

VOL. 11 NO. 30 JULI 2023

ISSN: 2253-4765



JURNAL RESTI

REKAYASA SISTEM DAN
TEKNOLOGI INFORMASI



Indexed in Scopus
Indexed in the Journal Index of
Higher Education Research and Technology
(IJERT) (2023)



Indexed in Garuda
Indexed in the Garuda database
(2023)

Indexed in Dimensions
Indexed in the Dimensions database
(2023)

Google Scholar, Garuda, Dimensions, Crossref, Scopus, WJCI, SJ, IJERT, and others.

Sources

ISSN

Enter ISSN or ISSNs

Find sources

ISSN: 2580-0760 x

CiteScore 2024 has been released. [View CiteScore methodology](#)

Filter refine list

Apply Clear filters

Display options

☐ Display only Open Access journals

Counts for 4-year timeframe

☒ No minimum selected

☐ Minimum citations

☐ Minimum documents

CiteScore highest quartile

☐ Show only titles in top 10 percent

☐ 1st quartile

☐ 2nd quartile

☒ 3rd quartile

☐ 4th quartile

1 result

[Download Scopus Source List](#) [Learn more about Scopus Source List](#)

☐ All [Export to Excel](#) [Save to source list](#)

View metrics for year: 2024

	Source title ↓	CiteScore ↓	Highest percentile ↓	Citations 2021-24 ↓	Documents 2021-24 ↓	% Cited ↓
☐ 1	Jurnal RESTI: Open Access	1.8	33% 318/474 Information Systems	973	547	61

[Top of page](#)

SERTIFIKAT

Direktorat Jenderal Pendidikan Tinggi, Riset dan Teknologi
Kementerian Pendidikan, Kebudayaan, Riset dan Teknologi Republik Indonesia



Kutipan dari Keputusan Direktorat Jenderal Pendidikan Tinggi, Riset dan Teknologi
Kementerian Pendidikan, Kebudayaan, Riset, dan Teknologi Republik Indonesia

Nomor 158/E/KPT/2021

Peringkat Akreditasi Jurnal Ilmiah Periode 1 Tahun 2021

Nama Jurnal Ilmiah

Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)

E-ISSN: 25800760

Penerbit: Ikatan Ahli Informatika Indonesia (IAII)

Ditetapkan Sebagai Jurnal Ilmiah

TERAKREDITASI PERINGKAT 2

Akreditasi Berlaku selama 5 (lima) Tahun, yaitu
Volume 5 Nomor 2 Tahun 2021 Sampai Volume 10 Nomor 1 Tahun 2026

Jakarta, 09 December 2021

Plt. Direktur Jenderal Pendidikan Tinggi,
Riset, dan Teknologi



Prof. Ir. Nizam, M.Sc., DIC, Ph.D., IPU, ASEAN Eng
NIP. 196107061987101001



[Back](#)

A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids

[Jurnal RESTI](#) • Article • Open Access • 2025 • DOI: 10.29267/resti.v9i3.6533 
[Latifa Riyana Putri, Alim^a !\[\]\(034433b90593e82e5460e34e3ed48e9b_img.jpg\) !\[\]\(5f24500834b50a8307ffe63e419281a9_img.jpg\) !\[\]\(8502790a2fc8970abb55aa7f7a7a6bae_img.jpg\) !\[\]\(7f7375e819602f97f5594a28879b3e09_img.jpg\) !\[\]\(f947a10a261625f35f2b4defe4317e0a_img.jpg\) !\[\]\(56046418e888bdf1ca31370e2b118597_img.jpg\) !\[\]\(d93a559e2bd97e29d4d4efb058c4f350_img.jpg\) !\[\]\(df876734583bc80790adabc02c754ee7_img.jpg\) !\[\]\(ef6654fc5586f194bc1389c0ae8e18df_img.jpg\) !\[\]\(4bcf021fe01691f91b324a77ad27a4d1_img.jpg\) !\[\]\(7a0b7e977c220ef478244081a7c8bb28_img.jpg\) !\[\]\(0089474fb99fa1cfb945e8dddb63dd4c_img.jpg\) !\[\]\(b1b9829b54c0d6a26865841344606da3_img.jpg\) !\[\]\(2e6c6a42d8cae155608272f69cc9c3da_img.jpg\) !\[\]\(e9ce0f999b4b6b71c25622a51d4a9513_img.jpg\) !\[\]\(c5c9f8ebe39c1befc3a8bf99d4b2b978_img.jpg\) !\[\]\(09eba22b72ed3e15470cd6055e931a88_img.jpg\) !\[\]\(51f1b3ebfa05b9b49b44219cffa26038_img.jpg\) !\[\]\(6e9330a86001d6d32ac4e97412d2b61b_img.jpg\) !\[\]\(a04f46fc0412dc6de230f8a3b66c17bc_img.jpg\) !\[\]\(8b7b940528e960e072bc6cad630baf84_img.jpg\) !\[\]\(ac9454f5ffa76a68e2251033f1f74348_img.jpg\) !\[\]\(fe6ac115033689df4c4fa5568ca6f7b0_img.jpg\) !\[\]\(bbe5bafde4f7e6c040bb6724edc51618_img.jpg\) !\[\]\(205da0096b84be7192c682c55ab438e7_img.jpg\) !\[\]\(bad4295508e019fd9991fad0860fa9fd_img.jpg\) !\[\]\(207d585292fe4689674637880b2be6a0_img.jpg\) !\[\]\(8597c5193f43442791e3200acce2d571_img.jpg\) !\[\]\(fc8fced5951e0646d8170dd4ec15612b_img.jpg\) !\[\]\(f1ab1c3feae6acc13e31c1d9addb8710_img.jpg\) !\[\]\(f62d310dba1b7f7b1ba14fa002ecbe7f_img.jpg\) !\[\]\(de2c5574202a161639d60876eee6ac36_img.jpg\) !\[\]\(0726bf9efac954cc42a99454a79ca34e_img.jpg\) !\[\]\(90cf3d7343dffc377fbe17029ba182e5_img.jpg\)](#) [Eni Puliastruti, Christina^b !\[\]\(bdd5c53a89a3cc6a44750bd18c63d942_img.jpg\) !\[\]\(cbb6b357ea6cb537b07bff3b51c991e3_img.jpg\) !\[\]\(bafb1f88e2ae382fd6dc4f762f8885ef_img.jpg\) !\[\]\(0a4fa8454464bf3c89cdf8140fbeaaef_img.jpg\) !\[\]\(8f1b3539954c45aea9797eb973f0b38f_img.jpg\) !\[\]\(7584f580d1b3802b79a212b5d30e5236_img.jpg\) !\[\]\(ec3d29b34aa6c77e7126193651816158_img.jpg\) !\[\]\(eaa9809ae2171979392f5830c0053dfe_img.jpg\) !\[\]\(f2ed3b16cb2e11723352ced0cecfdf22_img.jpg\) !\[\]\(b690be2b18d7136c1e91c6e6258e892f_img.jpg\) !\[\]\(d9364e8cf9ca632f0ee10b80519e7839_img.jpg\) !\[\]\(ced058e174dd7667ca54b21ba810f35a_img.jpg\) !\[\]\(1030075b462ac060a4e3875ac0e35dea_img.jpg\) !\[\]\(9cd71421df0c6869c47844ffc67fd1ef_img.jpg\) !\[\]\(1570abdf2fe4454b87a79c1b8f347da6_img.jpg\) !\[\]\(43d3eaa4dc8587047e496c4aef6eeb29_img.jpg\) !\[\]\(c90116d44a71b40452e8d57450b536d0_img.jpg\) !\[\]\(e91a96ec81cba51272fb19b36149d678_img.jpg\) !\[\]\(de60486c4306d1b59071923892ce5b99_img.jpg\) !\[\]\(7adea749849490023db7b4fe9eaf5026_img.jpg\) !\[\]\(ba9061a9b458924089ed95160037f680_img.jpg\) !\[\]\(6a085937a2fbe09f4127d7a692c560d3_img.jpg\) !\[\]\(4d916582786af621b5911134ebd4c9d7_img.jpg\) !\[\]\(e9713469bbed47fdea7b8c90f85d08e2_img.jpg\) !\[\]\(6e4d8863c6525dbac2775706f2a46e2c_img.jpg\) !\[\]\(89d326cc0e3583ae9c9f42bd37b4efef_img.jpg\) !\[\]\(87c8ebb6c1cc1091db16c11b4e962cb0_img.jpg\) !\[\]\(ebc7c873738aa13ab78b45c229b78f4d_img.jpg\)](#) [Supriyandi^b !\[\]\(2bc570634568fb49116c65924bdad735_img.jpg\) !\[\]\(e3a441992ff8b8d780773aa329918e3e_img.jpg\) !\[\]\(fecb17514be1741ca02da7edcde7a87b_img.jpg\) !\[\]\(e41ff91b0a8e63e68408939ed99e4a56_img.jpg\) !\[\]\(c4f264624451976d27411b309951a0e2_img.jpg\) !\[\]\(1445434b3b3e1fb2d86dd546519f7f43_img.jpg\) !\[\]\(8102c0525dabc5b4708b008913433b7a_img.jpg\) !\[\]\(6d5d8e294d4beb13401085b274a0fb60_img.jpg\) !\[\]\(98d5978a0748433762b662501353b0e2_img.jpg\) !\[\]\(7ec44d55b910e1980b423923b5c4f2a3_img.jpg\) !\[\]\(caec2c8c53e096a6cb4c6c35482dbd43_img.jpg\) !\[\]\(f777f6205c58db93cbe3c4ecb344ceae_img.jpg\) !\[\]\(d13d537c5b6ae136074628d20fbeb2db_img.jpg\) !\[\]\(0046465f42956e180dbb1b3270aba071_img.jpg\) !\[\]\(63aacf6b120de6f9bb5ec670360ebfcc_img.jpg\) !\[\]\(085c988b025ebaf1955763efadbae8fd_img.jpg\) !\[\]\(212089ac707c2d7f4e8253e44fe5148a_img.jpg\) !\[\]\(bd628018d553b101f7f1b4f89959356d_img.jpg\) !\[\]\(e169849ce099ca3ff986e5042f65bbaa_img.jpg\) !\[\]\(0d4e963ca0bfa2d85901b9134f4c83eb_img.jpg\) !\[\]\(d9e60ee787a7174316a2f73c923b6dc1_img.jpg\) !\[\]\(17a7a828c256f1a0a210871e1cf64193_img.jpg\) !\[\]\(32edec5c549a5714382771ae5b8512b4_img.jpg\) !\[\]\(4029ced160133fdef7e141bf5557768a_img.jpg\) !\[\]\(d68096c048a6a9c36c8edc3a7273940d_img.jpg\) !\[\]\(e25d92feb8de5686041c986cd3e602f9_img.jpg\)](#)

^aDato Science, Telkom University, Purwokerto, Indonesia

[Show all information](#)

[View PDF](#)[Full text](#)[Export](#)[Save to list](#)

[Document](#) [Impact](#) [Cited by \(0\)](#) [References \(23\)](#) [Similar documents](#)

Abstract

The K-Means algorithm is a popular unsupervised learning method used for data clustering. However, its performance heavily depends on centroid initialization and the distribution shape of the data, making it less effective for datasets with complex or non-linear cluster structures. This study evaluates the performance of the standard K-Means algorithm and proposes a Multiobjective Particle Swarm Optimization K-Means (MOPSO-K-Means) approach to improve clustering accuracy. The evaluation was conducted on five benchmark datasets: Atom, Chainlink, EngTime, Target, and TwoDiamonds. Experimental results show that K-Means is effective only on datasets with clearly separated clusters, such as EngTime and TwoDiamonds, achieving accuracies of 95.6% and 100%, respectively. In contrast, MOPSO-K-Means achieved a substantial accuracy improvement on the complex Target dataset, increasing from 0.26% to 59.2%. The TwoDiamonds dataset achieved the most desirable trade-off: it had the lowest SSW (1521.32), relatively high SSB (2863.34), and lowest standard deviation values, indicating compact clusters, good separation, and high consistency across runs. These findings highlight the potential of swarm-based optimization to achieve consistent and accurate clustering results on datasets with varying structural complexity. © 2025, Ikatan Ahli Informatika Indonesia. All rights reserved.

Author keywords

centroid; k-means; multiobjective particle swarm optimization; the sum of square between; the sum of square within

Funding details

Details about financial support for research, including funding sources and grant numbers as provided in academic publications.

Funding sponsor	Funding number	Acronym
Universitas Telkom		
See opportunities		

Funding text

This work was supported by Telkom University.

Corresponding authors

Corresponding author	A. Latifa Riyana Putri
Affiliation	Dato Science, Telkom University, Purwokerto, Indonesia
Email address	airaap@telkomuniversity.ac.id

© Copyright 2025 Elsevier B.V. All rights reserved.

About Scopus

[What is Scopus](#)
[Content coverage](#)
[Scopus blog](#)
[Scopus API](#)
[Privacy matters](#)

Language

[日本語版を表示する](#)
[查看繁体中文版本](#)
[查看繁體中文版本](#)
[Просмотр версии на русском языке](#)

Customer Service

[Help](#)
[Tutorials](#)
[Contact us](#)

0

[Citations](#) 

Abstract

[Author keywords](#)[Funding details](#)[Corresponding authors](#)

Abstract

[Author keywords](#)[Funding details](#)[Corresponding authors](#)



JURNAL RESTI

Rekayasa Sistem dan Teknologi Informasi

ISSN Media Elektro
2580-0760

www.jurnal.iaii.or.id

[Home](#)[Current Issue](#)[Tutorial ▾](#)[About Us ▾](#)[Announcements](#)[Archives](#)[Dashboard 0](#)[Search](#)[Logout](#)[Home](#) / [Editorial Team](#)

Editor-in-Chief:

Yuhefizar

Politeknik Negeri Padang, Padang-Indonesia

[Scopus](#) | [Google Scholar](#) | [Orcid](#)

Managing Editor:

Ronal Watrianthos

Politeknik Negeri Padang, Padang-Indonesia

[ResearcherID](#) | [Scopus](#) | [Google Scholar](#) | [Orcid](#)

Rita Komalasari

Politeknik LP31, Bandung-Indonesia

[ResearcherID](#) | [Scopus](#) | [Google Scholar](#) | [Orcid](#)

Budi Sunaryo

Universitas Bung Hatta, Padang-Indonesia

[ResearcherID](#) | [Scopus](#) | [Google Scholar](#) | [Orcid](#)

Associates Editor:

Ikhwan Arief



1.8

2024
CiteScore

33rd percentile

Powered by [Scopus](#)



..MAIN MENU..

