method to be effective, it needs to be hidden and robust against common image manipulations such as compression, filtering, rotation, scaling, collusion attacks, and so on. Many other digital signal processing operations also need to be guaranteed. Current digital image watermarking techniques can be classified into two groups: spatial domain watermarking techniques and frequency domain watermarking techniques [14]. Compared to the spatial domain watermarking technique, the frequency domain watermarking technique proves more effective in achieving the hiddenness and robustness requirements of the digital watermarking algorithm.

Commonly used frequency domain transforms include DWT, Discrete Cosine Transform (DCT) and Discrete Fourier Transform (DFT). The proposed sustainable watermarking techniques on the frequency domain of the image have attracted attention in recent years because of their feasibility and effectiveness in copyright protection applications such as watermarking techniques on frequency domain DCT[1-3], frequency domain DWT [4-7], and frequency

domain DFT [8-9]. Frequency domain watermarking solutions have shown advantages in creating sustainable watermarking solutions, suitable for digital product copyright protection applications. Frequency domain watermarking techniques are often chosen to suit products affected by techniques such as compression, which will be applied to the DCT frequency domain watermarking solution to improve the image quality. However, DWT is used in digital image watermarking more often because of its spatial localization and multi-resolution characteristics. Its solution is similar to the theoretical model of the human visual system (HVS). Better performance improvements in DWT-based digital image watermark algorithms can be achieved in many ways to suit different digital product objects such as digital music, digital movies, and digital images.

To balance the improvement of digital image quality after watermark embedding and the stability of watermark information, research on extended frequency domains such as q -DCT [10], q -DWT [11], or q -SVD [12] is also proposed to balance these two goals. However, adjusting the parameter value q to satisfy the condition of balancing both criteria is difficult to achieve. The authors had to evaluate the image quality after embedding and the durability of the watermark to choose the most appropriate value for the parameter q . Although the computational complexity of these proposals still needs improvement, it opens up a new frequency domain research direction applied in digital watermarking proposals.

Another technique to protect digital image copyright is zero watermarking technique, in which the watermark information will not be directly embedded in the digital image, but the stable features of the digital image will be extracted for encoding with copyright information to generate information registered with the Copyright Office [13]. However, zero watermarking depends heavily on the properties and characteristics of the digital image. For images with large “complexity”, the extracted feature will be resistant to image attacks.

In this paper, we focus our research on embedding and extracting watermark algorithms based on QIM combined with VSS to improve the inequalities. Our proposed method can recognize and enhance the stability of the watermark after it is extracted, and at the same time ensure that the watermark can be sufficiently resistant to normal image transformations.

2. Our contribution

We summarize the problem to be solved in the proposal of the new digital watermarking technique as follows:

Most watermarking application solutions are used to protect digital product and copyright verification. However,

copyright information, after being embedded in digital products, requires resistance to image attacks to verify ownership when a dispute occurs. Our paper proposes an idea to combine visual secret sharing (VSS) with watermarking techniques to improve the stability of watermarks after being hidden in digital images. The use of an intuitive encryption system combined with scrambled encryption has helped our proposal improve the security of copyright information, creating reliability for users when participating in the solution.

3. Roadmap

In this paper, we organize the content as follows: The knowledge related to our proposed method is presented in Section II. The architecture and detailed description of the proposed model are presented in Section III. The experimental and evaluation results are presented in Section IV. Section V presents conclusions and future research proposals.

II. RELATED WORKS

The watermarking method is based on the DWT, which has been widely used in digital signal processing applications. The method in [23] effectively utilizes the property of discrete wavelet transform to improve the robustness of watermark information. However, when extracting watermarks from hidden images, it is necessary to have the same initial data as the original image. In fact, we are not always qualified to collect such data. Therefore, in this paper, we propose a better solution for the DWT transformation, which is to use QIM.

Visual Secret Sharing, is an encryption technique in which visual information (images, text, etc.) is encrypted so that decryption can be performed based on the human visual system. In this section, we briefly introduce these two techniques and outline their relevance to the implementation of watermarking on digital images.

1. Visual secret sharing - VSS

The overview of VSS can be shown in Figure 1. VSS may share confidential information to protect the confidentiality of copyright information [15].

The idea of secret sharing was proposed by Adi Shamir [16] and G. Blakley [18] in 1979. Four years later, another method of secret sharing was proposed by Asmuth. and Bloom [19] in 1983. Accordingly, Shamir’s scheme is based on polynomial interpolation, Blakley’s scheme is based on hyperplane geometry, and Asmuth and Bloom’s scheme is based on the Chinese remainder theorem.

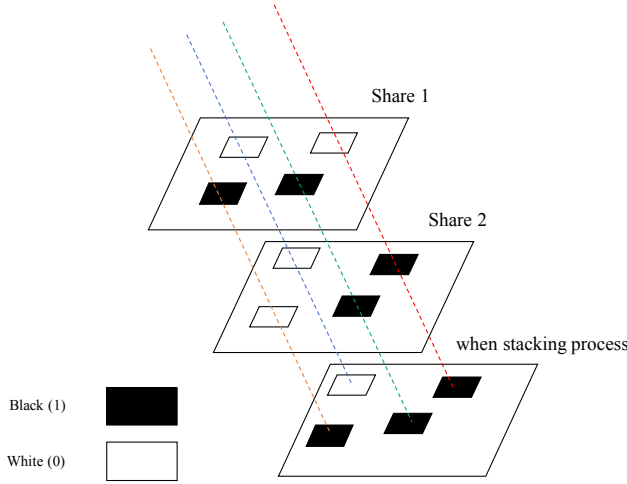


Figure 1. The human visual system functions as an OR-operator

In particular, the secret sharing scheme of Adi Shamir is considered the most potential in the field of image copyright protection. The basis of this proposal is to divide the image into secret shares (also known as Shares or Shadows). Each of these shares is treated as a random image generated from the original image. To reconstruct the original image, we simply combine these shares together as shown in Figure 2.

In this paper, we use the (k, n) secret sharing scheme based on random sequence (Random Sequence) [20]. In the schema (k, n) , suppose we have a set P consisting of n elements, a secret image S (Secret Image) encrypted in these n elements. For each element in set P , we get a secret share. For any k secret share fragment (with $2 < k < n$) stacked, the secret image S will be reconstructed.

The secret sharing scheme (k, n) is based on a random sequence, which is a secret image scattering technique using randomly arranged strings of $[0, 1]$ values. Where, a random

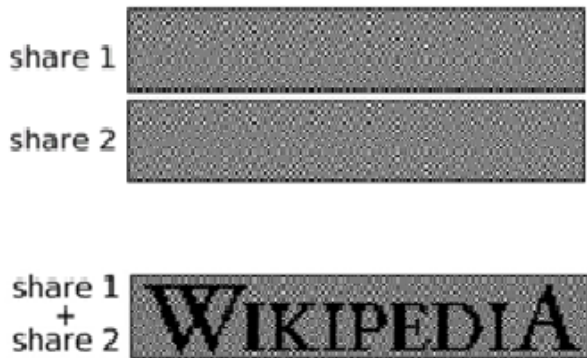


Figure 2. Adi Shamir's Secret Sharing Scheme

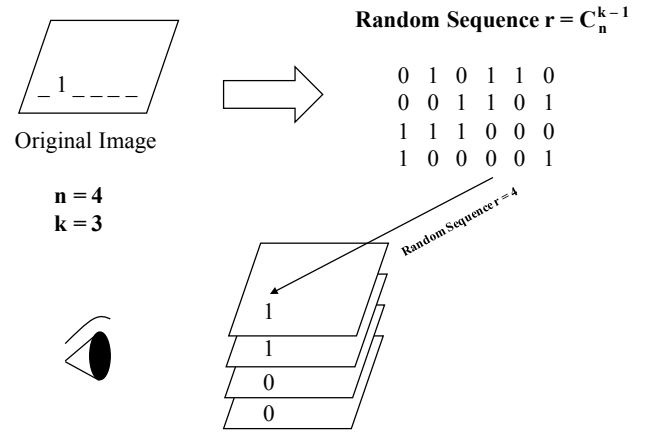

 Figure 3. String value at position $pix(w, h)$

 Figure 4. Sample shares created with $k = 6$ and $n = 10$

string is composed of $(n - k + 1)$ value "1" and $(k - 1)$ the value "0". For a secret image distributed with the scheme (k, n) , the number of random strings (ns) is calculated by the convolution $(k - 1)$ of the n part death.

When viewed from the side with the secret shares stacked, the string value at $pix(w, h)$ (with $pix(w, h)$ on the original image having the value "1") is made up of $(n - k + 1)$ value "1" and $(k - 1)$ value "0" as visually described in Figure 3. The result after scattering the watermark with $k = 6$ and $n = 10$ is shown in Figure 4.

2. Discrete wavelet transform - DWT

The DWT on the digital image [17, 21-23] is shown in Figure 5. When first two-dimensional DWT transform on image I , the image I will be analyzed into frequency domains LL (Low-Low), LH (Low-High), HL (High-Low) and HH (High-High). These frequency domains will be used to embed watermark information based on the purpose of data hiding, invisible watermarking, or visible watermarking. Normally, the DWT frequency transform domain is used to embed copyright information, after the reverse transformation of IDWT (Inverse DWT), the copyright information will be distributed over the entire image surface, creating a persistence for copyright information hidden in images. Thus, from 4 subimages (LL, HL, LH, HH) shown in Figure 5, through the IDWT inverse transformation, we receive image I .

In this paper, we apply the DWT frequency transformation to obtain the frequency domain of pixels. After that,

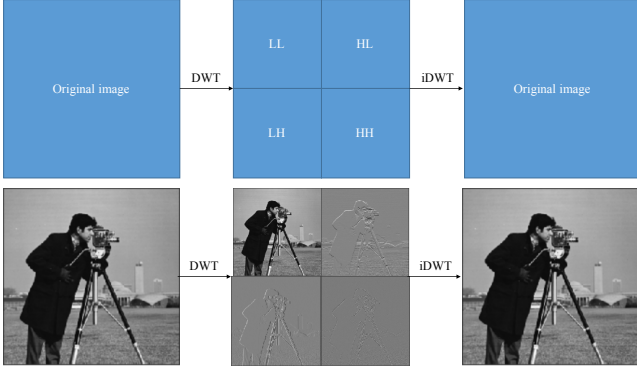


Figure 5. Discrete wavelet transform on digital image

we proceed to embed the distributed watermark information on the DWT frequency region to ensure the stability of the embedding method.

3. DWT based watermarking methods

Recently, studies on frequency domain watermarking techniques are being widely proposed with the aim of overcoming the limitations of traditional frequency domain watermarking methods. Research in this field has thought of a solution that is to combine multiple watermarking methods.

In a study by Nawaf Hazim Barnouti, Zaid Saeb Sabri and Khandoun L. Hameed [25], the authors proposed a watermarking method based on the combination of two transformations - DWT and DCT. In the paper, the authors have shown that the combination of two transformations, DWT and DCT, has brought a watermarking method that has the advantages of both transformations and has improved the performance compared to the simple methods. However, the implementation process of the authors has converted the original image from color image to gray image. This reduces the practicality of the proposal. Also, it has not been proven that this method can improve the robustness of the watermark against conventional image attacks.

Our most recent proposal [23] also proposed a watermarking method that combines DWT transformation and Visual Secret Sharing. That study also evaluated the effect of the above combination through two properties "Unrecognizability" and "Sustainability". The output helps to ensure the stability of the watermark against common attacks, and at the same time, it is undetectable to the human eye after the watermark is embedded into the digital image. However, we realize that the above watermarking method is still not perfect, specifically, we still need to use the original data when extracting the watermark. This also reduces the level of applicability of the method.