[About Journal](#)

[Editorial Team](#)

[Peer Review Process](#)

[Focus & Scope](#)

Ambassador to Indonesia - Directory of Open
Access Journals (DOAJ), Roskilde, **Denmark**

Research areas: Information Systems, Engineering, Manufacturing

[Scopus](#) | [Orcid](#)

Dini Oktarina Dwi Handayani

International Islamic University Malaysia, Kuala
Lumpur-**Malaysia**

Research areas: Artificial Intelligence

[Scopus](#) | [Orcid](#)

Windu Gata

Universitas Nusa Mandiri, Jakarta-Indonesia

Research areas: IOT; Data Science

[Scopus](#) | [Orcid](#)

Shoffan Saifullah

AGH University of Krakow, **Poland**

Research areas: image processing; medical image analysis

[Scopus](#) | [Orcid](#)

Teddy Mantoro

Universitas Nusa Putra, Sukabumi-Indonesia

*Research areas: Artificial Intelligence, Deep Learning, Information
Security, Mobile Computing*

[Scopus](#) | [Orcid](#)

Editorial Board Members:

Media Anugerah Ayu

Sampoerna University, Indonesia

[Scopus](#) | [Orcid](#)

Zairi Ismael Rizman

Universiti Teknologi MARA, Shah Alam, **Malaysia**

[Scopus](#) | [Orcid](#)

Diki Arisandi

Universitas Muhammadiyah Riau, Indonesia

[Scopus](#) | [Orcid](#)

Hendrick

Politeknik Negeri Padang, Indonesia

Author's Guide

Authors Fee

Publication Ethics

Open Access Statement

Plagiarism Policy

Archive Policy

Copyright and License

Author's Statement

F.A.Q

Download RESTI Templates

...VISITOR STATISTICS...

...:INFORMATION:...

For Reader

For Author

...:Find Us:...

Google Scholar

Garuda

SINTA

CrossRef

Dimensions

...:SUPERVISED BY:...





[Scopus](#) | [Orcid](#)

Tri Apriyanto Sundara

Universiti Kebangsaan Malaysia, **Malaysia**

[Scopus](#) | [Orcid](#)

Yance Sonatha

Politeknik Negeri Padang, Indonesia

[Scopus](#) | [Orcid](#)

Mohd Helmy Abd Wahab

Universiti Tun Hussein Onn, Batu Pahat, **Malaysia**

[Scopus](#) | [Orcid](#)

Madjid Eshaghi Gordji

Semnan University, Semnan, **Iran**

[Scopus](#) | [Orcid](#)

Shipra Shivkumar Yadav

RTM Nagpur University, **India**

[Scopus](#) | [Orcid](#)

Roheen Qamar

Quaid-e-Awam University of Engineering,
Nawabshah, **Pakistan**

[Scopus](#) | [Orcid](#)

Hari M. Srivastava

University of Victoria, **Canada**

[Scopus](#) | [Orcid](#)

Sasalak Tongkaw

Songkhla Rajabhat University, Songkla-**Thailand**

[Scopus](#) | [Orcid](#)

Ambrose Agbon Azeta

Namibia University of Science and Technology,
Windhoek- **Namibia**

[Scopus](#) | [Orcid](#)

Wasswa Shafik

Dig Connectivity Research Laboratory, Health,
Ecology and Technology Research, **Uganda**

[Scopus](#) | [Orcid](#)

Haydar Abdulameer Marhoon

Al-Ayen Iraqi University, AUIQ, An Nasiriyah, **Iraq**

[Scopus](#) | [Orcid](#)

Yeshi Jamtsho

Royal University of Bhutan, Thimphu, **Bhutan**

[Scopus](#) | [Orcid](#)

Mutaz Rasmi Abu Sara

Palestine Ahliya University, Bethlehem, **Palestine**

[Scopus](#) | [Orcid](#)

Arif Bramantoro

Universiti Teknologi Brunei, Gadong, **Brunei**

Darussalam

[Scopus](#) | [Orcid](#)

Basanta Joshi

Institute of Engineering Pulchowk, Kathmandu,

Nepal

[Scopus](#) | [Orcid](#)

Nayma Cepero-Pérez

Universidad Tecnológica de la Habana José

Antonio Echeverría **Cuba**

[Scopus](#) | [Orcid](#)

Neema Mduma

Nelson Mandela African Institution of Science and

Technology, Arusha, **Tanzania**

[Scopus](#) | [Orcid](#)

Haruna Jallow

The Pan African University Institute for Basic

Sciences, Technology and Innovation (PAUSTI),

Nairobi, **Kenya**

[Scopus](#) | [Orcid](#)

Christian Utama

Freie Universität Berlin, Berlin, **Germany**

[Scopus](#) | [Orcid](#)

Inés López-Baldominos

Universidad de Alcalá, Alcala de Henares, **Spain**

[Scopus](#) | [Orcid](#)

Dafina Hyka

Polytechnic University of Tirana, Tirana, **Albania**

[Scopus](#) | [Orcid](#)

Technical Support:

Roni Putra

Politeknik Negeri Padang, Indonesia

[Scopus](#)

Copy Editor:

Gilang Surendra

Politeknik Negeri Padang, Indonesia

Click [here](#) for more details

RESTI (Rekayasa Sistem dan Teknologi Informasi) Journal indexed by:



Published by Professional Organization Ikatan Ahli Informatika Indonesia (IAII)/Indonesian Informatics Experts Association

Website Publisher : www.iaii.or.id - Website Journal : www.jurnal.iaii.or.id/index.php/RESTI - Email : jurnal.resti@gmail.com

Halaman Facebook : <http://fb.me/jurnalRESTI>





JURNAL RESTI

Rekayasa Sistem dan Teknologi Informasi

ISSN Media Elektro
2580-0760

www.jurnal.iaii.or.id

[Home](#)
[Current Issue](#)
[Tutorial ▾](#)
[About Us ▾](#)
[Announcements](#)
[Archives](#)
[Dashboard 0](#)
 Search

[Logout](#)
[Home](#) / [Archives](#) / Vol 9 No 3 (2025): June 2025


This issue features 28 scholarly articles from 32 distinguished institutions in Indonesia, two from Malaysia, and one from Australia. The editors thank all the researchers who chose the RESTI Journal as a platform to spread the fruits of their labor. Your contributions enrich not only our publication but also the broader scientific community, particularly in the ever-evolving field of information technology.

The cover, foreword, and table of contents can be downloaded [here](#).



1.8 2024 CiteScore

33rd percentile

Powered by Scopus



..MAIN MENU..

[About Journal](#)

[Editorial Team](#)

[Peer Review Process](#)

[Focus & Scope](#)

DOI: <https://doi.org/10.29207/resti.v9i3>

Published: 2025-05-24

Computer Science Applications

Classification of Red Foxes: Logistic Regression and SVM with VGG-16, VGG-19, and Inception V3

Brian Sabayu, Imam Yuadi 435 - 443



Comparative Evaluation of Preprocessing Methods for MobileNetV1 and V2 in Waste Classification

Aulia Afifah, Endah Ratna Arumi, Maimunah Maimunah, Setiya Nugroho 444 - 452



University Students Stress Detection During Final Report Subject by Using NASA TLX Method and Logistic Regression

Alfita Khairah, Melinda, Iskandar Hasanuddin, Didi Asmadi, Riski Arifin, Rizka Miftahujjannah 465 - 476



Expertise Retrieval Using Adjusted TF-IDF and Keyword Mapping to ACM Classification Terms

Lyla Ruslana Aini, Evi Yulianti 497 - 505



Development of MongoDB-based Gait System with Interactive Visualization for Clinical Analysis

Author's Guide

Authors Fee

Publication Ethics

Open Access Statement

Plagiarism Policy

Archive Policy

Copyright and License

Author's Statement

F.A.Q

Download RESTI Templates

...VISITOR STATISTICS...

...INFORMATION...

For Reader

For Author

...Find Us...

Google Scholar

Garuda

SINTA

CrossRef

Dimensions

...SUPERVISED BY...





Rizal Rahman Rizkika, Helisyah Nur Fadhilah, Tanzilal Mustaqim, Rifdatun Ni'mah 554 - 563



XGBoost Algorithm for Cervical Cancer Risk Prediction: Multi-dimensional Feature Analysis

Sudi Suryadi, Masrizal 619 - 625



Improving Frame-based Engagement Classification in E-Learning Using EfficientNet and Normalized Loss Weighting

Joseph Ananda Sugihdharma, Fitra Bachtiar, Novanto Yulistira 635 - 645



Artificial Intelligence

Forecasting Stock Returns Using Long Short-Term Memory (LSTM) Model Based on Inflation Data and Historical Stock Price Movements

Nur Faid Prasetyo, Wina Witanti, Asep Id Hadiana 453 - 464



The Impact of Cancer on Poverty: An Analytical Study Using Big Data and OLS Regression

Heny Pratiwi, Muhammad Ibnu Sa'ad, Wahyuni Wahyuni, Syamsuddin Mallala 487 - 496



Obesity Status Prediction Through Artificial Intelligence and Balanced Label Distribution Using SMOTE

Arif Riyandi, Mahazam Afrad, M Yoka Fathoni, Yogo Dwi Prasetyo 519 - 524



Benchmarking Metaheuristic Algorithms Against Optimization Techniques for Transportation Problem in Supply Chain Management

Felicia Lim Xin Ying, Suliadi Firdaus Sufahani 525 - 534



Health Risk Classification Using XGBoost with Bayesian Hyperparameter Optimization

Syaiful Anam, Imam Nurhadi Purwanto, Dwi Mifta Mahanani, Feby Indriana Yusuf, Hady Rasikhun 535 - 543



Prediction of Financial Distress in Retail Companies Using Long-Short Term Memory (LSTM)

Wahyuni Windasari, Tuti Zakiyah 564 - 569



Enhancing Tomato Leaf Disease Detection via Optimized VGG16 and Transfer Learning Techniques

Sandy Putra Siregar, Imam Akbari, Poningsih Poningsih, Anjar Wanto, Solikhun Solikhun 570 - 580



UDAWA Gadadar: Agent-based Cyber-physical System for Universal Small-scale Horticulture Greenhouse Management System

I Wayan Aditya Suranata, Ketut Elly Sutrisni, I Made Surya Adi Putra 581 - 593



Performance Comparison of Monolithic and Microservices Architectures in Handling High-Volume Transactions

Mastura Diana Marieska, Arya Yunanta, Harisatul Aulia, Alvi Syahrini Utami, Muhammad Qurhanul Rizqie 594 - 600



Optimizing Sensitivity in Machine Learning Models for Pediatric Post-operative Kyphosis Prediction

Raja Ayu Mahessya, Dian Eka Putra, Rostam Ahmad Efendi, Rayendra, Rozi Meri, Riyan Ikhbal Salam, Dedi Mardianto, Ikhsan, Ismael, Arif Rizki Marsa 601 - 608



Development of IoT-based Automatic Water Drainage System on Fishing Boat to Improve Operational Efficiency

Zulfachmi Zulfachmi, Zulkipli Zulkipli, Vita Rahayu, Aggry Saputra, Muthiah As Saidah 609 - 618



A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids

Aina Latifa Riyana Putri, Joko Riyono, Christina Eni Pujiastuti, Supriyadi 626 - 634



Enhancing Stroke Prediction with Logistic Regression and Support Vector Machine Using Oversampling Techniques

Syamsul Risal, Fajar Apriyadi, A. Sumardin, Andini Dani Achmad, Annisa Nurul Puteri 646 - 658



A New Framework for Dynamic Educational Marketing Segmentation in Student Recruitment: Optimizing Fuzzy C-Means with Metaheuristic Techniques

Rizal Bakri, Bobur Sobirov, Niken Probondani Astuti, Ansari Saleh Ahmar, Pawan Kumar Singh 659 - 669



Stunting Prediction Modeling in Toddlers Using a Machine Learning Approach and Model Implementation for Mobile Application

Eko Abdul Goffar, Rosa Eliviani, Lili Ayu Wulandhari 670 - 676



Minangkabau Language Stemming: A New Approach with Modified Enhanced Confix Stripping

Fadhli Almu'iini Ahda, Aji Prasetya Wibawa, Didik Dwi Prasetya, Danang Arbian Sulistyo, Andrew Nafalski 677 - 687



Automatic Classification of Multilanguage Scientific Papers to the Sustainable Development Goals Using Transfer Learning

Lya Hulliyyatus Suadaa, Anugerah Karta Monika, Berliana Sugiarti Putri, Yeni Rimawati 688 - 696



Enhancing Areca Nut Detection and Classification Using Faster R-CNN: Addressing Dataset Limitations with Haar-like Features, Integral Image, and Anchor Box Optimization

Yovi Pratama, Errissya Rasywir, Suyanti, Agus Siswanto, Fachruddin 697 - 705



Computer Vision and Pattern Recognition

Classification of Retinoblastoma Eye Disease on Digital Fundus Images Using Geometric Features and Machine Learning

Arif Setiawan 477 - 486



The Effect of Hyperparameters on Faster R-CNN in Face Recognition Systems

Jasman Pardede, Khairul Rijal

506 - 518

**Automated Indonesian Plate Recognition: YOLOv8 Detection and TensorFlow-CNN Character Classification**

Windu Gata, Dwiza Riana, Muhammad Haris, Maria Irminda Prasetyowati, Dika Putri Metalica

544 - 553



RESTI (Rekayasa Sistem dan Teknologi Informasi) Journal indexed by:



Published by Professional Organization Ikatan Ahli Informatika Indonesia (IAII)/Indonesian Informatics Experts Association

Website Publisher : www.iaii.or.id - Website Journal : www.jurnal.iaii.or.id/index.php/RESTI - Email : jurnal.resti@gmail.comHalaman Facebook : <http://fb.me/jurnalRESTI>License: [Creative Commons Attribution 4.0 International License](https://creativecommons.org/licenses/by/4.0/).