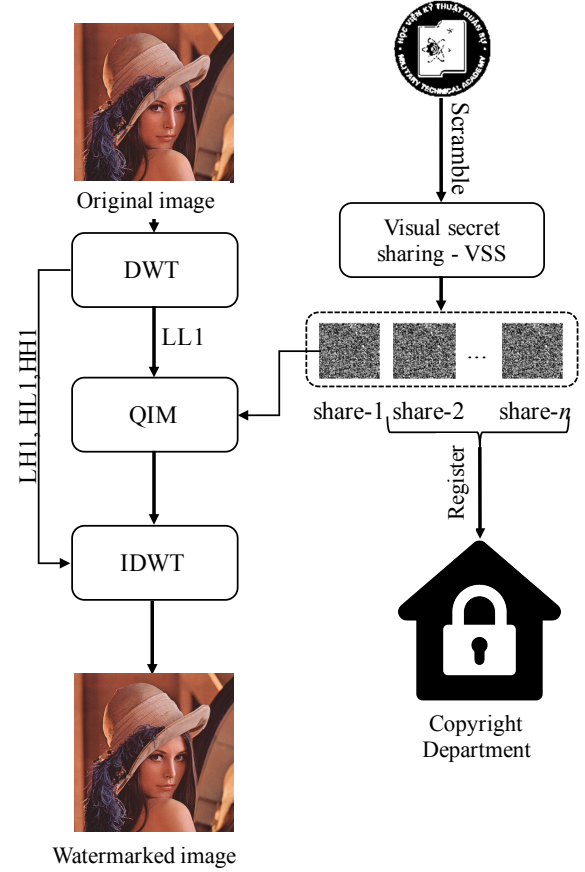


Figure 6. Watermark embedding algorithm combined with VSS

III. PROPOSED QIM WATERMARKING ALGORITHM COMBINED WITH VSS

In this section, we present a combination that is considered to be smarter and more practical with the combination of DWT-QIM and VSS. QIM combined with watermark scheme on DWT frequency domain is not a new idea, but in our paper, watermarking combined with QIM on DWT frequency domain will be applied to embedding distributed watermark using VSS to form a schema with the properties of zero watermark to overcome the limitations mentioned in section I.2.2. With this combination, we have overcome the limitations of the proposed method of watermarking combined with traditional DWT and VSS.

1. Watermark embedding algorithm combined with VSS

The input data of the embedding process includes: 01 cover image and 01 binary image to make the copyright watermark. The watermarked image is processed with scramble and visual cryptographic techniques to disperse the scrambled watermark image into different shares. Therefore, each share will not display any information. One share

will be selected to embed in the cover image, and the remaining shares will be sent to a 3rd party in charge of copyright protection. After the watermark embedding process, we obtain the watermarked image. The watermark embedding process is visually described in Figure 6. The watermark embedding process is carried out according to the following steps:

Step 1: Input: Cover image, copyright logo.

Step 2: Scramble the copyright logo image by randomly shuffling the pixel positions with the key being an array containing $(width * height)$ zero values duplicates. The elements of this array are the result of randomly shuffling the numbers from 1 to $(width * height)$, where $width$ and $height$ are width and height height of watermarked image, respectively.

Step 3: Distribute the scrambled logo using visual cryptography using the (k, n) secret sharing scheme and select a secret share to be watermarked.

Step 4: Perform a DWT-1-level transformation on any color (RGB) channel of the containing image to obtain 4 subbands (LL1 - LH1 - HL1 - HH1).

Step 5: Choose one subband (Here, we choose LL band) and embed watermark with logic based on the method in the paper [24]:

$$S_{ij}^* = \begin{cases} S_{ij} - (S_{ij} \bmod Q) + 0.75 * Q, & \text{if } W(i, j) = "1" \\ S_{ij} - (S_{ij} \bmod Q) + 0.25 * Q, & \text{if } W(i, j) = "0" \end{cases} \quad (1)$$

Given:

- S_{ij} is the value at (i, j) on the sub-band of the original image.
- S_{ij}^* is the value at the corresponding position on the image after embedding the watermark.
- Q is the quantization step.

After embedding the $W(i, j)$ information, the value of S_{ij} will be shifted around the origin of S_{ij} based on the quantum step value Q . The quality of the image after embedding the information will depend on the magnitude of the quantum step value Q . At the same time, the stability of the watermark will also be affected by the quantum step value Q .

Step 6: Perform Inverse DWT transformation (IDWT) to get image after embedding watermark.

We use the LL subband because after performing the DWT transform, the values of the DWT coefficients carrying the "energy" of the whole image are concentrated mainly in the LL subband. The selection of LL sub-band for information embedding will affect the quality of the image after embedding, but the value domain used for embedding will be richer because the values of non-zero

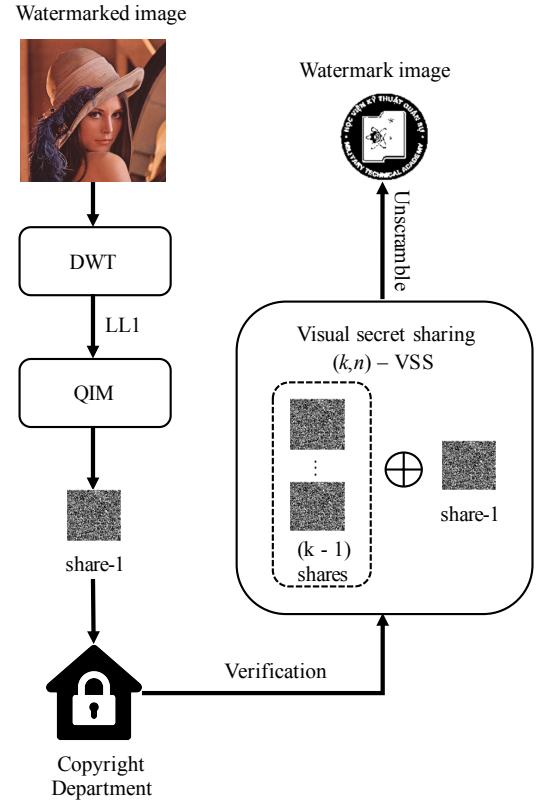


Figure 7. Watermark extraction algorithm combined with VSS

DWT coefficients are mostly concentrated in the sub-band LL [27, 28]. On the other hand, we use QIM embedding algorithm, so the effect of image quality is minimized after embedding information.

2. Watermark extraction algorithm combined with VSS

The input data for the watermark extraction process is the watermarked image, the output of the watermark extraction process is the original watermark image (the logo before being distributed). The watermark extraction process is described in Figure 7.

The watermark extraction process is carried out according to the following steps:

Step 1: Input: The watermarked image

Step 2: Perform a discrete wavelet transformation on the watermarked color channel and select subband LL1.

Step 3: Extract the watermark according to the following equation:

$$\begin{aligned} \text{If } (S_{ij}^* \bmod Q) > Q/2 \text{ then } W^*(i, j) &= "1", \\ \text{If } (S_{ij}^* \bmod Q) < Q/2 \text{ then } W^*(i, j) &= "0", \end{aligned}$$

where:

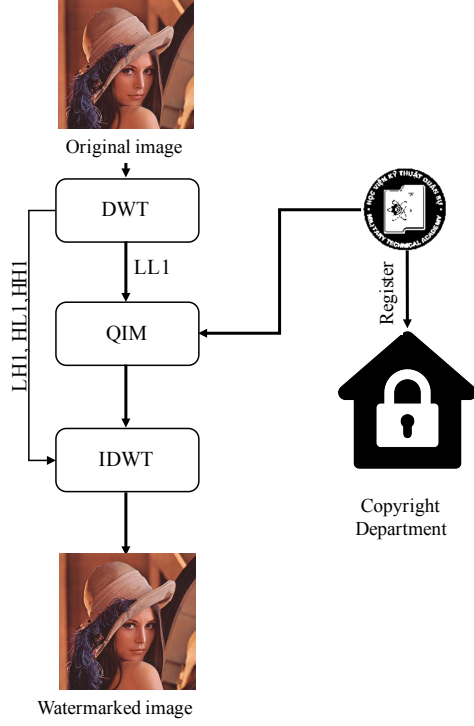


Figure 8. Watermark embedding using only DWT-QIM

- S_{ij}^* is the value at position (i, j) on the watermarked subband.
- $W^*(i, j)$ is the value at (i, j) of the extracted watermark.
- Q is the quantization step.

Step 4: Combine the watermark extracted in Step 3 with another $(k - 1)$ shared secret fragments to reproduce the logo.

Step 5: Perform unscramble recovery to obtain the complete copyright logo by re-positioning the pixels according to the key generated during the watermark embedding process.

3. Watermarking/embedding algorithm using only DWT-QIM

To compare the effectiveness of the proposed solution, we perform a traditional DWT domain watermarking method (called DWT-QIM solution) to evaluate and determine the highlights and feasibility of the method.

a. Watermarking embedding process:

The process of embedding the watermark using only DWT-QIM is described in Figure 8. The process of embedding the watermark is done in the same way as in III.1, but steps 2 and 3 are omitted.

b. Watermark extraction process:

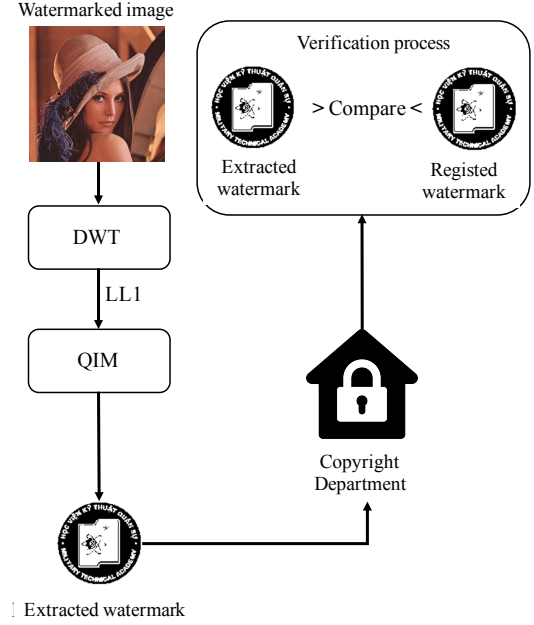


Figure 9. Watermark extraction using only DWT-QIM

The process of watermark extraction using only DWT-QIM is described in Figure 9. The process of watermark extraction is done in the same way as in III.2, but steps 4 and 5 are omitted to recover the watermark image after extraction.

IV. EXPERIMENTAL RESULTS AND EVALUATION

1. Evaluation measure

We conduct experiments and compare the results of two methods: the embedding/watermarking method combining DWT-QIM with VSS and the embedding/watermarking method using only DWT-QIM to compare the effectiveness of our method.

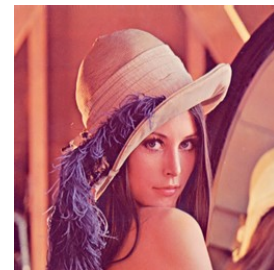


Figure 10. Lena image

We evaluate the results of the experiments through two properties: "Unrecognizability" represented by the peak signal/noise ratio - PSNR (Peak Signal to Noise Ratio) measured in decibel with formula (2) and "Stability" of the watermark after extraction, represented by the value NC (Normalization Correlation) calculated by the formula (4).

$$PSNR = 20 \log_{10} \frac{MAX(I)}{\sqrt{MSE}}, \quad (2)$$

where, $MAX(I) = 255$ and the value of MSE is calculated by formula (3):

$$MSE = \frac{1}{p \times q} \sum_{i=1}^p \sum_{j=1}^q (I(i, j) - I_w(i, j))^2, \quad (3)$$

where $I(i, j)$, $I_w(i, j)$ is the pixel value of the original and watermarked image at position (i, j) . The value $p \times q$ is the size of the image.

For PSNR, the higher the value of the PSNR, the more "unrecognizability" is guaranteed. If the value of PSNR ≥ 35 dB, the human eye cannot distinguish between the original image and the image that has been watermarked.

$$NC = \frac{\sum_{i=1}^{M_1} \sum_{j=1}^{M_2} W(i, j) \cdot W'(i, j)}{\sqrt{\sum_{i=1}^{M_1} \sum_{j=1}^{M_2} W(i, j)^2} \sqrt{\sum_{i=1}^{M_1} \sum_{j=1}^{M_2} W'(i, j)^2}}, \quad (4)$$

where, $M_1 \times M_2$ is the size of the logo, $W(i, j)$, $W'(i, j)$ is the value of the elements of the original logo and the restored logo. The cross-correlation coefficient normalizes NC and has a value between 0 and 1. The closer the value of NC is to 1, the better the stability of the watermarking solution. However, NC greater than or equal to 0.75 is an acceptable threshold.