License: [Creative Commons Attribution 4.0 International License](https://creativecommons.org/licenses/by/4.0/).





J. Riyono <jokoriyono@trisakti.ac.id>

[RESTI] Editor Decision

1 message

chief <jurnal@iaii.or.id>

Wed, Jun 4, 2025 at 12:48 PM

To: Aina Latifa Riyana Putri <ainaqp@telkomuniversity.ac.id>, Joko Riyono <jokoriyono@trisakti.ac.id>, Christina Eni Pujiastuti <christina.eni@trisakti.ac.id>

Aina Latifa Riyana Putri, Joko Riyono, Christina Eni Pujiastuti:

We have reached a decision regarding your submission to Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi), "A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids".

Our decision is to: Accept Submission with minor revision

Please follow : <http://jurnal.iaii.or.id/index.php/reSTI/accepted> for requirement publish.

chief
Ikatan Ahli Informatika Indonesia (IAII) Nusantara
jurnal@iaii.or.id

Reviewer A:
Recommendation: Accept Submission

Is the title appropriate and aligned with the journal's scope?

Yes

Does the abstract concisely inform readers and summarize the research?

Yes

Are relevant prior studies adequately contextualized, especially in the background?

Yes

Does the author provide in-depth analysis, particularly in the discussion and results sections?

Yes

Does the author demonstrate sufficient scientific reasoning, argumentation and interpretation?

Yes

Are descriptions and explanations presented clearly and understandably?

Yes

Are figures, tables, and numbers of adequate quality and easy to interpret?

Yes

Are visuals like diagrams and images used appropriately and clearly?

Yes

Does the manuscript present an original contribution?

Yes

Is the writing style suitably academic with clear and proper language?

Yes

Give constructive feedback:

1. Enhance the abstract by including more specific numerical results, particularly the most significant improvements (e.g., the dataset with the greatest accuracy gain or deviation reduction).
2. Ensure all figures and diagrams are included in the final manuscript submission. Several visual references (e.g., Figure 1A, 1B, etc.) are mentioned in the text,

but the actual images are missing from the document reviewed.

3. Add a brief discussion on computational cost or time complexity, especially since MOPSO-based methods may involve additional overhead compared to standard K-Means.

4. Consider expanding the conclusion to include reflections on potential real-world applications of the proposed method (e.g., biomedical clustering, IoT sensor grouping, or document clustering), which could increase the relevance and appeal of the study.

5. Optional: Include a comparison with other metaheuristic optimization algorithms (e.g., Genetic Algorithm, Differential Evolution) either in the discussion or related works section. Even a brief comparative remark based on the literature or prior studies would strengthen the contribution.

Reviewer B:

Recommendation: Accept Submission

Is the title appropriate and aligned with the journal's scope?

Yes

Does the abstract concisely inform readers and summarize the research?

Yes

Are relevant prior studies adequately contextualized, especially in the background?

Yes

Does the author provide in-depth analysis, particularly in the discussion and results sections?

Yes

Does the author demonstrate sufficient scientific reasoning, argumentation and interpretation?

Yes

Are descriptions and explanations presented clearly and understandably?

Yes

Are figures, tables, and numbers of adequate quality and easy to interpret?

Yes

Are visuals like diagrams and images used appropriately and clearly?

Yes

Does the manuscript present an original contribution?

Yes

Is the writing style suitably academic with clear and proper language?

Yes

Give constructive feedback:

The authors have fulfilled all required revisions; the article can now be accepted and processed for publication.

Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)



A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids

Aina Latifa Riyana Putri^{1*}, Joko Riyono², Christina Eni Pujiastuti³, Supriyadi⁴

¹Data Science, Telkom University, Purwokerto, Indonesia

^{2,3,4} Universitas Trisakti, Jakarta, Indonesia

¹ainaqp@telkomuniversity.ac.id, ²jokoriyono@trisakti.ac.id, ³christina.eni@trisakti.ac.id, ⁴supri@trisakti.ac.id

Abstract

The K-Means algorithm is a popular unsupervised learning method used for data clustering. However, its performance heavily depends on centroid initialization and the distribution shape of the data, making it less effective for datasets with complex or non-linear cluster structures. This study evaluates the performance of the standard K-Means algorithm and proposes a Multiobjective Particle Swarm Optimization K-Means (MOPSO+K-Means) approach to improve clustering accuracy. The evaluation was conducted on five benchmark datasets: Atom, Chainlink, EngyTime, Target, and TwoDiamonds. Experimental results show that K-Means is effective only on datasets with clearly separated clusters, such as EngyTime and TwoDiamonds, achieving accuracies of 95.6% and 100%, respectively. In contrast, MOPSO+K-Means achieved a substantial accuracy improvement on the complex Target dataset, increasing from 0.26% to 59.2%. The TwoDiamonds dataset achieved the most desirable trade-off: it had the lowest SSW (1323.32), relatively high SSB (2863.34), and lowest standard deviation values, indicating compact clusters, good separation, and high consistency across runs. These findings highlight the potential of swarm-based optimization to achieve consistent and accurate clustering results on datasets with varying structural complexity.

Keywords: centroid; k-means; multiobjective particle swarm optimization; the sum of square within; the sum of square between

How to Cite: A. Latifa Riyana Putri, J. Riyono, and C. Eni Pujiastuti, "A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids", *J. RESTI (Rekayasa Sist. Teknol. Inf.)*, vol. 9, no. 3, pp. 542 - 550, Jun. 2025.

Permalink/DOI: <https://doi.org/10.29207/resti.v9i3.6533>

Received: April 11, 2025

Accepted: June 4, 2025

Available Online: June 21, 2025

This is an open-access article under the CC BY 4.0 License

Published by Ikatan Ahli Informatika Indonesia

1. Introduction

Clustering is one of the data mining techniques aimed at grouping data into several clusters based on similarities in their characteristics. One of the most popular clustering methods is k-Means clustering [1], which is a non-hierarchical clustering algorithm that works by initializing centroids randomly. The objects are then grouped into k clusters based on their distances to the centroids, and the centroids' positions are iteratively updated until convergence is reached.

The advantages and disadvantages of this algorithm have been widely discussed in various studies. The k-Means algorithm is well-known for being efficient and scalable in processing large datasets [2]. Previous research, such as that conducted by [3] and [4], has shown that k-Means can produce more compact clusters compared to hierarchical clustering methods. However, k-Means has some limitations, particularly related to

the random initialization of centroids, which can cause the clustering results to vary each time the algorithm is run [5]. Additionally, k-Means tends to get stuck in local optima, leading to suboptimal cluster assignments [6]. Its sensitivity to outliers is also a major concern, as extreme values can significantly shift the centroids' positions [7]. Furthermore, the assumption that clusters are spherical [8] and of uniform size makes k-Means less effective when dealing with datasets that have complex cluster shapes or varying densities.

To understand the quality of clustering results generated by k-Means, it is important to review the objectives and evaluation metrics. In k-Means clustering, the main goal is to form optimal clusters, where members of each cluster are highly like one another but significantly different from members of other clusters [9]. To achieve this, two primary metrics commonly used are the Sum of Squares Within-cluster (SSW) and the Sum of Squares Between-cluster (SSB). SSW measures the

density of the cluster, indicating how tightly the data points within a cluster are grouped around the centroid [10]. Smaller SSW values indicate greater similarity between data points within the same cluster. On the other hand, SSB measures the distance between centroids, reflecting the separation between clusters [11]. Larger SSB values indicate greater distance between clusters, making the clustering more effective at distinguishing between different groups of data.

To address these issues, this study proposes an optimization approach based on Multi-Objective Particle Swarm Optimization (MOPSO) to determine more optimal centroids. MOPSO is a variant of Particle Swarm Optimization (PSO) [12] developed to solve problems with multiple objectives, making it suitable for clustering tasks that involve balancing two criteria simultaneously: minimizing the Sum of Squares Within-cluster (SSW) and maximizing the Sum of Squares Between-cluster (SSB), thereby producing clusters that are balanced in terms of both homogeneity and separation. The main advantage of MOPSO over other optimization algorithms lies in its convergence speed, efficient global exploration capabilities, and ease of implementation. Unlike evolutionary algorithms that use complex selection and mutation processes, MOPSO relies on a simple particle interaction mechanism in the search space. Additionally, MOPSO uses an external archive to store the best non-dominated solutions (Pareto optimal), facilitating decision-making based on trade-offs between cluster homogeneity and separation.

This approach will be evaluated using several benchmark datasets commonly used in clustering studies [13], namely Atom, Chainlink, Engytime, Target, and Two Diamonds. Unlike previous studies that often use classic datasets such as Iris, this study specifically selects these five datasets because they present more complex challenges in the data clustering process. The Atom dataset challenges the algorithm in separating very close clusters, while Chainlink has a topological structure that is interconnected, which is difficult for centroid-based methods to handle. Engytime has an uneven density distribution, which may complicate the identification of proper cluster boundaries. The Target dataset presents a non-linear pattern that is hard for standard k-Means to capture, while Two Diamonds involves clusters that are very close together, making optimal separation difficult.

The performance evaluation was conducted by comparing the clustering results against the ground truth, which refers to the true labels that are known beforehand and used as a reference to assess the accuracy of the clustering performed by the algorithm. This comparison will be made using accuracy metrics and by testing the MOPSO-based approach against conventional clustering methods such as standard k-Means. With MOPSO-based optimization, this method is expected to be able to produce more separated and uniform clusters, thus outperforming data complexity. The findings of this study are anticipated to contribute

significantly to the development of more optimal clustering methods for various applications in data mining and machine learning.

2. Methods

The methodology section will sequentially present the analytical methods employed in this study.

2.1 Experimental Testing and Simulation of the Proposed Multi-Objective Particle Swarm Optimization (PSO) Algorithm

As part of the empirical analysis, this study also evaluates the effectiveness of the proposed method using a set of benchmark datasets. These datasets are employed in the testing process to assess the capability of MOPSO in determining the optimal centroids for the K-Means algorithm, which is a critical step in enhancing clustering quality across various data scenarios. The first step in this study is the selection of datasets to be used for testing the effectiveness of the proposed method. The datasets used in this research are Atom, Chainlink, Engytime, Target, and Two Diamonds.

The Atom dataset [14] is a commonly used benchmark for evaluating clustering algorithms under complex spatial conditions. It exists in a three-dimensional space ($\mathbb{R}^3$) and consists of two main clusters: a dense core cluster with 100 data points located at the center, and a larger, more dispersed outer hull cluster with 400 data points that geometrically encloses the core. This structure creates what is known as an overlapping convex hull, where the outer cluster fully surrounds the inner one. The significant difference in density and spatial arrangement poses challenges for traditional clustering algorithms, especially those that rely on distance measures such as K-Means.

The Chainlink dataset [15], [16] is a benchmark designed to evaluate the ability of clustering algorithms to handle complex and interrelated data structures. It consists of two clusters, each containing 300 data points, forming an interlocked chain-like structure in three-dimensional space ($\mathbb{R}^3$). Each cluster is shaped like a ring, and the two rings are intertwined, creating a configuration known as linear nonseparable entanglement. This refers to a condition where the clusters cannot be linearly separated due to their intertwined spatial arrangement. Although the clusters are globally distinct, many points from one cluster are locally closer to points from the other, which introduces a conflict between global separability and local proximity. Additionally, both clusters have nearly identical densities and inter-point distances, making it difficult to distinguish them based solely on size or distribution. This makes Chainlink particularly challenging for distance-based algorithms such as K-Means.

The EngyTime dataset is a benchmark used to evaluate the effectiveness of clustering algorithms in handling

overlapping clusters with varying densities [17]. It contains 2,000 data points grouped into two clusters in a two-dimensional space ($\mathbb{R}^2$), based on two variables: “Engy” and “Time”. This dataset represents a simplified yet realistic density-based clustering problem, like those encountered in applications such as flow cytometry and sonar signal processing. EngyTime is generated from a mixture of two 2D Gaussian distributions, making it a suitable test case for evaluating how well clustering algorithms can distinguish overlapping groups. The clusters in this dataset are not separated by empty space, and they differ in density, which creates a significant challenge for traditional centroid-based methods like K-Means. These algorithms rely primarily on distance metrics and often ignore local density, which may lead to incorrect groupings when faced with overlapping, unevenly distributed data.

The Target dataset [18] is a benchmark designed to test the robustness of clustering algorithms when faced with overlapping clusters and the presence of outliers. This dataset exists in two-dimensional space ($\mathbb{R}^2$) and comprises 743 data points, divided into two main clusters and four small outlier groups. The first cluster is a dense spherical structure with 365 data points, while the second cluster forms an enclosing ring with 395 data points. This circular arrangement results in overlapping convex hulls, a geometric configuration that is particularly difficult to resolve using centroid-based algorithms like K-Means, which assume well-separated, linearly distinguishable clusters. The dataset also includes four corner-located outlier groups, each containing four data points. These outliers introduce additional complexity by potentially skewing the centroid calculations or being misclassified as independent clusters. The combination of dense core–ring overlap and peripheral noise makes the Target dataset a comprehensive benchmark for evaluating both the accuracy and stability of clustering methods.