2. Experimental environment

We have built a program to simulate embedding and extracting watermarks using MATLAB R2013a (8.1.0.604) - 64bit. The program is installed on Laptop with system parameters such as: Windows 10 Pro operating system, Intel(R) Core(TM) processor i5-8625U CPU @ 1.60GHz 1.80GHz, 8GB RAM.

In this paper, we choose Lena image (Figure 10) as the image to be protected by copyright and a binary image as



Figure 11. Watermark logo



Figure 12. Experimental Lena image when embedding watermark by DWT-QIM method combined with VSS

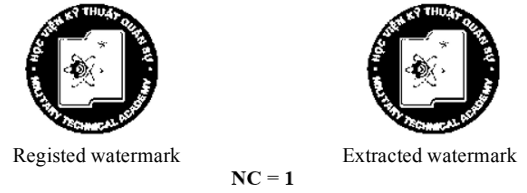


Figure 13. Extraction and reconstruction results of watermark by DWT-QIM method combined with VSS.

the copyright logo (Figure 11). We use Haar algorithm for DWT transformation. The parameters used for VSS in the test are $k = 4$ and $n = 10$.

3. Results of watermarking with combined VSS and without combined VSS

a. Experimental results with DWT-QIM method combined with VSS (Proposed Method)

Figure 12 shows the result after watermark embedding using the proposed method.

Figure 13 shows the results of watermark extraction using the proposed method.

b. Experimental results using only DWT-QIM method



Figure 14. Watermark embedding results in DWT-QIM method

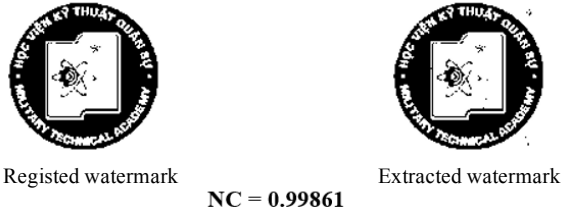


Figure 15. Watermark extraction results when using only DWT-QIM method

Figure 14 shows the results after embedding the watermark when using only the DWT-QIM method (not combining with VSS).

Figure 15 shows the results of watermark extraction when using only DWT-QIM method (not combining with VSS).

4. Determine the value of Q in QIM

In our proposal, the value of Q (Quantization step) plays a decisive role in the two criteria evaluating the efficiency of the algorithm, namely, the "unrecognizability" of the image after watermark embedding is represented by the ratio $PSNR$, and the "sustainability" of the watermark after it is extracted is expressed by the value NC . The smaller Q is, the larger $PSNR$ is and the smaller NC is. This means the smaller Q is, the greater the imperceptibility and the lower the stability of the watermark is.

As stated at the beginning of Section IV, each of the above values has its standard threshold. With a $PSNR$ of 35dB and a NC of 0.75. Therefore, to balance between the above two values when embedding and extracting watermarks, we have tested on different Q values to give the most reasonable value. The values of $PSNR$ and NC with different Q are shown in Table I.

TABLE I
 $PSNR$ AND NC VALUE WITH DIFFERENT Q

Quantization step (Q)	NC	$PSNR$ (dB)
5	0.75	48.09
10	1	42.72
15	1	39.02
20	1	36.57
25	1	34.59
30	1	32.97

Accordingly, we decided to choose the quantization step $Q = 20$ for our proposal because there the values of $PSNR$ and NC are closest to the standard level.

5. Image quality evaluation

To evaluate the image quality after watermarking of the proposed method, we use a dataset of 05 images: Lena,

Pepper, Couple, Mandrill and Parrots with size 512×512 and choose 01 binary logo size 64×64 to be our copyright logo. Simultaneously, the experiment was conducted with both methods: the proposed method (DWT-QIM-VSS) and the conventional watermarking method (DWT-QIM) on the Blue channel of the test image to compare the rating results. The results are shown in Figure 16 and Table II.

With the obtained results, we find that using the proposed watermarking method gives better $PSNR$ results than using the conventional method.

6. Robustness evaluation

As mentioned above, the value NC represents the stability of the watermark after it is extracted. In this section, we conduct tests and comparisons with watermarked images using both methods: using only DWT-QIM based on the idea of [25, 26] and DWT-QIM watermarking method combines VSS (recommended method) in the condition of being unaffected and affected by common attacks. Next, we extract the watermark and calculate the NC value to evaluate the sustainability of the watermark under the above conditions.

a. Watermark extraction results when not affected by attack spells

The results of watermark extraction on the dataset of 05 images: Lena, Pepper, Couple, Mandrill, Parrots, in the unattacked condition when using 2 different watermarking methods to embed the watermark are shown in Figure 17 and Table III. From the results of this experiment, we can see that the proposed method has better results even when the watermarked image has not been attacked by any image attacks.

b. Watermark extraction results under the influence of normal attacks

Figure 18 and Table IV show the visual results as well as concrete data when extracting watermark from container image using different methods for watermark embedding under conditions affected by normal attacks. The reason our proposed solution has a better result than the DWT-QIM algorithm is because of the following reasons: In the

TABLE II
 $PSNR$ IN THE TRADITIONAL WATERMARK METHOD AND THE PROPOSED METHOD

Image	$PSNR$ (DWT-QIM)	$PSNR$ (DWT-QIM & VSS)
Lena	36.3811	36.5217
Pepper	36.3224	36.4191
Couple	36.2893	36.6512
Mandrill	36.3538	36.5417
Parrots	36.5033	36.7141



Figure 16. Test results of 2 watermarking methods with a dataset of 05 images

 TABLE III
NC VALUES IN UNATTACKED CONDITION

Image	NC (DWT-QIM)	NC (DWT-QIM & VSS)
Lena	0.997711	1
Pepper	1	1
Couple	0.999852	0.999481
Mandrill	0.999718	1
Parrots	0.995578	0.999574

 TABLE IV
VALUE OF NC WHEN THE WATERMARKED IMAGE IS AFFECTED BY ATTACKS

Type of attacks	NC (DWT-QIM)	NC (DWT-QIM & VSS)
Gaussian noise (0.01)	0.27106	0.75892
Poisson noise	0.262291	0.756114
Salt & Pepper (0.02)	0.977128	0.983195
JPEG (QF = 50)	0.309221	0.787013
Histogram Equalization	0.394625	0.786501
Median filter (3×3)	0.833469	0.842374
Laplacian Sharpening	0.26666	0.777764
Blur (3×3)	0.272942	0.760467

DWT-QIM solution combining VSS, we disperse the original watermark into n . fragment and use only 1 fragment to embed information like DWT-QIM solution. However, because the remaining $(n - 1)$ fragment is used to register

with the Copyright Office, when extracting information a watermark fragment from the watermark image is attacked and restored with $(k - 1)$ fragment taken from $(n - 1)$ piece stored by Copyright Office will recover watermark information close to original watermark. Thereby proving the copyright even though the watermarked image has been attacked to change the image. This point is the difference of the proposed algorithm combining encryption and watermarking to secure watermark information and strengthen its stability against some common attacks. Thus, the experimental results show that the proposed solution is effective in copyright protection on digital images.

c. Proving DWTT-QIM-VSS is more sustainable than DWT-QIM

To prove the stability of the proposed solution, DWT-QIM-VSS, with the method using only DWT-QIM, we assume that the watermark images of both methods are attacked by any image attack. Then, the DWT-QIM method extracts the watermark image W' , the DWT-QIM method & VSS extracts the watermark image S'_1 , with $S_1, S_2, \dots, S_n$ is n shared fragment generated in VSS algorithm in II.1 and S'_1 is watermark extracted from hacked image of watermark S_1 . Without loss of generality, we assume that the change of W' relative to W and the change of S'_1 relative to S_1 are the same. Use VSS to recover watermark W'' according to the following formula:

Experimental extracted watermark comparison						
Original image	Lena	Pepper	Couple	Mandrill	Parrots	
 Registered watermark	 DWT-QIM method	 DWT-QIM method	 DWT-QIM method	 DWT-QIM method	 DWT-QIM method	NC =
 Our method	 Our method	 Our method	 Our method	 Our method	 Our method	NC =
	0.997711	1	0.999852	0.999718	0.995578	
	1	1	0.999481	1	0.999574	

Figure 17. Result of watermark extraction in unattacked condition

$$W'' = S'_1 \text{ OR } S_2 \text{ OR } \dots \text{ OR } S_n, \quad (5)$$

We need to prove $NC(W, W'') > NC(W, W')$.

Proof: Assume the same variation for W' and S'_1 , ie:

Assume the same variation for W' and S'_1 , ie:

$$TS(W, W') = TS(S_1, S'_1) = P > 0, \quad (6)$$

where, $TS(W, W')$ is the sum of the errors between W and W' , and $TS(S_1, S'_1)$ is the sum of the errors between S_1 and S'_1 .

Set

$$Z_0 = \{(i, j) | W(i, j) = 0\}, Z_1 = \{(i, j) | W(i, j) = "1"\} \quad (7)$$

Then, according to the properties of VSS, we have:

$$S_1(i, j) = 0, \forall (i, j) \in Z_0 \quad (8)$$

Assume the error is evenly divided by Z_0 and Z_1 . That means if the notation X_0 is the set of errors over Z_0 , X_1 is the set of errors over Z_1 of W' and Y_0 is the set of errors over Z_0 , Y_1 is the error set on Z_1 of S'_1 , then:

$$Sp(X_0) = Sp(X_1) = Sp(Y_0) = Sp(Y_1) = \frac{1}{2}P, \quad (9)$$

where $Sp(U)$ is the number of elements of the set U .

According to the properties of VSS (assuming $2 \leq k < n$):

$$\forall (i, j) \in Z_0 \implies S_t(i, j) = 0, t = 2, \dots, n$$

$$\forall (i, j) \in Z_1 \implies S_2(i, j) \text{ OR } \dots \text{ OR } S_n(i, j) = 1, t = 2, \dots, n$$

From that and (5), (7), (8) deduce:

$$W''(i, j) = 1, \forall (i, j) \in Z_1 \text{ (false fixed on } Z_1)$$

$$W''(i, j) = 0, \forall (i, j) \in Z_0, Y_0$$

$$W''(i, j) = 1, \forall (i, j) \in Y_1 \text{ (cannot be corrected on } Z_0)$$

Contrast this with (7), we get:

$$TS(W, W'') = Sp(Y_0) = \frac{1}{2}P \quad (10)$$

Combine this formula with (6), and get:

$$TS(W, W'') = \frac{1}{2}TS(W, W') \quad (11)$$

Thus, the error of W'' is halved compared to W' .

Now if we put: $T = Sp(Z_1)$, $q = Sp(X_0) = Sp(X_1) = Sp(Y_0)$, it's easy to see:

Above Z_1 : W and W'' both have T pixels equal to "1"; W' only $(T - q)$ pixels equals "1".

Above Z_0 : W has $Sp(Z_0)$ pixels equal to "0"; W' and W'' have $Sp(Z_0) - q$ pixels equal to "0".

From there it is easy to calculate:

$$NC(W, W'') = \frac{T}{\sqrt{T}\sqrt{T+q}} \quad (12)$$

$$NC(W, W') = \frac{T-q}{\sqrt{T}\sqrt{T}}$$

Thus, to prove that $NC(W, W'') > NC(W, W')$, it is necessary to show:

$$\frac{T}{\sqrt{T+q}} > \frac{T-q}{\sqrt{T}} \quad (13)$$

This is equivalent to:

$$T\sqrt{T} > (T-q)\sqrt{T+q}$$

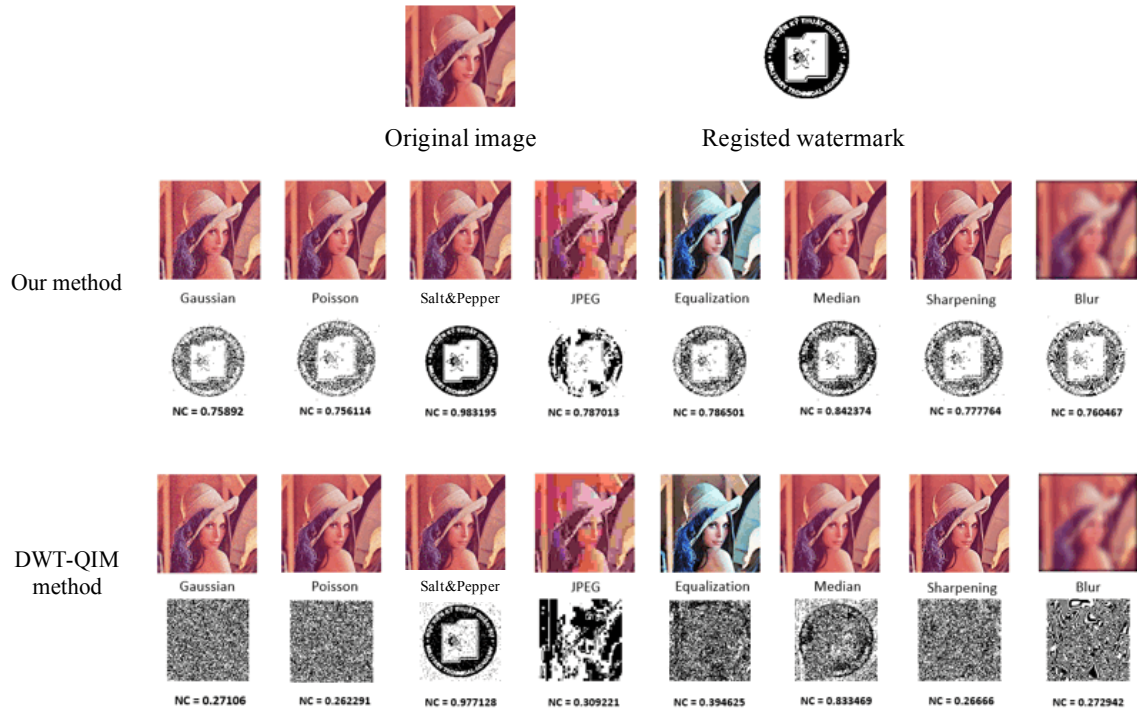


Figure 18. Watermark extraction results under many attacks

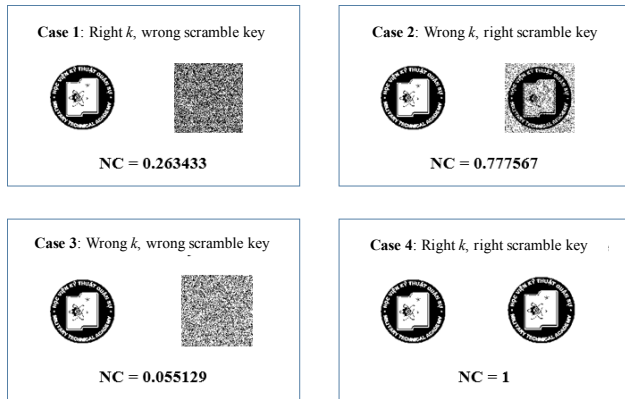


Figure 19. Copyright logo is reproduced in different cases

Square both sides:

$$T^2T > (T - q)^2(T + q) = (T^2 - q^2)(T - q)$$

This is obviously true, so the proof is complete.

Thus, with the condition that it is not affected by normal attacks, the value of NC between the original watermark and the extracted watermark when using both methods is relatively similar. It is difficult to conclude that our proposed watermarking method is more stable than the method using only DWT-QIM in this condition. Therefore, we conduct experiments with conventional image attacks on

images containing Lena watermark to once again evaluate the effectiveness of the algorithm. The results show that the proposed method of watermarking using DWT-QIM in combination with VSS can ensure the stability of the watermark against common attacks.

7. Security evaluation

In this paper, to enhance the security of the copyright logo, in addition to the visual cryptographic technique (VSS), we also use the scramble technique on the copyright logo before distributing it into n shared secret fragments. We use Arnold Transform [29] and shuffle key = 10 to shuffle the logo image before scattering using VSS. Thus, to restore the copyright logo, it is necessary to have 02 elements, namely: “Minimal share k value to restore copyright logo” and “key scramble” to recover scramble.

TABLE V
NC VALUES FOR EACH KEY

Test case	NC
Right k , wrong scramble key	0.263433
Wrong k , right scramble key	0.777567
Wrong both k và scramble key	0.055129
Right both k và scramble key	1

In this section, we evaluate the security of the algorithm against copyright information in the cases shown in Figure 19 and Table V.

Thus, through tests in each specific case, we have shown that the algorithm we propose can completely secure the information of the copyright logo.

V. CONCLUSION

Our paper proposes a digital image copyright protection solution based on the method of combining digital watermarking with visual cryptography to create a more sustainable and effective method in the digital image protection and digital product copyright.

From the results obtained through various tests with both the algorithm when using DWT-QIM in combination with VSS (recommended method) and using only DWT-QIM, we found that when using only DWT-QIM, the image quality after watermark embedding is worse than that of the proposed method. However, with tests under the influence of attacks, using the proposed method will ensure more robustness of the watermark than the traditional method of using only DWT-QIM.

Thus, in this paper, we have shown that: with the application of visual cryptographic techniques to disperse watermarks, and at the same time combine the watermarking algorithm based on the discrete wavelet transform DWT using DWT quantization index (QIM), has contributed to improving the imperceptibility as well as the stability of the watermark against conventional attacks.

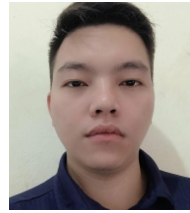
ACKNOWLEDGMENT

This research is funded by the Vietnam National Foundation for Science and Technology Development (NAFOS-TED) under grant number 102.01-2019.12

REFERENCES

- [1] A. K. Singh, "Robust and distortion control dual watermarking in LWT domain using DCT and error correction code for color medical image," *Multimedia Tools and Applications*, vol. 78, no. 19, pp. 30523–30533, 2019.
- [2] M. Charfeddine, M. El'arbi, and C. B. Amar, "A new DCT audio watermarking scheme based on preliminary MP3 study," *Multimedia Tools Appl.*, vol. 70, no. 3, pp. 1521–1557, 2014. DOI:https://doi.org/10.1007/s11042-012-1167-0
- [3] Roy, S., & Pal, A. K., "An indirect watermark hiding in discrete cosine transform-singular value decomposition domain for copyright protection," *Royal Society open science*, 4(6), 2017. https://doi.org/10.1098/rsos.170326
- [4] Barni M., Bartolini F., Piva, A., "Improved wavelet-based watermarking through pixel-wise masking," *IEEE transactions on image processing*, 10(5), pp.783–791, 2001.
- [5] Singh A. K., Kumar B., Dave M., Mohan A., "Robust and imperceptible dual watermarking for telemedicine applications," *Wireless Personal Communications*, 80(4), pp. 1415–1433, 2015.
- [6] Singh A. K., Dave M., Mohan A., "Robust and secure multiple watermarking in wavelet domain," *Journal of medical imaging and health informatics*, 5(2), pp. 406–414, 2015.
- [7] Singh AK, Dave M, Mohan A, "Hybrid technique for robust and imperceptible multiple watermarking using medical images," *Multimedia Tools Appl.*121, 2015.
- [8] Urvoy M., Goudia D., Autrusseau F., "Perceptual DFT watermarking with improved detection and robustness to geometrical distortions," *IEEE Transactions on Information Forensics and Security*, 9(7), pp. 1108–1119, 2014.
- [9] Cedillo-Hernandez M., Garcia-Ugalde F., Nakano-Miyatake M., Perez-Meana H., "Robust digital image watermarking using interest points and DFT domain," In: *35th IEEE International Conference on Telecommunications and Signal Processing (TSP)*, pp. 715–719, 2014.
- [10] T. M. Thanh, K. Tanaka, "The novel and robust watermarking method based on q -logarithm frequency domain," *Multimedia Tools and Applications*, volume 75, pp. 11097–11125, 2016.
- [11] T. M. Thanh, K. Tanaka, "A proposal of novel q -DWT for blind and robust image watermarking," *Proceeding of IEEE 25th International Symposium on Personal, Indoor and Mobile Radio Communications - (PIMRC)*, Washington D.C., pp. 2066–2070, 2014.
- [12] T. M. Thanh, P. T. Hiep, T. M. Tam, "A New Spatial q -log Domain for Image Watermarking," *IJIIP: International Journal of Intelligent Information Processing*, Vol. 5, No. 1, pp.12–20, ISSN 2093-1964, 2014.
- [13] T. M. Thanh, K. Tanaka, "An image zero-watermarking algorithm based on the encryption of visual map feature with watermark information," *International Journal of Multimedia Tools and Applications (MTAP)*, ISSN: pp. 1573–7721, 2016.
- [14] Mohanarathinam, A., Kamalraj, S., Prasanna Venkatesan, G.K.D. et al. "Digital watermarking techniques for image security: a review," *J Ambient Intell Human Comput* 11, pp.3221–3229 (2020).
- [15] M. Naor and A. Shamir, "Visual Cryptography," *Advances in Cryptology (EUROCRYPTO 94)*, (Lecture Notes in Computer Science), vol. 950, A. De Santis, Ed. Berlin, Germany: Springer-Verlag, pp. 1–12, 1995.
- [16] A. Shamir: "How to share a secret ?", *Comm ACM*, 22(11): pp. 612–613, 1979.
- [17] Mukesh Dalal, Mamta Juneja, "Steganography and Steganalysis (in digital forensics): a Cybersecurity guide," *International Journal of Multimedia Tools and Applications (MTAP)*, ISSN: 80, pp. 5723–5771, 2021.
- [18] G. Blakley, "Safeguarding cryptographic keys," *Proc. of AFIPS National Computer Conference*, 1979.
- [19] C. Asmuth and J. Bloom, "A modular approach to key safeguarding," *IEEE transaction on Information Theory*, 29(2), pp. 208–210, 1983.
- [20] Shyamalendu Kandar and Bibhas Chandra Dhara, "k-n Secret Sharing Visual Cryptography Scheme on Color Image using Random Sequence," *International Journal of Computer Applications* (0975 – 8887), Volume 25, No.11, July 2011.
- [21] E. Elbasi, A. M. Eskicioglu, "Naïve Bayes Classifier Based Watermark Detection in Wavelet Transform," *Multimedia Content Representation, Classification and Security (MRCS)*, Lecture Notes in Computer Science, vol 4105, 2006.
- [22] S. Huang, W. Zang, "Blind Watermarking Scheme Based on Neural Network," *The 7th World Congress on Intelligent Control and Automation*, Chongqing, pp. 5985–5989, 2008.
- [23] Giang Ngọc Dân, Tạ Minh Thanh, "Giải pháp bảo vệ bản quyền ảnh số dựa trên kỹ thuật kết hợp thủy vân số và mã hóa trực quan," *Journal of Science and Technology on Information and Communications*, 2020.
- [24] R. Sun, H. Sun, T. Yao, "A SVD and Quantization Based Semi_fragile Watermarking Technique for Image Authentification"

- cation,” *ICSP’02 Proceedings*, pp. 1592–1595, 2002.
- [25] A. Akter, Nur-E-Tajina, M. A. Ullah, “Digital image watermarking based on DWT-DCT: Evaluate for a new embedding algorithm,” *2014 International Conference on Informatics, Electronics & Vision (ICIEV)*, Dhaka, 2014, pp. 1–6, doi: 10.1109/ICIEV.2014.6850699.
- [26] A. Benoraira, K. Benmahammed, N. Boucenna, “Blind image watermarking technique based on differential embedding in DWT and DCT domains,” *EURASIP Journal on Advances in Signal Processing*, vol 55, 2015.
- [27] O.Jane, E. Elbasi, H.G.Ilk, “Hybrid Non-Blind Watermarking Based on DWT and SVD,” *Journal of Applied Research and Technology*, vol. 12, Issue 4, pp. 750–761, 2014.
- [28] Cox, Ingemar J., Joe Kilian, F. Thomson Leighton, Talal Shamoon. “Secure spread spectrum watermarking for multimedia,” *IEEE transactions on image processing*, vol. 6, no. 12, pp. 1673–1687, 1997.
- [29] M. Li, T. Liang, Y. He, “Arnold Transform Based Image Scrambling Method,” *3rd International Conference on Multimedia Technology*, pp. 1309–1316, 2013.