The TwoDiamonds dataset [19], [20] is a benchmark commonly used to evaluate the ability of clustering algorithms to distinguish between weakly connected yet distinct cluster structures. It consists of 400 data points distributed evenly across two clusters, each shaped like a diamond, and located in a two-dimensional space ($\mathbb{R}^2$). The two clusters are positioned in adjacent square regions that almost touch at one side, forming a configuration that resembles two diamonds placed side by side. The primary challenge of this dataset lies in the weak connection between the two clusters. While the clusters are globally distinct, the narrow gap separating them can mislead clustering algorithms—particularly those based on local distance metrics like K-Means—into interpreting them as a single elongated cluster. Successfully separating the clusters in this dataset requires the algorithm to capture the overall geometric structure rather than relying purely on inter-point distances.

2.2 Standard K-Means Implementation as a baseline Comparison

As a baseline for performance comparison, the next step involves applying the standard K-Means clustering algorithm to each benchmark dataset. K-Means begins by randomly initializing centroids, followed by an iterative process of assigning data points to the nearest centroid based on Euclidean distance and updating the centroids until convergence. The number of clusters (k) is predetermined based on the ground truth of each dataset.

Due to the random nature of centroid initialization, K-Means may produce different clustering results in different runs. To address this, multiple independent runs are performed to assess consistency. The resulting cluster assignments are then evaluated using a confusion matrix, which enables the calculation of clustering accuracy by comparing the predicted clusters to the true class labels.

This baseline evaluation provides a reference point for comparing the clustering performance of the proposed MOPSO-KMeans method. By analyzing both standard K-Means and the optimized version under the same conditions and metrics, a more comprehensive assessment of the benefits and improvements introduced by the proposed approach can be achieved.

3. Results and Discussions

3.1 Design and Workflow of the Proposed MOPSO Method

To improve the quality of clustering results, this study implements the Multi-Objective Particle Swarm Optimization (MOPSO) algorithm to optimize the selection of centroids in the K-Means algorithm. This approach simultaneously considers two objectives: minimizing the Sum of Squared Within-Cluster (SSW) and maximizing the Sum of Squared Between-Cluster (SSB).

The first objective function aims to minimize the Sum of Squared Within-Cluster (SSW) shown in Equation 1.

$$f_1 = \min \left(\sum_{j=1}^k \sum_{x_i \in C_j} \|x_i - \mu_j\|^2 \right) \quad (1)$$

k is the number of clusters; x_i is the i -th data point; μ_j is the centroid of cluster C_j ; $\|x_i - \mu_j\|^2$ is the squared Euclidean distance between the data point and the cluster centroid.

The second objective function shown in Equation 2 aims to maximize the Sum of Squared Between-Cluster (SSB), which is expressed as the minimization of its negative:

$$f_2 = -\min \left(-\sum_{j=1}^k n_j \|\mu_j - \mu\|^2 \right) \quad (2)$$

n_j is the number of data points in cluster j ; μ is the global centroid of the entire dataset; $\|\mu_j - \mu\|^2$ is the squared

distance between the cluster centroid and the global centroid.

The complete procedure of the proposed Multi-Objective Particle Swarm Optimization (PSO) algorithm can be summarized as follows, with its pseudocode presented on Figure 1.

Algorithm 1: Pseudocode of the proposed MOPSO+K-Means Algorithm	
Input:	Dataset D with n data points; Number of particles N; Maximum number of iterations T; Objective functions: $f1$ = intra-cluster distance (minimize) and $f2$ = inter-cluster separation (maximize)
Process:	<pre> 1. Initialization Phase: 2. For i = 1 to N Do: 3. Randomly initialize centroids_i with K positions in the data space. 4. Initialize velocity_i = 0. 5. Assign each data point in D to the nearest centroid in centroids_i. 6. Evaluate objective functions f1 and f2. 7. Set best_position_i = centroids_i. 8. Set best_objective_i = [f1, f2]. 9. End For 10. Evaluate dominance among all particles. 11. Save non-dominated solutions into repository. 12. Search Phase: 13. For iteration = 1 to T Do: 14. For i = 1 to N Do: 15. Select global_best from repository using crowding distance. 16. Update velocity_i using PSO velocity update rule. 17. Update centroids_i using velocity_i. 18. Assign data points to the nearest centroid in centroids_i. 19. Evaluate objective functions f1 and f2. 20. Apply mutation (optional). 21. If centroids_i dominates best_position_i Then: 22. best_position_i = centroids_i. 23. End If 24. End For 25. Evaluate dominance among all particles. 26. Update repository with new non-dominated solutions. 27. Remove dominated solutions from repository. 28. Update inertia weight w (if used). 29. End For 30. Return repository as the Pareto front of optimal centroid configurations. </pre>
Output:	Repository of non-dominated centroid sets (Pareto-optimal solutions)

Figure 1. Pseudocode of the Proposed MOPSO+K-Means Algorithm

The proposed MOPSO+K-Means algorithm integrates Multi-Objective Particle Swarm Optimization (MOPSO) with the traditional K-Means clustering method to overcome the limitations of random centroid initialization. As outlined in Figure 1, the process begins with the initialization phase, where each particle represents a potential solution in the form of a set of cluster centroids. Objective functions are defined to minimize intra-cluster distance (SSW) and maximize inter-cluster separation (SSB), enabling a balanced evaluation of cluster compactness and separation.

The fitness of each particle is assessed based on these two objectives, and the non-dominated solutions are stored in a Pareto-based repository. In the search phase, the particle positions are updated using both personal and global bests selected from the repository using the crowding distance. This iterative process continues until the stopping criteria are met, gradually refining the solutions toward the optimal trade-offs between the two objectives.

At the end of the MOPSO optimization, the repository contains a set of Pareto-optimal centroid configurations. From these, one or more candidate solutions can be selected to initialize the K-Means algorithm. This hybrid approach allows K-Means to begin with well-optimized centroid positions, potentially resulting in more stable and accurate clustering outcomes compared to traditional random initialization.

In addition to the Particle Swarm Optimization (PSO) algorithm, previous studies have also explored other metaheuristic approaches to improve clustering quality. Among the most widely used are Genetic Algorithm [21] (GA) and Differential Evolution [22] (DE). However, findings from several prior works suggest that PSO tends to offer advantages in terms of convergence speed and simplicity of parameters. PSO only requires a few parameters to be configured, such as inertia weight and learning coefficients. In contrast, GA and DE involve more complex parameter settings, including crossover rate, mutation rate, and selection strategies [23] which can significantly influence performance if not properly adjusted. Furthermore, PSO is known for its ability to maintain a good balance between exploration and exploitation during the search process [24], making it especially suitable for clustering problems with high complexity.

3.2 Experimental Results and Analysis

To evaluate the performance of the proposed Multi-Objective PSO, a series of experiments were conducted on a set of well-known benchmark datasets. These five datasets have been described in the Research Method.

The experimental setup in this study is defined as follows: the swarm size is set to $N=40$, and each benchmark dataset is tested independently 30 times, with each execution consisting of 100 iterations. All PSO-based algorithms are terminated upon reaching this maximum number of iterations. The performance of the proposed MOPSO+K-Means method is evaluated using standard clustering metrics, namely the best value, average value, and standard deviation of SSW, SSB, and accuracy. These metrics are used to assess the effectiveness and stability of the proposed method in comparison to standard K-Means across various benchmark datasets.

The performance of MOPSO+K-Means was evaluated using commonly used optimization metrics, namely the average solution and standard deviation. These metrics were used to assess the effectiveness of MOPSO+K-Means in solving the benchmark clustering tasks.

Table 1. Clustering With MOPSO+K-Means

Dataset	Item	SSW	SSB
Atom	Avg.	1191.22	1414.52
	Std.	230.85	238.22
ChainLink	Avg.	1531.74	1711.26
	Std.	167.04	168.519
EngyTime	Avg.	49122.71	85124.33
	Std.	2951.153	3913.554
Target	Avg.	4126.614	7658.074
	Std.	892.147	1006.303
TwoDiamonds	Avg.	1323.322	2863.340
	Std.	42.822	44.191

Although the numerical results demonstrate that the proposed MOPSO+K-Means algorithm performs competitively across various datasets, a deeper analysis provides further insights into its behavior and performance dynamics. From a clustering quality perspective on Table 1, the ideal objective is to

minimize the Sum of Squared Within-cluster distances (SSW) while maximizing the Sum of Squared Between-cluster distances (SSB). In this context, the EngyTime dataset exhibits the highest SSW (49122.71) and SSB (85124.33), indicating that the dataset likely has a large or widely spread structure. Despite this complexity, the algorithm successfully maintains strong inter-cluster separation. In contrast, the TwoDiamonds dataset achieves the most desirable trade-off: it has the lowest SSW (1323.32), relatively high SSB (2863.34), and the lowest standard deviation values, indicating compact clusters, good separation, and high consistency across runs.

For other datasets, such as Target and Atom, the SSW and SSB values fall in a mid-range category, but the relatively high standard deviations, particularly in Atom, highlight variability in the clustering results, suggesting sensitivity to initial conditions or swarm dynamics. Similarly, ChainLink yields results comparable to Atom but also shows considerable variability (Std SSW: 167.04, Std SSB: 168.52), likely reflecting the complexity or overlap in its cluster structure. These findings suggest that while MOPSO+K-Means can handle both simple and

complex datasets, its stability may be challenged under certain data conditions.

Overall, TwoDiamonds emerges as the most stable dataset for this method, whereas EngyTime, despite strong average performance, reveals higher variability possibly due to noise or dispersed data points. These behavioral differences point to both strengths and limitations of the method. The ability of MOPSO+K-Means to find competitive clustering solutions is evident, but further enhancements could improve its robustness. Future research should explore adaptive parameter strategies, improved initialization techniques, or hybrid models to increase the method's reliability across diverse datasets. Additionally, expanding testing to high-dimensional or real-world datasets would help evaluate its scalability and broader applicability.

In this section, an analysis is conducted on the clustering outcomes derived from the implementation of the standard K-Means algorithm and the MOPSO-enhanced K-Means across a range of benchmark datasets. These results highlight the capabilities and limitations of both approaches when confronted with datasets exhibiting varying structural complexities.

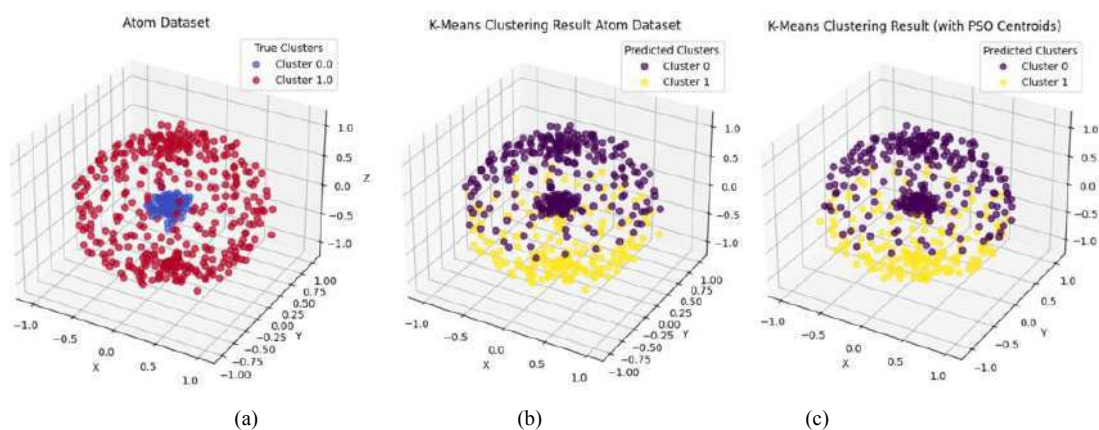


Figure 2. Clustering Result (a) Atom Dataset, (b) K-Means for Atom Dataset, (c) MOPSO+K-Means for Atom Dataset

The clustering results on the Atom dataset highlight the limitations of The Atom dataset illustrates the challenge of clustering data with a concentric structure, comprising a dense core surrounded by a shell. As depicted in Figure 2(a), the ground truth clearly reflects this core-shell configuration. However, the standard K-Means algorithm on Figure 2(b) produces a vertical partition, disregarding the radial nature of the data. This misalignment stems from K-Means assumption of spherical and convex clusters, which proves inadequate for capturing non-linear distributions. The integration of MOPSO+K-Means, as shown in Figure 2(c), results in a more accurate division that successfully distinguishes between the core and the surrounding shell. This demonstrates the ability of MOPSO to guide K-Means toward more structure-aware clustering outcomes in complex spatial configurations.

A similar pattern is observed in the ChainLink dataset, which features two intertwined, non-convex clusters on Figure 3(a). The standard K-Means algorithm on Figure 3(b) again fails to separate the data meaningfully, as it assigns points based on straight-line distances, ignoring the dataset's intricate shape. Conversely, the MOPSO+K-means on Figure 3(c) more effectively untangles the two chains, maintaining their topological distinction. This improved result underscores the role of MOPSO in adapting centroid placement to fit non-linear geometries that would otherwise confound traditional methods.

The EngyTime dataset, in contrast, presents a linearly separable structure, offering a more favorable scenario for K-Means. In Figure 4(a), the data exhibits two well-defined, adjacent clusters. The standard K-Means output on Figure 4(b) broadly captures the separation but misclassifies several points near the decision

boundary. With optimized centroids on Figure 4(c), the clustering becomes cleaner, with reduced boundary ambiguity. This shows that even in simpler datasets,

MOPSO+K-Means contributes by refining the clustering precision and minimizing convergence to suboptimal solutions.

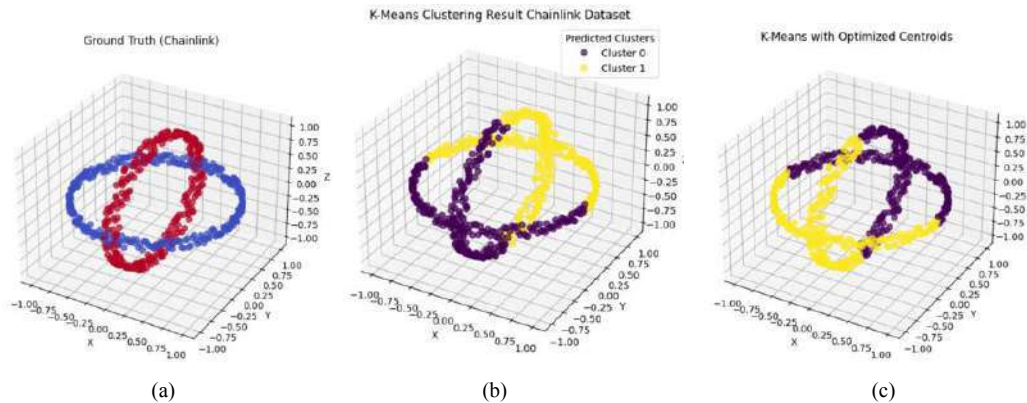


Figure 3. Clustering Result (a) Chainlink Dataset, (b) K-Means for Chainlink Dataset, (c) MOPSO+K-Means for Chainlink Dataset

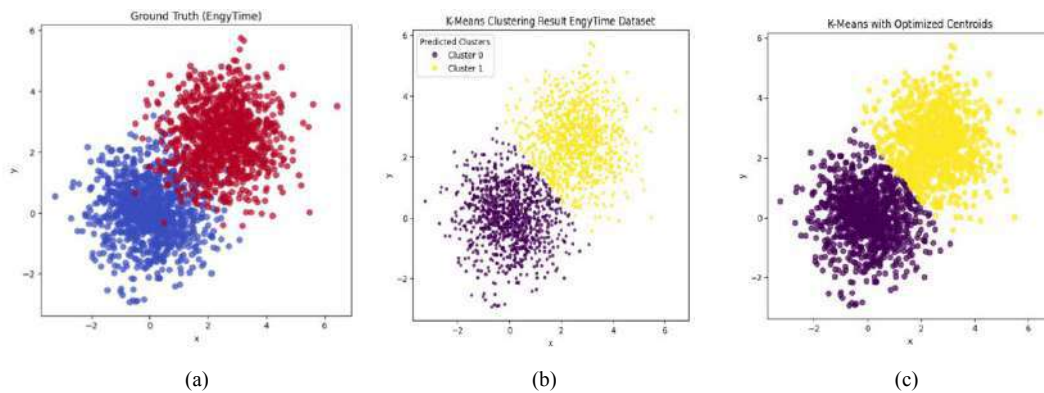


Figure 4. Clustering Result (a) EngyTime Dataset, (b) K-Means for EngyTime Dataset, (c) MOPSO+K-Means for EngyTime Dataset

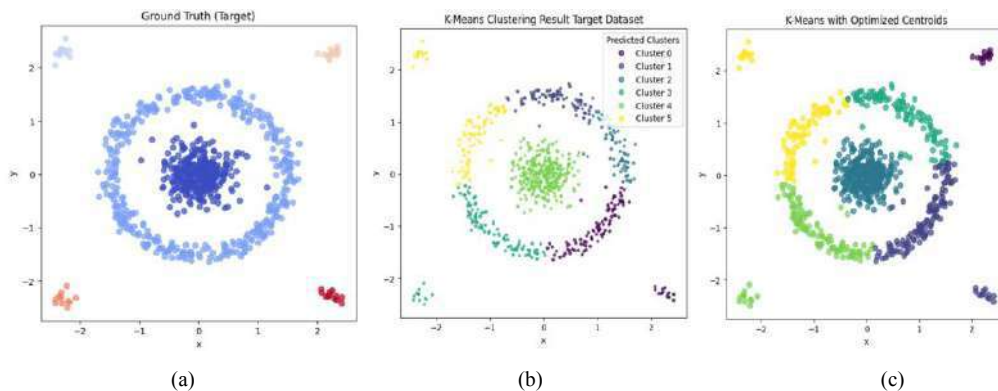


Figure 5. Clustering Results (a) Target Dataset, (b) K-Means for Target Dataset, (c) MOPSO+K-Means for Target Dataset

The analysis of the Target dataset provides further insight into the impact of complex geometries on clustering performance. This dataset contains concentric circular clusters and scattered peripheral groups on Figure 5(a). The clustering produced by the standard K-Means algorithm on Figure 5(b) results in fragmented groupings that poorly reflect the actual layout. Its tendency to impose spherical boundaries leads to significant structural mismatches. The MOPSO+K-Means on Figure 5(c) successfully aligns with the circular patterns and isolates the outlying groups more accurately, reinforcing the value of

optimization in adapting to irregular spatial distributions.

Lastly, the TwoDiamonds dataset poses a moderate challenge due to its diamond-shaped, linearly separated clusters on Figure 6(a). While the standard K-Means on Figure 6(b) performs reasonably well, some inconsistencies are evident along the boundary, suggesting less-than-optimal centroid positioning. By contrast, the optimized version on Figure 6(c) yields a more symmetric and faithful clustering result, demonstrating how MOPSO can enhance K-Means even in datasets that are linearly separable but geometrically unconventional.

Table 2 presents the performance evaluation results of the MOPSO-K-Means algorithm on five benchmark datasets: Atom, ChainLink, EngyTime, Target, and TwoDiamonds. The evaluation was carried out using commonly used optimization metrics, namely the average solution and standard deviation of the SSW

(Sum of Squares Within) and SSB (Sum of Squares Between), along with the best accuracy achieved for each dataset. The objective of this evaluation is to assess the effectiveness of the MOPSO-K-Means algorithm in producing optimal cluster partitions.

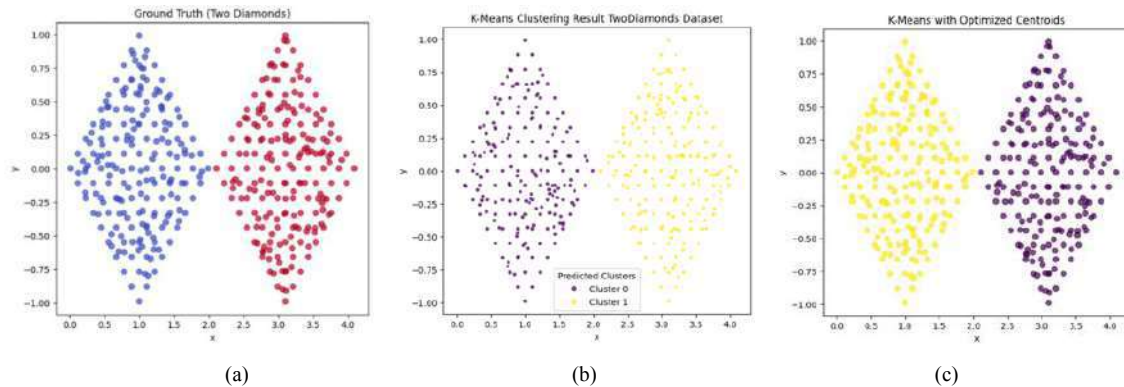


Figure 6. Clustering Result (a) TwoDiamonds Dataset, (b) K-Means Clustering Result TwoDiamonds Dataset, (c) MOPSO+K-Means Clustering Result TwoDiamonds Dataset

Table 2. Comparing Accuracy and Time Computational

Dataset	Accuracy K-means	Best Accuracy MOPSO+K-Means	Average Time Computational MOPSO+K-Means
Atom	54.4%	52.8%	5.1137 seconds
ChainLink	50%	50.2%	6.6889 seconds
EngyTime	95.6%	95.7%	7.617 seconds
Target	0.2692%	59.2%	8.539 seconds
TwoDiamonds	100%	100%	0.844 seconds

Compared to the conventional K-Means algorithm, the results indicate that MOPSO-K-Means generally performs better on most datasets. On the Atom dataset, K-Means achieved an accuracy of 54.4%, while MOPSO-K-Means recorded an accuracy of 52.8%. Although there was a slight decrease, the SSW and SSB values obtained by MOPSO-K-Means still reflect a good and stable cluster distribution, with relatively low standard deviations. For the ChainLink dataset, K-Means achieved 50% accuracy, while MOPSO-K-Means achieved 50.2%, suggesting a slightly better performance in separating the clusters.

Next, on the EngyTime dataset, K-Means reached an accuracy of 95.60%, while MOPSO-K-Means achieved 95.7%. The difference is very small, indicating that both algorithms are equally effective in clustering data with clear cluster structures. However, the most significant improvement was observed on the Target dataset. K-Means achieved only 0.26% accuracy, while MOPSO-K-Means improved the accuracy to 59.2%. This demonstrates that MOPSO-K-Means is more capable of handling datasets with complex or non-linearly separable cluster structures. Lastly, on the TwoDiamonds dataset, both K-Means and MOPSO-K-Means achieved perfect accuracy (100%), indicating that this dataset has a very clear structure that can be easily separated by both algorithms.

The Average Time Computational *MOPSO+K-Means* column in Table 2 presents the average computational

time required by the algorithm to complete one clustering process for each dataset. This time is measured from the beginning of the centroid optimization using the MOPSO algorithm to the final clustering result produced by K-Means. The values represent the average time taken across 30 independent trials. They indicate the efficiency of the algorithm in solving clustering tasks, which is influenced by the complexity of the data patterns and the algorithm's ability to converge toward optimal solutions.

For example, the TwoDiamonds dataset demonstrates the lowest average computational time (0.844 seconds), which can be attributed to its well-separated and clearly structured clusters. This allows MOPSO to converge quickly, requiring minimal exploration. In contrast, the Target dataset shows the highest computational time (8.539 seconds), suggesting that the algorithm needed significantly more effort to explore the solution space due to the highly irregular and overlapping nature of the data clusters. Similarly, EngyTime and Chainlink also required more computational time, which reflects the greater complexity in their data distributions and the increased difficulty in identifying distinct clusters. Meanwhile, Atom shows a moderate computational time, indicating that while the data is not entirely straightforward, it is still manageable for the algorithm to optimize efficiently.

Overall, the evaluation results show that MOPSO+K-Means has advantages in terms of flexibility and effectiveness in identifying complex cluster structures that conventional K-Means struggles to handle. The relatively small standard deviations across most datasets also indicate that this algorithm can produce stable and consistent solutions in each optimization run. Therefore, MOPSO+K-Means can be considered a more reliable alternative for clustering tasks involving datasets with diverse characteristics.

However, one key limitation of the current study is the assumption that the number of clusters (k) is known beforehand. Although this facilitates benchmarking against ground truth labels, it does not reflect the realities of unsupervised learning tasks, where k must be inferred from the data. In practice, k is a hyperparameter that must be estimated carefully. Common strategies include the elbow method, which identifies diminishing returns in intra-cluster variance as k increases, the silhouette score, which quantifies cluster cohesion and separation, and the gap statistic, which compares clustering performance against that of random reference distributions.

To overcome this limitation, future work could explore the extension of MOPSO to simultaneously optimize both the number of clusters and the centroid positions. This could be framed as a multi-objective optimization problem—balancing cluster compactness, separation, and model complexity—or as a constrained optimization task where k is bounded within a reasonable range. Such an approach would enhance the applicability of the method in real-world scenarios, allowing it to autonomously discover both the optimal clustering structure and its corresponding parameters without relying on prior knowledge.

4. Conclusions

Based on the analysis and evaluation of five benchmark datasets, it can be concluded that the performance of the K-Means algorithm is highly dependent on the shape and structural characteristics of the clusters in the data. On datasets with simple and linearly separable structures, such as TwoDiamonds and EngyTime, K-Means performs very well, achieving high accuracy—up to 100%. However, on datasets with non-linear or complex structures, such as Atom, ChainLink, and Target, the algorithm fails to properly separate clusters, resulting in low accuracy and poor alignment with the ground truth.

To address these limitations, the MOPSO+K-Means approach was introduced as an alternative solution. Based on the experimental results, this algorithm shows significant performance improvement on datasets with complex structures—most notably on the Target dataset, where the accuracy increased from 26% (K-Means) to 59.2% (MOPSO+K-Means). In addition, the obtained SSW and SSB values, along with relatively low standard deviations, indicate that MOPSO-K-Means can produce stable and consistent clustering solutions.

Overall, MOPSO+K-Means has proven to be more flexible and reliable in handling various types of cluster structures, making it a more suitable choice for clustering tasks involving non-convex or non-linearly separable data distributions. Such characteristics are often found in real-world applications, including biomedical data analysis where patient subgroups may form irregular patterns, sensor grouping in IoT

environments with overlapping signal zones, and document clustering tasks where semantic relationships are not linearly separable. These application domains can benefit from algorithms that offer both structural flexibility and solution stability, as demonstrated by MOPSO+K-Means in this study.

Acknowledgements

This work was supported by Telkom University.