Dan Giang Ngoc is currently studying Network Technology at Le Quy Don University. His field of research is digital watermarking for digital multimedia.
Email: giangngocdan3110@gmail.com.



Thanh Ta Minh is currently an associate professor and vice dean of Faculty of Information Technology in Le Quy Don Technical University Vietnam. He is also a Postdoctoral Fellow of the Department of Mathematical and Computing Sciences at Tokyo Institute of Technology. He received his B.S. and M.S in Computer Science from

National Defense Academy, Japan, in 2005 and 2008 and his Ph.D. from Tokyo Institute of Technology, Japan, in 2015, respectively. He is the member of IPSJ Japan and IEEE. His research interests lie in the area of watermarking, network security, and computer vision.

Email: thanhhtm@lqdtu.edu.vn.

An Integrated Approach for Table Detection and Structure Recognition

Hai-Hong Phan, Ngo Dai Duong

Le Quy Don Technical University, Ha Noi, Viet Nam

Correspondence: hongph@lqdtu.edu.vn; daiduong28789@hotmail.com

Communication: received 4 April, revised 29 May, accepted 31 May

Digital Object Identifier: 10.32913/mic-ict-research.v2021.n1.974

Abstract: Detecting and identifying the table structure is an important issue in document digitization. Although there have been many great strides based on current deep learning techniques, table structure identification is still a difficult and arduous problem, especially when solving the problem of digitizing text in practice. The paper proposes a solution to digitize table documents based on the Cascade R-CNN HRNet network to detect, classify tables and integrate image processing algorithms to improve table data identification results. The proposed algorithm proved effective on real data - the hydrometeorological station record book contains tables including simple and complex structures tables with over 98% accuracy.

Keywords: Document structure analysis, table detection, detect the table structure, hydrometeorology.

I. INTRODUCTION

Currently, digital transformation is one of the development goals of each country. The digital transformation brings great benefits to management, economic and social development. In particular, the use of digital documents instead of traditional documents is being rapidly developed. Digitizing tables in documents makes a lot of sense because tables often contain a lot of important and useful information. There are two main tasks of table digitization: table detection and table understanding. Table detection is to identify the table container, assign an identifier, and distinguish it from the text areas in the document. Table structure identification is to identify cells, rows, columns, and hierarchical relationships between cells.

This paper proposes a method to identify the table structure for the real-world problem - digitizing several types of hydrogeological records. The model is classified into tabular forms: old technical books before 2000, new technical books since 2000 and normal tables. In particular, the table usually has a simple structure, the technical book table has a complex structure.

The article's contributions include:

- Proposing to use Cascade mask R-CNN HRNetv2pW40 deep learning network for table detection and table classification, an improvement over using W=32. Achieve up to 98.74% accuracy for the table detection task with the Hydrometeorology set.

- Proposing a total solution to recognize the table structure, using techniques to extract the structure of simple tables. Use templates to handle complex structured tables encountered in practice.

The solution has achieved high accuracy for the task of detecting and classifying tables based on deep learning techniques and image processing on the Hydrometeorology dataset.

II. RELATED WORK

In 1997, P.Pyreddy and, W.B.Croft [1] first proposed a table detection method based on character alignment, holes and spaces.

Naganjaneyulu et al. [2] in 2016 proposed a model for table detection using two Hough algorithms for line detection and Harris for corner detection. The author applied it on many different types of documents and achieved an accuracy of up to 89%. This method fails with tables with uneven spatial delineation and table formatting.

Dhiran et al. [3] proposed a method to detect and extract table information from document images, using OCR for preprocessing, enhancement and contour noise removal. The author uses 1200 document images which are divided into 3 categories: type with line, type with line with column whitespace and type with only white space. This method achieves an accuracy of up to 88.7%. However, this method does not work well with documents that are divided into multiple columns.

Gilani et al. [4], propose a method with two modules of image transformation and table detection. In the article, the author proposes to use the Abbey Cloud OCR SDK and tesseract OCR to display the comparison results.

Shafait et al. [5], propose an approach to detect tables from heterogeneous documents with different table layouts. Suggest your approach to document analysis via tab-stop. Huynh-Van et al. [6], proposed a three-step combined method to detect tablespaces: 1) Zone classification, 2) Table detection based on intersection of vertical and horizontal lines and 3) define a table made up of parallel lines only. The authors use data set UW-III for the experiment to obtain the results. Achieve precision (82%) and recall (80%).

Kleber et al. [7], proposed a method to match the sample table structure using a linked graph for handwritten/printed historical documents. Table matching is performed by detecting the maximum group in the association graph, which represents the matching of the sample and document of interest line information.

Along with the development of deep learning techniques. In 2017, Schreiber et al. [8] proposed a deep learning-based approach for table detection and table structure recognition in documents or images. The author uses Faster R-CNN network for table discovery task and FCN network for table structure recognition task (specifically, table row and column recognition). This method is tested on the ICDAR 2013 dataset consisting of 67 documents with 238 pages, with the F1 measure reaching 96.77% for the table detection task and 91.44% for the table structure recognition task.

Devashish Prasad et al. [9] in 2020 proposed an end-to-end deep learning model CascadeTabNet used to detect and recognize the table structure. The article's tests are on the data sets ICDAR 2013, ICDAR 2019, Tablebank [10]. With the F1 measure achieved 90.1% for the ICDAR 2019 dataset, 94.33% for the Tablebank dataset for the table discovery task, 23.2% for the ICDAR 2019 dataset branch B2 (modern tables) for the identification table structure task.

In object detection, a certain IoU (Intersection over Union) threshold is used to determine positive and negative. Usually will be trained with a low IoU threshold eg 0.5, this causes the model to always generate a lot of noise in the detection. However, if the IoU threshold is increased, the detector performance decreases. The underlying cause for this phenomenon is over-fitting in the training process, due to an exponential decline with positive samples as the IoU threshold is increased. The multi-stage object detection architecture, Cascade R-CNN proposed by Zhaowei Cai et al. in 2017 [11] aims to solve these problems. Cascade R-CNN uses several stages (usually 3), each stage resampling, thereby changing the distribution of the input hypothesis. Cascade R-CNN has relatively high accuracy with test benchmark data.

Most of the networks developed used for classification

such as Alexnet, VGGNet, GoogleNet, ResNet are usually followed by Lenet-5's rule of reducing the spatial size of feature maps, connecting convolutions from degrees high resolution to low resolution in a sequence and using these low resolution representations for classification. Meanwhile, high-resolution representations are needed for tasks such as semantic segmentation, human action computation, and object detection. Jingdong Wang et al. proposed that HRNet (High-Resolution Net) [12] is capable of maintaining high resolution representations throughout the processing. The HRNetV2p instance is made up of a feature tower from HRNetV2 representatives used for the object detection task and evaluated against the COCO dataset.

These methods, which mostly deal with photos of construction documents for competitions, have not yet processed and proven the method's effectiveness with real-world data. Second, the data is usually of a historical or modern type only. While the data to be digitized in practice can include both. The article focuses on proposing algorithms to recognize table structure on actual data from hydrometeorological stations. This data has a high difficulty because it contains many different table structures, besides there are both old and new documents, so the image quality is much different.

III. PROPOSED METHODS

In this paper, we propose an algorithm intergrating deep learning and traditional to solve the real problem - digitizing tables of hydrometeorological documents. Towards this goal, the main strategy of the paper includes:

- Based on state-of-the-art CNN architecture for table segmentation and detection.
- Tranfer learning for Hydrometeorological data to improve performance.
- Use the template table to better identify the table structure in the document.

With the collected initial data, we divide the Hydrometeorology data into the following 3 types of tables: table_skt1_old, table_skt1_new, table_normal. The table_skt1_old is the old table of the meteorological observation book from 1969 to 1985. The table_skt1_new is the table of the basic meteorological observation book from 2000s. Table_normal is the remaining table, both old and new, including specific tables for each type of documents.

We propose a solution to recognize the table structure as shown in Figure 1. Accordingly, the original document image is included in the Cascade mask RCNN HRNetv2pW40 model to detect the tables in the document. In addition, the model also classifies the table into 3 types, corresponding to 3 output branches: Table_normal, Table_skt1_old and

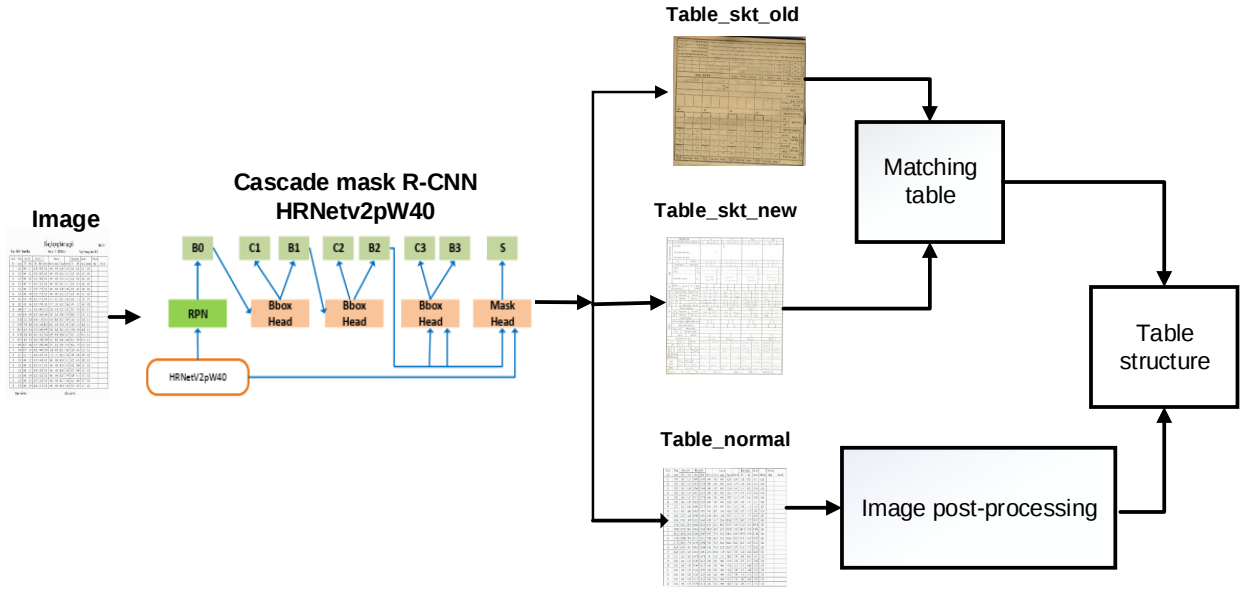


Figure 1: The proposed pipeline for identifying the table structure in the Hydrometeorology documents.

Table_skt1_new. Tables in the table_normal branch would be processed by traditional image processing methods to recognize the cell areas in the table and the relationship between them. Besides, the proposed method applies alignment techniques and template tables to identify the structure of the table for table_skt1_new and table_skt1_old.

1. The proposed table detection model

In this paper, the proposed method integrates Cascade mask R-CNN and HRNetV2pW40. Cascade mask R-CNN is very similar to Cascade R-CNN, this model extends the ability to segment objects by adding a segmentation branch as done with Mask RCNN.