References

- [1] T. M. Ghazal *et al.*, “Performances of K-Means Clustering Algorithm with Different Distance Metrics,” *Intelligent Automation & Soft Computing*, vol. 30, no. 2, pp. 735–742, Aug. 2021, doi: 10.32604/IASC.2021.019067.
- [2] M. Ahmed, R. Seraj, and S. M. S. Islam, “The k-means Algorithm: A Comprehensive Survey and Performance Evaluation,” *Electronics* 2020, Vol. 9, Page 1295, vol. 9, no. 8, p. 1295, Aug. 2020, doi: 10.3390/ELECTRONICS9081295.
- [3] A. Abdulhafedh, “Incorporating K-means, Hierarchical Clustering and PCA in Customer Segmentation,” *Journal of City and Development*, vol. 3, no. 1, 2021.
- [4] P. Patel, B. Sivaiah, and R. Patel, “Approaches for finding Optimal Number of Clusters using K-Means and Agglomerative Hierarchical Clustering Techniques,” in *2022 International Conference on Intelligent Controller and Computing for Smart Power, ICICCSPP 2022*, 2022, doi: 10.1109/ICICCSPP53532.2022.9862439.
- [5] A. Vouros, S. Langdell, M. Croucher, and E. Vasilaki, “An empirical comparison between stochastic and deterministic centroid initialisation for K-means variations,” *Mach Learn*, vol. 110, no. 8, pp. 1975–2003, Aug. 2021, doi: 10.1007/S10994-021-06021-7/FIGURES/8.
- [6] A. Qtaish, M. Braik, D. Albashish, M. T. Alshammari, A. Alreshidi, and E. J. Alreshidi, “Optimization of K-means clustering method using hybrid capuchin search algorithm,” *Journal of Supercomputing*, vol. 80, no. 2, 2024, doi: 10.1007/s11227-023-05540-5.
- [7] Z. Zhang, Q. Feng, J. Huang, Y. Guo, J. Xu, and J. Wang, “A local search algorithm for k-means with outliers,” *Neurocomputing*, vol. 450, pp. 230–241, Aug. 2021, doi: 10.1016/J.NEUCOM.2021.04.028.
- [8] A. A. Wani, “Comprehensive analysis of clustering algorithms: exploring limitations and innovative solutions,” *PeerJ Comput Sci*, vol. 10, pp. 1–45, Aug. 2024, doi: 10.7717/PEERJ-CS.2286/FIG-14.
- [9] R. Gustriansyah, N. Suhandi, and F. Antony, “Clustering optimization in RFM analysis Based on k-Means,” *IJECS*, vol. 18, no. 1, pp. 470–477, Apr. 2020, doi: 10.11591/ijeecs.v18.i1.pp470-477.
- [10] R. Richard, H. Cao, and M. Wachowicz, “An Automated Clustering Process for Helping Practitioners to Identify Similar EV Charging Patterns across Multiple Temporal Granularities,” *International Conference on Smart Cities and Green ICT Systems*, pp. 67–77, 2021, doi: 10.5220/0010485000670077.
- [11] M. T. Guerreiro *et al.*, “Anomaly Detection in Automotive Industry Using Clustering Methods—A Case Study,” *Applied Sciences* 2021, Vol. 11, Page 9868, vol. 11, no. 21, p. 9868, Oct. 2021, doi: 10.3390/AP11219868.
- [12] M. Jain, V. Saihpal, N. Singh, and S. B. Singh, “An Overview of Variants and Advancements of PSO Algorithm,” *MDPI Applied Sciences*, 2022, doi: 10.3390/app12178392.
- [13] M. C. Thrun and A. Ultsch, “Clustering benchmark datasets exploiting the fundamental clustering problems,” *Data Brief*, vol. 30, p. 105501, Jun. 2020, doi: 10.1016/J.DIB.2020.105501.
- [14] A. Ultsch, “Strategies for an Artificial Life System to cluster high dimensional Data,” 2004, Accessed: Apr. 14,

2025. [Online]. Available: <https://www.researchgate.net/publication/228932819>
- [15] A. Ultsch, G. Guimaraes, D. Korus, and H. Li, "Knowledge Extraction from Artificial Neural Networks and Applications," *Parallele Datenverarbeitung mit dem Transputer*, pp. 148–162, 1994, doi: 10.1007/978-3-642-78901-4_11.
- [16] P. Mangiameli, S. K. Chen, and D. West, "A comparison of SOM neural network and hierarchical clustering methods," *Eur J Oper Res*, vol. 93, no. 2, pp. 402–417, Sep. 1996, doi: 10.1016/0377-2217(96)00038-0.
- [17] S. P. Chatzis and D. I. Kosmopoulos, "A variational Bayesian methodology for hidden Markov models utilizing Student's-t mixtures," *Pattern Recognit*, vol. 44, no. 2, pp. 295–306, Feb. 2011, doi: 10.1016/J.PATCOG.2010.09.001.
- [18] J. Poelmans, M. M. Van Hulle, S. Viaene, P. Elzinga, and G. Dedene, "Text mining with emergent self organizing maps and multi-dimensional scaling: A comparative study on domestic violence," *Appl Soft Comput*, vol. 11, no. 4, pp. 3870–3876, Jun. 2011, doi: 10.1016/J.ASOC.2011.02.026.
- [19] A. Ultsch, "U*-Matrix : a Tool to visualize Clusters in high dimensional Data," 2004.
- [20] A. Ultsch, "Density Estimation and Visualization for Data Containing Clusters of Unknown Structure," *Studies in Classification, Data Analysis, and Knowledge Organization*, pp. 232–239, 2005, doi: 10.1007/3-540-28084-7_25.
- [21] B. Khusul Khotimah, F. Irhamni, and T. Sundarwati, "A GENETIC ALGORITHM FOR OPTIMIZED INITIAL CENTERS K-MEANS CLUSTERING IN SMEs," *J Theor Appl Inf Technol*, vol. 15, no. 1, 2016, Accessed: Jun. 13, 2025. [Online]. Available: www.jatit.org
- [22] H. He, B. Sun, Y. Yang, and S. Liu, "An improved K-Means clustering based on differential evolution," *J Phys Conf Ser*, vol. 2595, no. 1, p. 012010, Sep. 2023, doi: 10.1088/1742-6596/2595/1/012010.
- [23] A. Hassanat, K. Almohammadi, E. Alkafaween, E. Abunawas, A. Hammouri, and V. B. S. Prasath, "Choosing Mutation and Crossover Ratios for Genetic Algorithms—A Review with a New Dynamic Approach," *Information 2019, Vol. 10, Page 390*, vol. 10, no. 12, p. 390, Dec. 2019, doi: 10.3390/INFO10120390.
- [24] M. Jain, V. Saijpal, N. Singh, and S. B. Singh, "An Overview of Variants and Advancements of PSO Algorithm," *Applied Sciences 2022, Vol. 12, Page 8392*, vol. 12, no. 17, p. 8392, Aug. 2022, doi: 10.3390/AP12178392.

A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids

by Joko Riyanto

Submission date: 06-Nov-2025 02:51PM (UTC+0700)

Submission ID: 2805088658

File name: 3._RESTI.pdf (474.4K)

Word count: 5719

Character count: 34564



A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids

Aina Latifa Riyana Putri^{1*}, Joko Riyono², Christina Eni Pujiastuti³, Supriyadi⁴

¹Data Science, Telkom University, Purwokerto, Indonesia

^{2,3,4} Universitas Trisakti, Jakarta, Indonesia

¹ainaqp@telkomuniversity.ac.id, ²jokoriyono@trisakti.ac.id, ³christina.eni@trisakti.ac.id, ⁴supri@trisakti.ac.id

Abstract

The K-Means algorithm is a popular unsupervised learning method used for data clustering. However, its performance heavily depends on centroid initialization and the distribution shape of the data, making it less effective for datasets with complex or non-linear cluster structures. This study evaluates the performance of the standard K-Means algorithm and proposes a Multiobjective Particle Swarm Optimization K-Means (MOPSO+K-Means) approach to improve clustering accuracy. The evaluation was conducted on five benchmark datasets: Atom, Chainlink, EngyTime, Target, and TwoDiamonds. Experimental results show that K-Means is effective only on datasets with clearly separated clusters, such as EngyTime and TwoDiamonds, achieving accuracies of 95.6% and 100%, respectively. In contrast, MOPSO+K-Means achieved a substantial accuracy improvement on the complex Target dataset, increasing from 0.26% to 59.2%. The TwoDiamonds dataset achieved the most desirable trade-off: it had the lowest SSW (1323.32), relatively high SSB (2863.34), and lowest standard deviation values, indicating compact clusters, good separation, and high consistency across runs. These findings highlight the potential of swarm-based optimization to achieve consistent and accurate clustering results on datasets with varying structural complexity.

Keywords: centroid; k-means; multiobjective particle swarm optimization; the sum of square within; the sum of square between

How to Cite: A. Latifa Riyana Putri, J. Riyono, and C. Eni Pujiastuti, "A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids", *J. RESTI (Rekayasa Sist. Teknol. Inf.)*, vol. 9, no. 3, pp. 542 - 550, Jun. 2025.
Permalink/DOI: <https://doi.org/10.29207/resti.v9i3.6533>

Received: April 11, 2025

Accepted: June 4, 2025

Available Online: June 21, 2025



This is an open-access article under the CC BY 4.0 License
Published by Ikatan Ahli Informatika Indonesia

1. Introduction

Clustering is one of the data mining techniques aimed at grouping data into several clusters based on similarities in their characteristics. One of the most popular clustering methods is k-Means clustering [1], which is a non-hierarchical clustering algorithm that works by initializing centroids randomly. The objects are then grouped into k clusters based on their distances to the centroids, and the centroids' positions are iteratively updated until convergence is reached.

The advantages and disadvantages of this algorithm have been widely discussed in various studies. The k-Means algorithm is well-known for being efficient and scalable in processing large datasets [2]. Previous research, such as that conducted by [3] and [4], has shown that k-Means can produce more compact clusters compared to hierarchical clustering methods. However, k-Means has some limitations, particularly related to

the random initialization of centroids, which can cause the clustering results to vary each time the algorithm is run [5]. Additionally, k-Means tends to get stuck in local optima, leading to suboptimal cluster assignments [6]. Its sensitivity to outliers is also a major concern, as extreme values can significantly shift the centroids' positions [7]. Furthermore, the assumption that clusters are spherical [8] and of uniform size makes k-Means less effective when dealing with datasets that have complex cluster shapes or varying densities.

To understand the quality of clustering results generated by k-Means, it is important to review the objectives and evaluation metrics. In k-Means clustering, the main goal is to form optimal clusters, where members of each cluster are highly like one another but significantly different from members of other clusters [9]. To achieve this, two primary metrics commonly used are the Sum of Squares Within-cluster (SSW) and the Sum of Squares Between-cluster (SSB). SSW measures the

density of the cluster, indicating how tightly the data points within a cluster are grouped around the centroid [10]. Smaller SSW values indicate greater similarity between data points within the same cluster. On the other hand, SSB measures the distance between centroids, reflecting the separation between clusters [11]. Larger SSB values indicate greater distance between clusters, making the clustering more effective at distinguishing between different groups of data.

To address these issues, this study proposes an optimization approach based on Multi-Objective Particle Swarm Optimization (MOPSO) to determine more optimal centroids. MOPSO is a variant of Particle Swarm Optimization (PSO) [12] developed to solve problems with multiple objectives, making it suitable for clustering tasks that involve balancing two criteria simultaneously: minimizing the Sum of Squares Within-cluster (SSW) and maximizing the Sum of Squares Between-cluster (SSB), thereby producing clusters that are balanced in terms of both homogeneity and separation. The main advantage of MOPSO over other optimization algorithms lies in its convergence speed, efficient global exploration capabilities, and ease of implementation. Unlike evolutionary algorithms that use complex selection and mutation processes, MOPSO relies on a simple particle interaction mechanism in the search space. Additionally, MOPSO uses an external archive to store the best non-dominated solutions (Pareto optimal), facilitating decision-making based on trade-offs between cluster homogeneity and separation.

This approach will be evaluated using several benchmark datasets commonly used in clustering studies [13], namely Atom, Chainlink, Engytime, Target, and Two Diamonds. Unlike previous studies that often use classic datasets such as Iris, this study specifically selects these five datasets because they present more complex challenges in the data clustering process. The Atom dataset challenges the algorithm in separating very close clusters, while Chainlink has a topological structure that is interconnected, which is difficult for centroid-based methods to handle. Engytime has an uneven density distribution, which may complicate the identification of proper cluster boundaries. The Target dataset presents a non-linear pattern that is hard for standard k-Means to capture, while Two Diamonds involves clusters that are very close together, making optimal separation difficult.

The performance evaluation was conducted by comparing the clustering results against the ground truth, which refers to the true labels that are known beforehand and used as a reference to assess the accuracy of the clustering performed by the algorithm. This comparison will be made using accuracy metrics and by testing the MOPSO-based approach against conventional clustering methods such as standard k-Means. With MOPSO-based optimization, this method is expected to be able to produce more separated and uniform clusters, thus outperforming data complexity. The findings of this study are anticipated to contribute

significantly to the development of more optimal clustering methods for various applications in data mining and machine learning.

2. Methods

The methodology section will sequentially present the analytical methods employed in this study.

2.1 Experimental Testing and Simulation of the Proposed Multi-Objective Particle Swarm Optimization (PSO) Algorithm

As part of the empirical analysis, this study also evaluates the effectiveness of the proposed method using a set of benchmark datasets. These datasets are employed in the testing process to assess the capability of MOPSO in determining the optimal centroids for the K-Means algorithm, which is a critical step in enhancing clustering quality across various data scenarios. The first step in this study is the selection of datasets to be used for testing the effectiveness of the proposed method. The datasets used in this research are Atom, Chainlink, Engytime, Target, and Two Diamonds.

The Atom dataset [14] is a commonly used benchmark for evaluating clustering algorithms under complex spatial conditions. It exists in a three-dimensional space (R^3) and consists of two main clusters: a dense core cluster with 100 data points located at the center, and a larger, more dispersed outer hull cluster with 400 data points that geometrically encloses the core. This structure creates what is known as an overlapping convex hull, where the outer cluster fully surrounds the inner one. The significant difference in density and spatial arrangement poses challenges for traditional clustering algorithms, especially those that rely on distance measures such as K-Means.