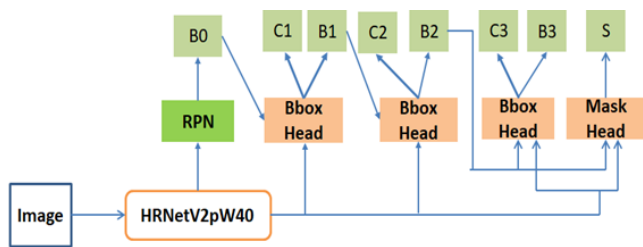


Figure 2: Illustration of using Cascade Mask model R-CNN HRNetV2pW40 for table detection in Hydrometeorological documents

As shown in Figure 2, HR_NetV2p_W40 model extracts feature map of input image I. The Region Proposal Network (RPN) predicts the original proposed features from the feature maps. “Bbox Head” takes RoI features as input and makes predictions for each RoI. Each “Bbox Head” makes

two predictions: predicting the bounding box and classifying. We indicate “B” as the bounding boxes predicted by the “head”. “C” is the classification prediction of the object. “Mask Head” predicts the mask of the object and “S” denotes the output segment. At the inference step, the object detected by “Bbox Head” is supplemented with the segment mask performed by “Mask Head”.

The model bases on the MMDetection library [13]. We use the default cascade mask configuration of rcnn hrnetv2p w32 20e for our model during the experiment and analysis.

2. Transfer learning

For table detection task, we use iterative transformation learning technique. First, we use the general data containing many different tables, like table with border, table without border, table lacking border. After training, we re-adjust with the hydrographic materials of the table, in order to create a model that can infer well with this special the data set.

Figure 4 illustrates the transformation process for problem release table accountants.

First, we use a model that has been trained with generic data. We only perform the task of detecting tables, all tables in the document will be annotated as a class (table class).

Second, we refine the model trained in the first step. With data obtained from documents of hydrometeorological centers of Vietnam. In this step, we classify the table into 3 types: table_skt_1_old, table_skt_1_new, table_normal. Ta-

(a) Table skt_1_old

(b) Table skt_1_new

Figure 3: Table data labeled table_skt_1_new and table_skt_1_old

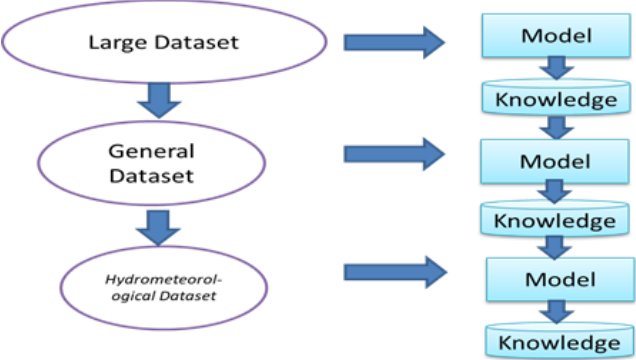


Figure 4: Two stage transfer learning for the Hydrometeorological table detection

ble_skt_1_old are old technical books, historical documents (since 1969) and complex structure.

Table_skt_1_new are new, complex tables. Table_normal are simple structured tables that can be processed by traditional or deep learning methods.

3. Analyze table structure based on sample table

For some types of tables with complex structure, historical documents, it is very difficult to automatically recognize the table structure. For special data with known information about some tables, several techniques can be used to identify the structure of the table as shown in Figure 5. First, after the table is detected in the document with the model in the previous figure, the tables table_skt_1_old and table_skt_1_new are truncated from the table container. Second, we use ORB (Binary Robust Independent Elementary Features) technique to extract features of the image, then use RANSAC algorithm to match key points in two images, find the identity matrix and transform the input image according to the identity matrix. We get the output image with a high match to the sample image. Third, we define the table template as xml. Enclose box coordinates around table cells and show their relationship. The cells are then OCR to detect the filled text in the defined table structure.

IV. RESULTS AND ANALYSIS

TABLE I: COMPARISON OF TABLE FINDING RESULTS ON GENERAL DATA.
(data compiled from three databases ICDAR19 [15], Marmot [15] and Github [16])

Model	IoU				Wags
	0.6	0.7	0.8	0.9	
Cascade mask RCNN HRNetv2pW32 [9]	0.9604	0.9606	0.9658	0.9696	0.9641
Cascade mask RCNN HRNetv2pW40	0.9667	0.9669	0.9674	0.9693	0.9675

TABLE II: COMPARISON TABLE FINDING RESULTS ON DATA SET TABLEBANK.

Model	Word			Latex			Word+Latex		
	P	R	F1	P	R	F1	P	R	F1
X101(Word) [17]	0.9352	0.9398	0.9375	0.9905	0.5851	0.7356	0.9579	0.7474	0.8397
X152(Word) [17]	0.9418	0.9415	0.9416	0.9912	0.6882	0.8124	0.9641	0.8041	0.8769
X101(Latex) [17]	0.8453	0.9335	0.8872	0.9819	0.9799	0.9809	0.9159	0.9587	0.9368
X152(Latex) [17]	0.8476	0.9264	0.8853	0.9816	0.9814	0.9815	0.9173	0.9562	0.9364
X101(Word+Latex) [17]	0.9178	0.9363	0.9270	0.9827	0.9784	0.9806	0.9526	0.9592	0.9559
X152(Word+Latex) [17]	0.9229	0.9266	0.9247	0.9837	0.9752	0.9795	0.9557	0.9530	0.9543
The proposed method (Word -3000 sample)	0.9563	0.9514	0.9538	0.9276	0.9826	0.9543	0.9365	0.9687	0.9523
The proposed method (Latex - 2000 sample)	0.9112	0.8698	0.8900	0.9845	0.9842	0.9844	0.9646	0.9466	0.9556
The proposed method (All 4000 sample)	0.9286	0.9211	0.9248	0.9749	0.9797	0.9773	0.9620	0.9588	0.9604

TABLE III: TABLE FINDING RESULTS AFTER FINETURNING MODEL AND DATA ON KTTV DATA.

Model	IoU				Wags
	0.6	0.7	0.8	0.9	
Cascade mask RCNN HRNetv2pW40 + small data is dilation	0.9784	0.9855	0.9855	0.9855	0.9837
Cascade mask RCNN HRNetv2pW40 fineturning general dataset	0.9723	0.9774	0.9794	0.9803	0.9773
Cascade mask RCNN HRNetv2pW40 fineturning general dataset + dilation	0.9865	0.9865	0.9883	0.9883	0.9874

TABLE IV: RESULTS OF CLASSIFICATION OF TYPES OF HYDROLOGICAL TABLES.

Table type name	Correct identification number	Amount to be identified	Probability
Table_normal	115	115	1.0
Table_skt1_old	233	233	1.0
Table_skt1_new	190	190	1.0

1934 photos, the documents in the photos are new, have good image quality, and are available in 2 languages: English and Chinese.

- Tablebank dataset [10]: Tablebank ...to evaluate the quality of the table detector. Tablebank is a relatively large data set of 278582 images divided into 2 main types, Latex and Word.

- Hydro-meteorological data: The first collected hydro-meteorological data set includes 780 high-resolution scanned images (360dbi), equivalent to the image width of 6 inches. Data is taken from 13 different types of technical books related to meteorology, for example: basic meteorological observation book, soil temperature monitoring book. The data is divided and labeled into 3 types of tables table_skt1_old, table_skt1_new, table_normal.

In Figure 3 (a) the skt1_old tables are the old tables of the meteorological observation books dating from 1969 to 1985. The pictures are often of poor quality, due to old books, faded ink, and bad paper. Figure 3 (b) skt1_new are the tables of the basic meteorological observation books from 2000 onwards. These document images are relatively good, but the table structure is very complicated, the vertical lines are short, not aligned from the end to the end, so it is difficult to identify the links from the table cells automatically. These two types of tables have complex structures.

Figure 6 is a description of the table_normal table type. Figure 6-a The table_normal table has a simple structure, these tables can be identified by the traditional methods of row and column identification, or using learning deep[14],

(a) Detect table table_normal

(b) Detect table table_stk1_new

(c) Detect table table_stk1_old

Figure 7: Some results of the hydrometeorological data set table detection.

Figure 6-b is a type of table with a complex structure that is highly similar to table_stk1_new. However, because the amount of data we currently have is too small to classify this table as a separate type, we have included it with the table_normal table type. In addition, this is also used to test the robustness of the model's classifier. In total, our data includes 780 images, in which table_normal has 243 images, table_stk1_old has 300 images, table_stk1_new has 237 images. We will rely on this data set to evaluate the quality of the model.

2. Results and analysis

We conduct model testing on the data sets presented in section IV.2.a. Compare some CNN network models, and the capabilities of our proposed model. We also tested some post-processing techniques related to output table alignment. Our experiments are performed on Google Colaboratory platform with P100 PCIE GPU with 16GB GPU memory, Intel(R) Xeon(R) CPU @2.30GHz and 12.72 GB Ram.

For each dataset, we use 60% for training set, 40% for testing set. Specifically with the general dataset, we use 1160 images for the training set, 774 images for the test set; With the KTTV dataset, we use 486 images for the training set, 294 images for the test set. First, we use common data to evaluate the panel detection of two models Cascade mask R-CNN HRNetv2pW32 (CascadeTabNet [9]) and Cascade mask R-CNN HRNetv2pW40 (W represents feature width).

a) Experimental results on the general data set

From the results of Table I, it is found that the Cascade mask R-CNN model HRNetv2pW40 has higher accuracy. This can be explained by the fact that we use more feature widths. Therefore, the Cascade mask R-CNN HRNetv2pW40 model will be selected to be used to conduct further evaluation tests. Table I compares the results between the models, based on the F1 score. The proposed model achieves higher results than the Cascade mask RCNN HRNetv2pW32 model in most IoU indexes with an average accuracy of 96.75%.

b) Experimental results on the TableBank data set

Second, we evaluate the table detection quality of the model on the Tablebank dataset. The results show that the model achieves the best results of 96.44% with the F1 score, with both Latex and Word data. While, we only use 4000 samples for training instead of using 260582 training samples of original setting. The test set uses the split ratio of the Tablebank dataset (5719 for latex, 2281 words, 8000 samples for both)

c) Experimental results on hydrometeorological data set

Third, we test the Cascade mask R-CNN HRNetv2pW40 model with hydrometeorological data, combined with the use of dilation enhancement technique. Here, we did not use the smudge technique to enhance the image as old hydrometeorological images are of poorer quality than the original modern images used in [9], the use of smudge

Giờ Hà Nội									
VV	Tiêu điểm	Khoảng cách	Mã số	1	7	13	19		
	Ký hiệu			119	200m	92	04	200m	02
	Thời điểm bắt đầu								
	Thời điểm kết thúc								
	ww	W ₁ W ₂		A5	A1	A4	A3	A1	02
	Tổng quan	Máy dưới		?	?	?	?	?	?
	Trên	C ₁₁		1	1	1	1	1	1
	Giữa	C ₁₂		1	1	1	1	1	1
	Dưới. Độ cao (m)	C ₁		1	1	1	1	1	1
	VILD	dd	ff						
	Từ báo (TG)								
	Nhiệt kế khô	Nhiệt kế ướt		188	184	185	182	174	216
	Mật rượu	Con trỏ		188	187	186	183	274	186
	Mật Hg	Sau vảy		220	188	188	188	276	307
	Td	e	U	181	208	96	180	206	97
	d	Nhiệt kế	Ấm kỹ	09	185	95	07	180	95
		Nhiệt kế thường		205	195	195	195	195	195
	Tg	Mật rượu	Con trỏ	205	204	195	180		225
	Tg	Mật Hg	Sau vảy	225	205	205	195	180	195
		Nhiệt kế phụ							
	Khí áp kế	Số đọc	H/c						
		Biến thiên khí áp							
		Hiệu chỉnh							
		Đã hiệu chỉnh							
		Áp ký	Δ P24h						
		Tên quan trắc viên		Chường	Chường	Chường	Chường	Chường	Chường
	E	7h 0E	19h 0E	ARXX	17181	AR101	31192	20000	10188
	8h	T 185	P ₀						
	Số đọc	19-7	7-19						
		22/20	18						
	Lượng	02	18						
	RRRR (mm)								
	TB ngày	T 218	T 264						
	ddfx	e 218	U 85						
		W 275	d 50						
	Tx: 307	Tp: 185	Tg: 175						

(a)