The Chainlink dataset [15], [16] is a benchmark designed to evaluate the ability of clustering algorithms to handle complex and interrelated data structures. It consists of two clusters, each containing 300 data points, forming an interlocked chain-like structure in three-dimensional space (R^3). Each cluster is shaped like a ring, and the two rings are intertwined, creating a configuration known as linear nonseparable entanglement. This refers to a condition where the clusters cannot be linearly separated due to their intertwined spatial arrangement. Although the clusters are globally distinct, many points from one cluster are locally closer to points from the other, which introduces a conflict between global separability and local proximity. Additionally, both clusters have nearly identical densities and inter-point distances, making it difficult to distinguish them based solely on size or distribution. This makes Chainlink particularly challenging for distance-based algorithms such as K-Means.

The EngyTime dataset is a benchmark used to evaluate the effectiveness of clustering algorithms in handling

overlapping clusters with varying densities [17]. It contains 2,000 data points grouped into two clusters in a two-dimensional space ($\mathbb{R}^2$), based on two variables: "Engy" and "Time". This dataset represents a simplified yet realistic density-based clustering problem, like those encountered in applications such as flow cytometry and sonar signal processing. EngyTime is generated from a mixture of two 2D Gaussian distributions, making it a suitable test case for evaluating how well clustering algorithms can distinguish overlapping groups. The clusters in this dataset are not separated by empty space, and they differ in density, which creates a significant challenge for traditional centroid-based methods like K-Means. These algorithms rely primarily on distance metrics and often ignore local density, which may lead to incorrect groupings when faced with overlapping, unevenly distributed data.

The Target dataset [18] is a benchmark designed to test the robustness of clustering algorithms when faced with overlapping clusters and the presence of outliers. This dataset exists in two-dimensional space ($\mathbb{R}^2$) and comprises 743 data points, divided into two main clusters and four small outlier groups. The first cluster is a dense spherical structure with 365 data points, while the second cluster forms an enclosing ring with 395 data points. This circular arrangement results in overlapping convex hulls, a geometric configuration that is particularly difficult to resolve using centroid-based algorithms like K-Means, which assume well-separated, linearly distinguishable clusters. The dataset also includes four corner-located outlier groups, each containing four data points. These outliers introduce additional complexity by potentially skewing the centroid calculations or being misclassified as independent clusters. The combination of dense core-ring overlap and peripheral noise makes the Target dataset a comprehensive benchmark for evaluating both the accuracy and stability of clustering methods.

The TwoDiamonds dataset [19], [20] is a benchmark commonly used to evaluate the ability of clustering algorithms to distinguish between weakly connected yet distinct cluster structures. It consists of 400 data points distributed evenly across two clusters, each shaped like a diamond, and located in a two-dimensional space ($\mathbb{R}^2$). The two clusters are positioned in adjacent square regions that almost touch at one side, forming a configuration that resembles two diamonds placed side by side. The primary challenge of this dataset lies in the weak connection between the two clusters. While the clusters are globally distinct, the narrow gap separating them can mislead clustering algorithms—particularly those based on local distance metrics like K-Means—into interpreting them as a single elongated cluster. Successfully separating the clusters in this dataset requires the algorithm to capture the overall geometric structure rather than relying purely on inter-point distances.

2.2 Standard K-Means Implementation as a baseline Comparison

As a baseline for performance comparison, the next step involves applying the standard K-Means clustering algorithm to each benchmark dataset. K-Means begins by randomly initializing centroids, followed by an iterative process of assigning data points to the nearest centroid based on Euclidean distance and updating the centroids until convergence. The number of clusters (k) is predetermined based on the ground truth of each dataset.

Due to the random nature of centroid initialization, K-Means may produce different clustering results in different runs. To address this, multiple independent runs are performed to assess consistency. The resulting cluster assignments are then evaluated using a confusion matrix, which enables the calculation of clustering accuracy by comparing the predicted clusters to the true class labels.

This baseline evaluation provides a reference point for comparing the clustering performance of the proposed MOPSO-KMeans method. By analyzing both standard K-Means and the optimized version under the same conditions and metrics, a more comprehensive assessment of the benefits and improvements introduced by the proposed approach can be achieved.

3. Results and Discussions

3.1 Design and Workflow of the Proposed MOPSO Method

To improve the quality of clustering results, this study implements the Multi-Objective Particle Swarm Optimization (MOPSO) algorithm to optimize the selection of centroids in the K-Means algorithm. This approach simultaneously considers two objectives: minimizing the Sum of Squared Within-Cluster (SSW) and maximizing the Sum of Squared Between-Cluster (SSB).

The first objective function aims to minimize the Sum of Squared Within-Cluster (SSW) shown in Equation 1.

$$f_1 = \min \left(\sum_{i=1}^k \sum_{x_i \in C_i} \|x_i - \mu_i\|^2 \right) \quad (1)$$

k is the number of clusters; x_i is the i -th data point; μ_j is the centroid of cluster C_j ; $\|x_i - \mu_j\|^2$ is the squared Euclidean distance between the data point and the cluster centroid.

The second objective function shown in Equation 2 aims to maximize the Sum of Squared Between-Cluster (SSB), which is expressed as the minimization of its negative:

$$f_2 = -\min \left(-\sum_{j=1}^k n_j \|\mu_j - \mu\|^2 \right) \quad (2)$$

n_j is the number of data points in cluster j ; μ is the global centroid of the entire dataset; $\|\mu_j - \mu\|^2$ is the squared

distance between the cluster centroid and the global centroid.

The complete procedure of the proposed Multi-Objective Particle Swarm Optimization (PSO) algorithm can be summarized as follows, with its pseudocode presented on Figure 1.

```

Algorithm 1: Pseudocode of the proposed MOPSO+K-Means Algorithm
Input:
Dataset D with N data points; Number of particles N; Maximum number of
iterations T; Objective functions:  $f1$  = Intra-cluster distance
(minimizing) and  $f2$  = Inter-cluster separation (maximizing)
Process:
1. Initialization Phase:
2. For  $i = 1$  to N Do:
3. Randomly initialize centroids_i with K positions in the data space.
4. Initialize velocity_i = 0.
5. Assign each data point in D to the nearest centroid in centroids_i.
6. Evaluate objective functions  $f1$  and  $f2$ .
7. Set best_position_i = centroids_i.
8. Set best_objective_i = [f1, f2].
9. End For
10. Evaluate dominance among all particles.
11. Save non-dominated solutions into repository.
12. Search Phase:
13. For iteration = 1 to T Do:
14. For  $i = 1$  to N Do:
15. Select global best from repository using crowding distance.
16. Update velocity_i using PSO velocity update rule.
17. Update centroids_i using velocity_i.
18. Assign data points to the nearest centroid in centroids_i.
19. Evaluate objective functions  $f1$  and  $f2$ .
20. Apply mutation (optional).
21. If centroids_i dominates best_position_i then:
22. best_position_i = centroids_i.
23. End If
24. End For
25. Evaluate dominance among all particles.
26. Update repository with new non-dominated solutions.
27. Remove dominated solutions from repository.
28. Update inertia weight  $w$  (if used).
29. End For
30. Return repository as the Pareto front of optimal centroid
configurations.
Output:
Repository of non-dominated centroid sets (Pareto-optimal solutions)

```

Figure 1. Pseudocode of the Proposed MOPSO+K-Means Algorithm

The proposed MOPSO+K-Means algorithm integrates Multi-Objective Particle Swarm Optimization (MOPSO) with the traditional K-Means clustering method to overcome the limitations of random centroid initialization. As outlined in Figure 1, the process begins with the initialization phase, where each particle represents a potential solution in the form of a set of cluster centroids. Objective functions are defined to minimize intra-cluster distance (SSW) and maximize inter-cluster separation (SSB), enabling a balanced evaluation of cluster compactness and separation.

The fitness of each particle is assessed based on these two objectives, and the non-dominated solutions are stored in a Pareto-based repository. In the search phase, the particle positions are updated using both personal and global bests selected from the repository using the crowding distance. This iterative process continues until the stopping criteria are met, gradually refining the solutions toward the optimal trade-offs between the two objectives.

At the end of the MOPSO optimization, the repository contains a set of Pareto-optimal centroid configurations. From these, one or more candidate solutions can be selected to initialize the K-Means algorithm. This hybrid approach allows K-Means to begin with well-optimized centroid positions, potentially resulting in more stable and accurate clustering outcomes compared to traditional random initialization.

In addition to the Particle Swarm Optimization (PSO) algorithm, previous studies have also explored other metaheuristic approaches to improve clustering quality. Among the most widely used are Genetic Algorithm [21] (GA) and Differential Evolution [22] (DE). However, findings from several prior works suggest that PSO tends to offer advantages in terms of convergence speed and simplicity of parameters. PSO only requires a few parameters to be configured, such as inertia weight and learning coefficients. In contrast, GA and DE involve more complex parameter settings, including crossover rate, mutation rate, and selection strategies [23] which can significantly influence performance if not properly adjusted. Furthermore, PSO is known for its ability to maintain a good balance between exploration and exploitation during the search process [24], making it especially suitable for clustering problems with high complexity.

3.2 Experimental Results and Analysis

To evaluate the performance of the proposed Multi-Objective PSO, a series of experiments were conducted on a set of well-known benchmark datasets. These five datasets have been described in the Research Method.

The experimental setup in this study is defined as follows: the swarm size is set to $N=40$, and each benchmark dataset is tested independently 30 times, with each execution consisting of 100 iterations. All PSO-based algorithms are terminated upon reaching this maximum number of iterations. The performance of the proposed MOPSO+K-Means method is evaluated using standard clustering metrics, namely the best value, average value, and standard deviation of SSW, SSB, and accuracy. These metrics are used to assess the effectiveness and stability of the proposed method in comparison to standard K-Means across various benchmark datasets.

The performance of MOPSO+K-Means was evaluated using commonly used optimization metrics, namely the average solution and standard deviation. These metrics were used to assess the effectiveness of MOPSO+K-Means in solving the benchmark clustering tasks.

Table 1. Clustering With MOPSO+K-Means

Dataset	Item	SSW	SSB
Atom	Avg.	1191.22	1414.52
	Std.	230.85	238.22
ChainLink	Avg.	1531.74	1711.26
	Std.	167.04	168.519
EngyTime	Avg.	49122.71	85124.33
	Std.	2951.153	3913.554
Target	Avg.	4126.614	7658.074
	Std.	892.147	1006.303
TwoDiamonds	Avg.	1323.322	2863.340
	Std.	42.822	44.191

Although the numerical results demonstrate that the proposed MOPSO+K-Means algorithm performs competitively across various datasets, a deeper analysis provides further insights into its behavior and performance dynamics. From a clustering quality perspective on Table 1, the ideal objective is to

minimize the Sum of Squared Within-cluster distances (SSW) while maximizing the Sum of Squared Between-cluster distances (SSB). In this context, the EngyTime dataset exhibits the highest SSW (49122.71) and SSB (85124.33), indicating that the dataset likely has a large or widely spread structure. Despite this complexity, the algorithm successfully maintains strong inter-cluster separation. In contrast, the TwoDiamonds dataset achieves the most desirable trade-off: it has the lowest SSW (1323.32), relatively high SSB (2863.34), and the lowest standard deviation values, indicating compact clusters, good separation, and high consistency across runs.

For other datasets, such as Target and Atom, the SSW and SSB values fall in a mid-range category, but the relatively high standard deviations, particularly in Atom, highlight variability in the clustering results, suggesting sensitivity to initial conditions or swarm dynamics. Similarly, ChainLink yields results comparable to Atom but also shows considerable variability (Std SSW: 167.04, Std SSB: 168.52), likely reflecting the complexity or overlap in its cluster structure. These findings suggest that while MOPSO+K-Means can handle both simple and

complex datasets, its stability may be challenged under certain data conditions.

Overall, TwoDiamonds emerges as the most stable dataset for this method, whereas EngyTime, despite strong average performance, reveals higher variability possibly due to noise or dispersed data points. These behavioral differences point to both strengths and limitations of the method. The ability of MOPSO+K-Means to find competitive clustering solutions is evident, but further enhancements could improve its robustness. Future research should explore adaptive parameter strategies, improved initialization techniques, or hybrid models to increase the method's reliability across diverse datasets. Additionally, expanding testing to high-dimensional or real-world datasets would help evaluate its scalability and broader applicability.

In this section, an analysis is conducted on the clustering outcomes derived from the implementation of the standard K-Means algorithm and the MOPSO-enhanced K-Means across a range of benchmark datasets. These results highlight the capabilities and limitations of both approaches when confronted with datasets exhibiting varying structural complexities.

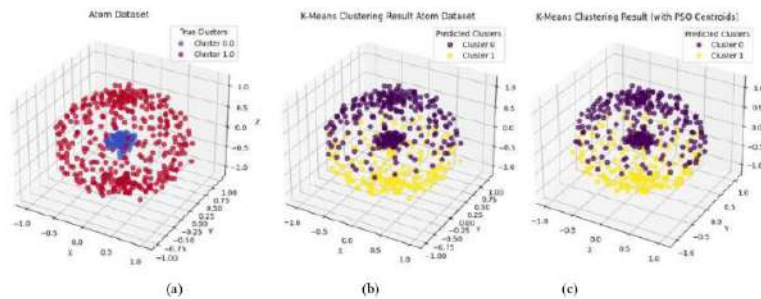


Figure 2. Clustering Result (a) Atom Dataset, (b) K-Means for Atom Dataset, (c) MOPSO+K-Means for Atom Dataset

The clustering results on the Atom dataset highlight the limitations of The Atom dataset illustrates the challenge of clustering data with a concentric structure, comprising a dense core surrounded by a shell. As depicted in Figure 2(a), the ground truth clearly reflects this core-shell configuration. However, the standard K-Means algorithm on Figure 2(b) produces a vertical partition, disregarding the radial nature of the data. This misalignment stems from K-Means assumption of spherical and convex clusters, which proves inadequate for capturing non-linear distributions. The integration of MOPSO+K-Means, as shown in Figure 2(c), results in a more accurate division that successfully distinguishes between the core and the surrounding shell. This demonstrates the ability of MOPSO to guide K-Means toward more structure-aware clustering outcomes in complex spatial configurations.