Ngày ... Tháng ... Năm ...									
VV	Tiêu điểm	Khoảng cách	Mã số	1	7	13	19		
	Ký hiệu			w	200m	92	cy	200m	02
	Thời điểm bắt đầu								
	Thời điểm kết thúc								
	ww	W ₁ W ₂		445	44	445	444	02	447
	Tổng quan	Máy dưới		1982	2	2	2	9	9
	Trên	CH		1	1	1	1	1	1
	Giữa	CM		1	1	1	1	1	1
	Dưới. Độ cao (m)	CL		03	F	1990			
	VILD	dd	ff						
	Từ báo (TG)								
	Nhiệt kế khô	Nhiệt kế ướt		146	140	132	26	208	162
	Tn	Mật rượu	Con trỏ	146	145	132	130	209	132
	Tx	Mật Hg	Sau vảy	174	146	1446	132	212	132
	Td	e	U	135	155	q3	121	1444	q3
	d	Nhiệt kế	Ấm kỹ	14	146	95	14	128	q6
		Nhiệt kế thường		165	19800198		34s		
	Tg	Mật rượu	Con trỏ	165	1644	US5	140	Trời nó	
	Tgx	Mật Hg	Sau vảy	180	165	465	155	3ss	345
		Nhiệt kế phụ							
	Khí áp kế	Số đọc	H/c						
		Biến thiên khí áp							
		Hiệu chỉnh							
		Đã hiệu chỉnh							
		Áp ký	Δ P24h						
		Tên quan trắc viên		03600000009	03600000009		Chường minh		
	E	7h 0E	19h 0E	ARV1	0281	48/01	31/92	96000	10146
	8h	T 136	P ₀						
	Số đọc	19-7	7-19						
		16/60	09	X	0310000001		323		
	Lượng	02	09	M	Nvx		03061	48/01	
	RRRR (mm)						Link		
	TB ngày	T 66	Tg 210	P ₀	ADXX		0312	48.51	
	ddfx	ess	U.844	d.36				000 18	
		0360100							
		000000							

(b)

Figure 9: Illustration of using sample table to identify cells of interest.

have another study to handle this handwriting.

V. CONCLUSION

This paper proposes an overall method based on the Cascade mask R-CNN HRNetw40 model to detect table objects, combining ORB and RANSAC techniques to digitize table documents. The results show that the proposed algorithm is capable of detecting and classifying document tables very well. The table detection results reached 98.74% on the F1 score and recognized different types of tables with a very high accuracy rate on the three test data sets. The results obtained in the study show the feasibility of applying deep learning techniques in the process of digitizing documents, especially tabular documents.

REFERENCES

- [1] P. Pyreddy and W. B. Croft, "Tintin: A system for retrieval in text tables," in *Proceedings of the Second ACM International Conference on Digital Libraries*, ser. DL '97. New York, NY, USA: Association for Computing Machinery, 1997, p. 193–200. [Online]. Available: <https://doi.org/10.1145/263690.263816>
- [2] G. V. S. S. K. R. Naganjaneyulu, N. V. Sathwik, and A. Narasimhadhan, "A multi clue heuristic based algorithm for table detection," in *2016 IEEE Region 10 Conference (TENCON)*, 2016, pp. 1246–1249.
- [3] T. Dhiran and R. Sharma, "Table detection and extraction from image document," *International Journal of Computer & Organization Trends*, vol. 3, no. 7, pp. 275–278, 2013.
- [4] A. Gilani, S. R. Qasim, I. Malik, and F. Shafait, "Table detection using deep learning," *IEEE*, 09 2017.
- [5] F. Shafait and R. Smith, "Table detection in heterogeneous documents," in *Proceedings of the 9th IAPR International Workshop on Document Analysis Systems*, ser. DAS '10. New York, NY, USA: Association for Computing Machinery, 2010, p. 65–72. [Online]. Available: <https://doi.org/10.1145/1815330.1815339>
- [6] T. Huynh-Van, K. Nguyen-An, T. L. B. Khanh, H.-J. Yang, T. A. Tran, and S.-H. Kim, "Learning to detect tables in document images using line and text information," in *Proceedings of the 2nd International Conference on Machine Learning and Soft Computing*, ser. ICMLSC '18. New York, NY, USA: Association for Computing Machinery, 2018, p. 151–155. [Online]. Available: <https://doi.org/10.1145/3184066.3184091>
- [7] F. Kleber, M. Diem, H. Dejean, J.-L. Meunier, and E. Lang, "Matching table structures of historical register books using association graphs," in *2018 16th International Conference on Frontiers in Handwriting Recognition (ICFHR)*, 2018, pp. 217–222.
- [8] S. Schreiber, S. Agne, I. Wolf, A. Dengel, and S. Ahmed, "Deepdesrt: Deep learning for detection and structure recognition of tables in document images," in *2017 14th IAPR International Conference on Document Analysis and Recognition (ICDAR)*, vol. 01, 2017, pp. 1162–1167.

- [9] D. Prasad, A. Gadpal, K. Kapadni, M. Visave, and K. Sultanpure, "Cascadetabnet: An approach for end to end table detection and structure recognition from image-based documents," in *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2020, pp. 2439–2447.
- [10] M. Li, L. Cui, S. Huang, F. Wei, M. Zhou, and Z. Li, "Tablebank: Table benchmark for image-based table detection and recognition," *CoRR*, vol. abs/1903.01949, 2019. [Online]. Available: <http://arxiv.org/abs/1903.01949>
- [11] Z. Cai and N. Vasconcelos, "Cascade r-cnn: High quality object detection and instance segmentation," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 43, no. 5, pp. 1483–1498, 2021.
- [12] K. Sun, B. Xiao, D. Liu, and J. Wang, "Deep high-resolution representation learning for human pose estimation," in *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019, pp. 5686–5696.
- [13] K. Chen, J. Wang, J. Pang, Y. Cao, Y. Xiong, X. Li, S. Sun, W. Feng, Z. Liu, J. Xu, Z. Zhang, D. Cheng, C. Zhu, T. Cheng, Q. Zhao, B. Li, X. Lu, R. Zhu, Y. Wu, J. Dai, J. Wang, J. Shi, W. Ouyang, C. C. Loy, and D. Lin, "Mmdetection: Open mmlab detection toolbox and benchmark," *CoRR*, vol. abs/1906.07155, 2019. [Online]. Available: <http://arxiv.org/abs/1906.07155>
- [14] L. Gao, Y. Huang, H. Déjean, J.-L. Meunier, Q. Yan, Y. Fang, F. Kleber, and E. Lang, "Icdar 2019 competition on table detection and recognition (ctdar)," in *2019 International Conference on Document Analysis and Recognition (ICDAR)*, 2019, pp. 1510–1515.
- [15] J. Fang, X. Tao, Z. Tang, R. Qiu, and Y. Liu, "Dataset, ground-truth and performance metrics for table detection evaluation," in *2012 10th IAPR International Workshop on Document Analysis Systems*, 2012, pp. 445–449.
- [16] sgrpanchal31, "sgrpanchal31/table-detection-dataset." [Online]. Available: <https://github.com/sgrpanchal31/table-detection-dataset>
- [17] Doc-Analysis, "doc-analysis/tablebank." [Online]. Available: <https://github.com/doc-analysis/TableBank>
- [18] S. A. Siddiqui, I. A. Fateh, S. T. R. Rizvi, A. Dengel, and S. Ahmed, "Deeptabstr: Deep learning based table structure recognition," in *2019 International Conference on Document Analysis and Recognition (ICDAR)*, 2019, pp. 1403–1409.



Hai-Hong Phan is currently a lecturer of Faculty of Information Technology in Le Quy Don Technical University. She received a Ph.D. degree in computer science from the University of CY Cergy Paris, France, in 2019. She received the B.S. and M.S. degrees in computer science from Le Quy Don Technical University, Vietnam, in 2009 and 2013, respectively. Her researches focus on computer vision, image processing, document analysis, action recognition, and machine learning.
Email: hongpth@lqdtu.edu.vn



Ngo Dai Duong is currently studying for master degree at Le Quy Don Technical University. He received the B.S. degree in Information Technology from Le Quy Don Technical University in 2014. His researches focus on computer vision, document analysis and machine learning.
Email: daiduong28789@hotmail.com

INFORMATION FOR AUTHORS

Journal of Research and Development on Information and Communication Technology - *The MIC Journal of ICT Research*, or *ICT Research* in short – is a scientific publication of the Journal of Information and Communication, published by Vietnam Ministry of Information and Communications (MIC) since 1999. *ICT Research* is peer-reviewed, open-access, and published four times a year: two issues in *English* (ISSN: 1859-3534), and two issues in *Vietnamese* (ISSN: 1859-3526) (Chuyên san Các công trình nghiên cứu phát triển Công nghệ Thông tin và Truyền thông).

The purpose of *ICT Research* is to provide a forum for researchers and professionals to disseminate original and innovative ideas in the fields of information technology, communications, and electronics in Vietnam and worldwide. Its Editorial Board is composed of renowned scientists from research institutions throughout Vietnam and other countries.

AUTHOR GUIDELINES

Authors are encouraged to follow good practice of technical paper writing, such as the guidelines on “How to write for technical periodicals & conferences” of the Institute of Electrical and Electronics Engineers.

Submissions must be made online on the website of *ICT Research* [<https://ictmag.vn>]. Authors should format the manuscript as closely as possible to the *ICT Research* manuscript template.

By submitting a manuscript, authors are required to ensure that the submission has not been previously published, nor is currently submitted for publication elsewhere. Submission also implies that the corresponding author has the consent of all authors. The submission file must be in PDF format. To facilitate double-blind review, authors are requested to submit two versions of the manuscript, one with full information about all authors, the other without.

Submissions of revised manuscripts should follow the same guidelines as for the initial manuscript, plus with a document entitled “Response to reviewers’ comments” which explains the modifications according to the comments of the anonymous reviewers.

Upon formal acceptance of a manuscript for publication, the authors are required to prepare the final manuscript in LaTeX, using the IEEE Transactions style.

PUBLICATION ETHICS

ICT Research is committed to following the rules, standards, and guidelines set forth by the Committee on Publication Ethics (COPE) in ensuring publication integrity and in handling cases of misconduct. All involved parties (Editors, Authors, Reviewers, and Publisher) are expected to be aware of and to conform to the respective rules, standards, and guidelines, available on COPE website.

The Editorial Board of *ICT Research* will immediately screen all articles submitted for publication. All submissions we receive are checked for plagiarism by using available online tools. Any type of plagiarism is unacceptable and is considered unethical publishing behavior. Such manuscripts will be rejected.

PEER REVIEW PROCESS AND PUBLICATION

An area editor of *ICT Research* (to be considered as the Editor from the point of view of the Authors) is assigned, by the Technical Editor-in-Chief, to each submission to coordinate the review process. The Editor pre-screens the submission and determines whether it is appropriate to the interests of readers of *ICT Research*. If appropriate, the Editor then initiates the review process. Each submission is then reviewed by at least two Reviewers. The review process is double blind, that is, identities of authors and reviewers are not revealed to each other.

Review decisions are: (A) *Accept submission* (i.e., no changes required), (B1) *Revision required* (i.e., minor revision), (B2) *Resubmit for review* (i.e., major revision, the submission needs to be re-worked with major changes and is then subject to a second round of review), and (C) *Decline submission* (i.e., reject). Each major round of review may take approximately 1 month. Should authors be requested by the Editor to revise the manuscript, the revised version should be submitted within 4 weeks for a major revision or 2 weeks for a minor revision. No more than 1 major revision and 2 minor revisions are allowed.

ICT Research is published with two issues a year in its *English Edition* (ISSN: 1859-3534). An accepted article will be immediately published online, with a DOI number, having been copyedited and laid out in compliance with *ICT Research* editing rules.

COPYRIGHT NOTICE

Upon acceptance for publication, transfer of copyright will be made to the Publisher of *ICT Research*, who guarantees that full content of the published article is freely distributed on the website of *ICT Research*. The copyright transfer gives the Publisher of *ICT Research* full authority to resolve any complaints of misuse or abuse (such as infringement or plagiarism) of the published article. The author has the freedom to redistribute and reuse the published article in any medium or format for any purpose, provided the original published article is properly cited. An article submission implies author agreement with this policy.



HỘI THẢO QUỐC GIA: MỘT SỐ VẤN ĐỀ CHỌN LỌC CỦA CÔNG NGHỆ THÔNG TIN VÀ TRUYỀN THÔNG

Viện Công nghệ thông tin & Trường Đại học Công nghệ thông tin và truyền thông

Thái Nguyên 13 – 14 / 12 / 2021. Website: <https://vnict.vn>



THÔNG BÁO SỐ 1

Hội thảo Quốc gia lần thứ XXIV “Một số vấn đề chọn lọc của Công nghệ thông tin và Truyền thông” (VNICT) do Viện Công nghệ thông tin – Viện Hàn lâm Khoa học và Công nghệ Việt Nam và Trường Đại học Công nghệ thông tin và truyền thông đồng tổ chức tại Thái Nguyên vào các ngày 13–14 tháng 12 năm 2021, và có sự tham gia phối hợp của Câu lạc bộ các Khoa–Trường–Viện CNTT – TT Việt Nam (FISU). Hội thảo là diễn đàn thường niên để những người làm công tác nghiên cứu, triển khai, giảng dạy và quản lý trong lĩnh vực công nghệ thông tin và truyền thông trao đổi học thuật, chia sẻ kinh nghiệm ứng dụng và tìm kiếm sự hợp tác. Đặc biệt hội thảo khuyến khích các nhà khoa học trẻ, nghiên cứu sinh, học viên cao học tham gia báo cáo trao đổi kết quả nghiên cứu và học tập của mình. Hội thảo khuyến khích trưng bày giới thiệu các sản phẩm khoa học công nghệ và những kết quả nghiên cứu ứng dụng vào thực tiễn. Chủ đề chính của hội thảo năm nay: **Trí tuệ nhân tạo trong chuyển đổi số.**

Một số thời hạn

Toàn văn báo cáo
30 / 7 / 2021

Chấp nhận đăng ký yêu
30 / 9 / 2021

Xuất bản kỷ yếu
30 / 11 / 2021

Đăng ký tham dự
25 / 11 – 5 / 12 / 2021

BAN TỔ CHỨC

Đồng Trưởng Ban

TS Nguyễn Trường Thắng, Viện Công nghệ thông tin
PGS. TS Phùng Trung Nghĩa, Trường Đại học Công nghệ thông tin và truyền thông
TS Nguyễn Văn Tảo, Trường Đại học Công nghệ thông tin và truyền thông

Thành viên

TS Đỗ Đình Cường, Trường Đại học Công nghệ thông tin và truyền thông
PGS. TS Nguyễn Văn Huân, Trường Đại học Công nghệ thông tin và truyền thông
TS Nguyễn Như Sơn, Viện Công nghệ thông tin
TS Bùi Thị Thanh Quyên, Viện Công nghệ thông tin
PGS. TS Bùi Thu Lâm, CLB các Khoa–Trường–Viện CNTT – TT VN (FISU)

CÁC CHỦ ĐỀ CHÍNH CỦA HỘI THẢO

- Trí tuệ nhân tạo
- Học máy
- Xử lý, phân tích dữ liệu lớn
- Xử lý ảnh và thực tại ảo
- Xử lý ngôn ngữ tự nhiên
- Các hệ thống thông minh
- Tin-sinh học
- Công nghệ đa phương tiện và mô phỏng
- Truyền thông và mạng máy tính
- An toàn thông tin
- Kỹ thuật điều khiển, tự động hóa, các hệ thống nhúng
- Cơ sở toán học cho tin học
- CNTT trong kinh tế-xã hội

Thông tin liên hệ

TS Ngô Hải Anh
Viện Công nghệ thông tin
ngohaianh@ioit.ac.vn
☎ 098 859 6582
TS Nguyễn Đức Bình
Trường Đại học CNTT&TT
ndbinh@ictu.edu.vn
☎ 094 701 9219

BAN CHƯƠNG TRÌNH

Đồng Trưởng Ban

PGS. TS Nguyễn Long Giang, Viện Công nghệ thông tin
TS Vũ Đức Thái, Trường Đại học Công nghệ thông tin và truyền thông

Thành viên

- ▶ GS. TS Đặng Quang Á, Viện Hàn lâm Khoa học và Công nghệ Việt Nam
- ▶ TS Nguyễn Thu Anh, Viện Công nghệ thông tin
- ▶ TS Nguyễn Việt Anh, Viện Công nghệ thông tin
- ▶ PGS. TS Huỳnh Thị Thanh Bình, Trường Đại học Bách Khoa Hà Nội
- ▶ PGS. TS Nguyễn Thanh Bình, Trường Đại học Bách khoa Đà Nẵng
- ▶ PGS. TS Phạm Việt Bình, Trường Đại học CNTT và truyền thông
- ▶ GS. TS Nguyễn Bường, Viện Công nghệ thông tin
- ▶ TS Nguyễn Văn Căn, Trường Đại học Kỹ thuật-Hậu cần CAND
- ▶ TS Phạm Quốc Chính, Sở Khoa học và Công nghệ Thái Nguyên
- ▶ TS Nguyễn Ngọc Cương, Cục ANM và PCTP sử dụng CNC – Bộ Công An
- ▶ PGS. TS Trịnh Viết Cường, Trường Đại học Hồng Đức
- ▶ TS Hà Mạnh Đào, Trường Đại học Công nghiệp Hà Nội
- ▶ TS Lương Khắc Định, Trường Đại học Hà Long
- ▶ PGS. TS Đặng Văn Đức, Viện Công nghệ thông tin
- ▶ PGS. TS Lê Bá Dũng, Viện Công nghệ thông tin
- ▶ PGS. TS Nguyễn Đức Dũng, Viện Công nghệ thông tin
- ▶ TS Ngô Quốc Dũng, Học viện Công nghệ Bưu chính Viễn thông
- ▶ PGS. TS Lương Thế Dũng, Học viện Kỹ thuật Mật mã
- ▶ PGS. TS Phạm Thanh Giang, Viện Công nghệ thông tin
- ▶ TS Nguyễn Thị Thu Hà, Trường Đại học Điện lực
- ▶ TS Trương Hà Hải, Trường ĐH Công nghệ thông tin và truyền thông
- ▶ PGS. TS Huỳnh Xuân Hiệp, Trường Đại học Cần Thơ
- ▶ TS Nguyễn Hữu Hòa, Trường Đại học Cần Thơ
- ▶ PGS. TS Nguyễn Ngọc Hóa, Trường Đại học Công nghệ – DHQG Hà Nội
- ▶ PGS. TSKH Nguyễn Xuân Huy, Viện Công nghệ thông tin
- ▶ PGS. TS Trần Đăng Hưng, Trường Đại học Sư phạm Hà Nội
- ▶ TS Nguyễn Quang Hưng, Tạp chí Thông tin Thông tin và Truyền thông
- ▶ PGS. TS Trần Văn Lãng, Viện Hàn lâm Khoa học và Công nghệ Việt Nam
- ▶ PGS. TS Ngô Thành Long, Học viện Kỹ thuật Quân sự
- ▶ PGS. TS Hoàng Việt Long, Trường Đại học Kỹ thuật-Hậu cần CAND
- ▶ TS Nguyễn Thị Lương, Trường Đại học Đà Lạt
- ▶ PGS. TS Lương Chi Mai, Viện Công nghệ thông tin
- ▶ PGS. TS Nguyễn Thị Hồng Minh, Trường Đại học KHTN – DHQG Hà Nội
- ▶ TS Nguyễn Duy Minh, Trường DH Công nghệ thông tin và truyền thông
- ▶ TS Nguyễn Hải Minh, Trường DH Công nghệ thông tin và truyền thông
- ▶ TS Phạm Ngọc Minh, Viện Công nghệ thông tin
- ▶ TS Lê Quang Minh, Đại học Quốc gia Hà Nội
- ▶ TS Trần Đức Nghĩa, Viện Công nghệ thông tin
- ▶ PGS. TS Lê Anh Phương, Trường Đại học Sư phạm Huế
- ▶ GS. TS Từ Minh Phương, Học viện Công nghệ Bưu chính Viễn thông
- ▶ TS Vũ Vinh Quang, Trường DH Công nghệ thông tin và truyền thông
- ▶ PGS. TS Nguyễn Hữu Quỳnh, Trường Đại học Thủy lợi
- ▶ PGS. TS. Ngô Hồng Sơn, Trường Đại học Phenikaa
- ▶ PGS. TS Lê Hoàng Sơn, Đại học Quốc gia Hà Nội
- ▶ PGS. TS Nguyễn Văn Tam, Viện Công nghệ thông tin
- ▶ TS Trần Minh Tân, Trường Đại học Giao thông vận tải
- ▶ PGS. TS Ngô Quốc Tào, Viện Công nghệ thông tin
- ▶ PGS. TS Huỳnh Quyết Thắng, Trường Đại học Bách Khoa Hà Nội
- ▶ TS Nguyễn Trường Thắng, Viện Công nghệ thông tin
- ▶ TS Trần Thiên Thành, Trường Đại học Quy Nhơn
- ▶ GS. TS Vũ Đức Thi, Trường Đại học Thăng Long
- ▶ TS Nguyễn Văn Thiên, Trường Đại học Công nghiệp Hà Nội
- ▶ GS. TS Nguyễn Thanh Thủy, Đại học Quốc gia Hà Nội
- ▶ TS Phạm Thị Thu Thủy, Trường Đại học Nha Trang
- ▶ PGS. TS Trần Minh Triết, Trường ĐHKH tự nhiên TPHCM
- ▶ TS Lê Anh Tú, Trường Đại học Hà Long
- ▶ TS Hoàng Đỗ Thanh Tùng, Viện Công nghệ thông tin
- ▶ PGS. TS Thái Quang Vinh, Viện Công nghệ thông tin
- ▶ TS Nguyễn Trần Quốc Vinh, Trường Đại học Sư phạm Đà Nẵng
- ▶ PGS. TS Lê Sỹ Vinh, Trường Đại học Công nghệ – DHQG Hà Nội
- ▶ PGS. TS Lê Trọng Vĩnh, Trường Đại học KHTN – DHQG Hà Nội
- ▶ PGS. TS Vũ Việt Vũ, Đại học Quốc gia Hà Nội

TÀI TRỢ

Viện Hàn lâm Khoa học và Công nghệ Việt Nam



Viện Công nghệ thông tin



Trường Đại học Công nghệ thông tin và truyền thông



Câu lạc bộ các Khoa–Trường–Viện CNTT – TT Việt Nam



Học viện Khoa học và Công nghệ



Tạp chí Thông tin và Truyền thông

THÔNG TIN & TRUYỀN THÔNG

BAN THƯ KÝ

TS Ngô Hải Anh, Viện Công nghệ thông tin
ThS Dương Thị Nhung, Trường Đại học Công nghệ thông tin và truyền thông
ThS Nguyễn Trần Ánh, Trường Đại học Công nghệ thông tin và truyền thông

KỶ YẾU HỘI THẢO & TUYỂN CHỌN BÀI ĐĂNG TẠP CHÍ

- ▶ Các bài gửi tham gia Hội thảo (theo định dạng của IEEE, gửi qua hệ thống EasyChair: <https://easychair.org/conferences/?conf=vnict2021>) có chất lượng tốt sau khi qua vòng phân biên sẽ được lựa chọn đăng trên kỷ yếu có mã số ISBN, được xuất bản bởi NXB Khoa học và Kỹ thuật (Bộ Khoa học và Công nghệ).
- ▶ Các bài xuất sắc (tiếng Anh/tiếng Việt) được chọn sẽ đăng trên Chuyên san Các công trình nghiên cứu, phát triển và ứng dụng Công nghệ Thông tin và Truyền thông (Tạp chí Thông tin và Truyền thông – Bộ Thông tin và Truyền thông), và vẫn báo cáo slide tại hội thảo.

Phí tham dự hội thảo: 1.000.000 (tác giả báo cáo), 500.000 (khách tham dự không báo cáo).