A similar pattern is observed in the ChainLink dataset, which features two intertwined, non-convex clusters on Figure 3(a). The standard K-Means algorithm on Figure 3(b) again fails to separate the data meaningfully, as it assigns points based on straight-line distances, ignoring the dataset's intricate shape. Conversely, the MOPSO+K-means on Figure 3(c) more effectively untangles the two chains, maintaining their topological distinction. This improved result underscores the role of MOPSO in adapting centroid placement to fit non-linear geometries that would otherwise confound traditional methods.

The EngyTime dataset, in contrast, presents a linearly separable structure, offering a more favorable scenario for K-Means. In Figure 4(a), the data exhibits two well-defined, adjacent clusters. The standard K-Means output on Figure 4(b) broadly captures the separation but misclassifies several points near the decision

boundary. With optimized centroids on Figure 4(c), the clustering becomes cleaner, with reduced boundary ambiguity. This shows that even in simpler datasets,

MOPSO+K-Means contributes by refining the clustering precision and minimizing convergence to suboptimal solutions.

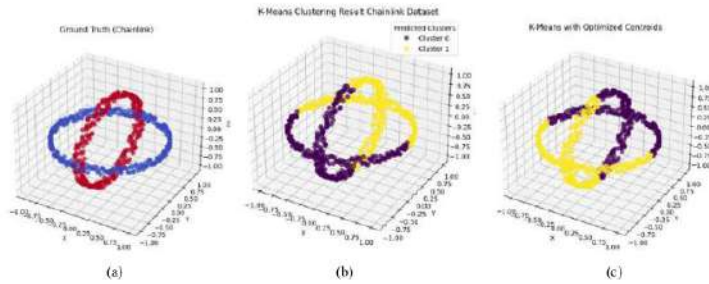


Figure 3. Clustering Result (a) Chainlink Dataset, (b) K-Means for Chainlink Dataset, (c) MOPSO+K-Means for Chainlink Dataset

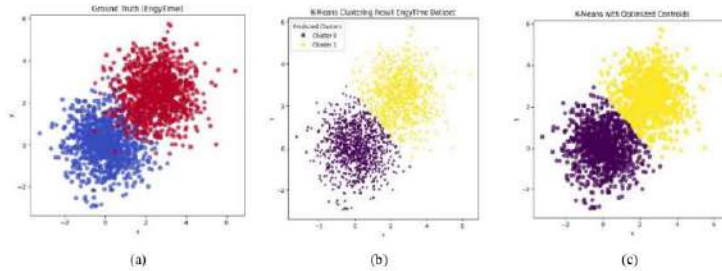


Figure 4. Clustering Result (a) EngyTime Dataset, (b) K-Means for EngyTime Dataset, (c) MOPSO+K-Means for EngyTime Dataset

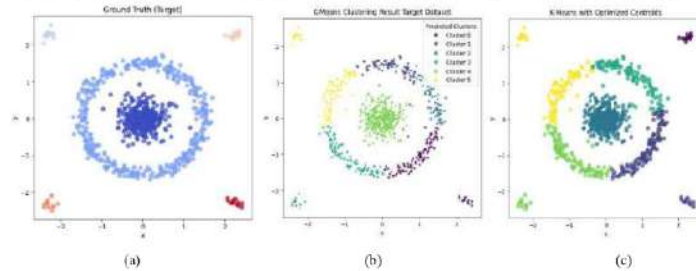


Figure 5. Clustering Results (a) Target Dataset, (b) K-Means for Target Dataset, (c) MOPSO+K-Means for Target Dataset

The analysis of the Target dataset provides further insight into the impact of complex geometries on clustering performance. This dataset contains concentric circular clusters and scattered peripheral groups on Figure 5(a). The clustering produced by the standard K-Means algorithm on Figure 5(b) results in fragmented groupings that poorly reflect the actual layout. Its tendency to impose spherical boundaries leads to significant structural mismatches. The MOPSO+K-Means on Figure 5(c) successfully aligns with the circular patterns and isolates the outlying groups more accurately, reinforcing the value of

optimization in adapting to irregular spatial distributions.

Lastly, the TwoDiamonds dataset poses a moderate challenge due to its diamond-shaped, linearly separated clusters on Figure 6(a). While the standard K-Means on Figure 6(b) performs reasonably well, some inconsistencies are evident along the boundary, suggesting less-than-optimal centroid positioning. By contrast, the optimized version on Figure 6(c) yields a more symmetric and faithful clustering result, demonstrating how MOPSO can enhance K-Means even in datasets that are linearly separable but geometrically unconventional.

Table 2 presents the performance evaluation results of the MOPSO-K-Means algorithm on five benchmark datasets: Atom, ChainLink, EngyTime, Target, and TwoDiamonds. The evaluation was carried out using commonly used optimization metrics, namely the average solution and standard deviation of the SSW

(Sum of Squares Within) and SSB (Sum of Squares Between), along with the best accuracy achieved for each dataset. The objective of this evaluation is to assess the effectiveness of the MOPSO-K-Means algorithm in producing optimal cluster partitions.

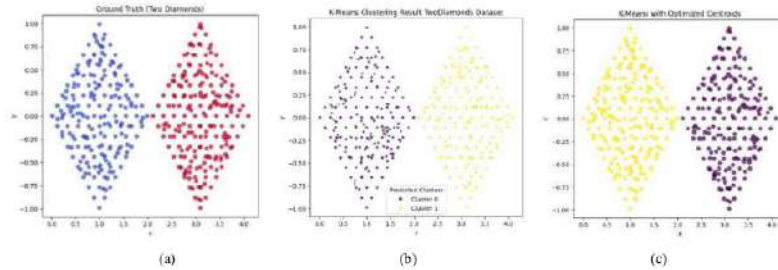


Figure 6. Clustering Result (a) TwoDiamonds Dataset, (b) K-Means Clustering Result TwoDiamonds Dataset, (c) MOPSO-K-Means Clustering Result TwoDiamonds Dataset

Table 2. Comparing Accuracy and Time Computational

Dataset	Accuracy K-means	Best Accuracy MOPSO+K-Means	Average Time Computational MOPSO+K-Means
Atom	54.4%	52.8%	5.1137 seconds
ChainLink	50%	50.2%	6.6889 seconds
EngyTime	95.6%	95.7%	7.617 seconds
Target	0.2692%	59.2%	8.539 seconds
TwoDiamonds	100%	100%	0.844 seconds

Compared to the conventional K-Means algorithm, the results indicate that MOPSO-K-Means generally performs better on most datasets. On the Atom dataset, K-Means achieved an accuracy of 54.4%, while MOPSO-K-Means recorded an accuracy of 52.8%. Although there was a slight decrease, the SSW and SSB values obtained by MOPSO-K-Means still reflect a good and stable cluster distribution, with relatively low standard deviations. For the ChainLink dataset, K-Means achieved 50% accuracy, while MOPSO-K-Means achieved 50.2%, suggesting a slightly better performance in separating the clusters.

Next, on the EngyTime dataset, K-Means reached an accuracy of 95.60%, while MOPSO-K-Means achieved 95.7%. The difference is very small, indicating that both algorithms are equally effective in clustering data with clear cluster structures. However, the most significant improvement was observed on the Target dataset. K-Means achieved only 0.26% accuracy, while MOPSO-K-Means improved the accuracy to 59.2%. This demonstrates that MOPSO-K-Means is more capable of handling datasets with complex or non-linearly separable cluster structures. Lastly, on the TwoDiamonds dataset, both K-Means and MOPSO-K-Means achieved perfect accuracy (100%), indicating that this dataset has a very clear structure that can be easily separated by both algorithms.

The Average Time Computational MOPSO-K-Means column in Table 2 presents the average computational

time required by the algorithm to complete one clustering process for each dataset. This time is measured from the beginning of the centroid optimization using the MOPSO algorithm to the final clustering result produced by K-Means. The values represent the average time taken across 30 independent trials. They indicate the efficiency of the algorithm in solving clustering tasks, which is influenced by the complexity of the data patterns and the algorithm's ability to converge toward optimal solutions.

For example, the TwoDiamonds dataset demonstrates the lowest average computational time (0.844 seconds), which can be attributed to its well-separated and clearly structured clusters. This allows MOPSO to converge quickly, requiring minimal exploration. In contrast, the Target dataset shows the highest computational time (8.539 seconds), suggesting that the algorithm needed significantly more effort to explore the solution space due to the highly irregular and overlapping nature of the data clusters. Similarly, EngyTime and Chainlink also required more computational time, which reflects the greater complexity in their data distributions and the increased difficulty in identifying distinct clusters. Meanwhile, Atom shows a moderate computational time, indicating that while the data is not entirely straightforward, it is still manageable for the algorithm to optimize efficiently.

Overall, the evaluation results show that MOPSO-K-Means has advantages in terms of flexibility and effectiveness in identifying complex cluster structures that conventional K-Means struggles to handle. The relatively small standard deviations across most datasets also indicate that this algorithm can produce stable and consistent solutions in each optimization run. Therefore, MOPSO-K-Means can be considered a more reliable alternative for clustering tasks involving datasets with diverse characteristics.

However, one key limitation of the current study is the assumption that the number of clusters (k) is known beforehand. Although this facilitates benchmarking against ground truth labels, it does not reflect the realities of unsupervised learning tasks, where k must be inferred from the data. In practice, k is a hyperparameter that must be estimated carefully. Common strategies include the elbow method, which identifies diminishing returns in intra-cluster variance as k increases, the silhouette score, which quantifies cluster cohesion and separation, and the gap statistic, which compares clustering performance against that of random reference distributions.

To overcome this limitation, future work could explore the extension of MOPSO to simultaneously optimize both the number of clusters and the centroid positions. This could be framed as a multi-objective optimization problem—balancing cluster compactness, separation, and model complexity—or as a constrained optimization task where k is bounded within a reasonable range. Such an approach would enhance the applicability of the method in real-world scenarios, allowing it to autonomously discover both the optimal clustering structure and its corresponding parameters without relying on prior knowledge.

4. Conclusions

Based on the analysis and evaluation of five benchmark datasets, it can be concluded that the performance of the K-Means algorithm is highly dependent on the shape and structural characteristics of the clusters in the data. On datasets with simple and linearly separable structures, such as TwoDiamonds and EngTime, K-Means performs very well, achieving high accuracy—up to 100%. However, on datasets with non-linear or complex structures, such as Atom, ChainLink, and Target, the algorithm fails to properly separate clusters, resulting in low accuracy and poor alignment with the ground truth.

To address these limitations, the MOPSO+K-Means approach was introduced as an alternative solution. Based on the experimental results, this algorithm shows significant performance improvement on datasets with complex structures—most notably on the Target dataset, where the accuracy increased from 26% (K-Means) to 59.2% (MOPSO+K-Means). In addition, the obtained SSW and SSB values, along with relatively low standard deviations, indicate that MOPSO+K-Means can produce stable and consistent clustering solutions.

Overall, MOPSO+K-Means has proven to be more flexible and reliable in handling various types of cluster structures, making it a more suitable choice for clustering tasks involving non-convex or non-linearly separable data distributions. Such characteristics are often found in real-world applications, including biomedical data analysis where patient subgroups may form irregular patterns, sensor grouping in IoT

environments with overlapping signal zones, and document clustering tasks where semantic relationships are not linearly separable. These application domains can benefit from algorithms that offer both structural flexibility and solution stability, as demonstrated by MOPSO+K-Means in this study.

Acknowledgements

This work was supported by Telkom University.

References

- [1] T. M. Ghazal et al., "Performances of K-Means Clustering Algorithm with Different Distance Metrics," *Intelligent Automation & Soft Computing*, vol. 30, no. 2, pp. 735–742, Aug. 2021, doi: 10.32604/IASC.2021.019067.
- [2] M. Ahmed, R. Seraj, and S. M. S. Islam, "The k-means Algorithm: A Comprehensive Survey and Performance Evaluation," *Electronics* 2020, Vol. 9, Page 1295, vol. 9, no. 8, p. 1295, Aug. 2020, doi: 10.3390/ELECTRONICS9081295.
- [3] A. Abdulhafedh, "Incorporating K-means, Hierarchical Clustering and PCA in Customer Segmentation," *Journal of City and Development*, vol. 3, no. 1, 2021.
- [4] P. Patel, B. Sivaiah, and R. Patel, "Approaches for finding Optimal Number of Clusters using K-Means and Agglomerative Hierarchical Clustering Techniques," in *2022 International Conference on Intelligent Controller and Computing for Smart Power, ICIICSP 2022*, 2022, doi: 10.1109/ICIICSP53532.2022.9862439.
- [5] A. Vouras, S. Langdell, M. Croucher, and E. Vasilaki, "An empirical comparison between stochastic and deterministic centroid initialization for K-means variations," *Mach. Learn.*, vol. 110, no. 8, pp. 1975–2003, Aug. 2021, doi: 10.1007/S10994-021-06021-7/FIGURES/8.
- [6] A. Qataish, M. Braik, D. Alheshish, M. T. Alshammari, A. Alreshidi, and E. J. Alreshidi, "Optimization of K-means clustering method using hybrid capuchin search algorithm," *Journal of Supercomputing*, vol. 80, no. 2, 2024, doi: 10.1007/s11227-023-05540-5.
- [7] Z. Zhang, Q. Feng, J. Huang, Y. Guo, J. Xu, and J. Wang, "A local search algorithm for k-means with outliers," *Neurocomputing*, vol. 450, pp. 230–241, Aug. 2021, doi: 10.1016/J.NEUCOM.2021.04.028.
- [8] A. A. Wani, "Comprehensive analysis of clustering algorithms: exploring limitations and innovative solutions," *PeerJ Comput. Sci.*, vol. 10, pp. 1–45, Aug. 2024, doi: 10.7717/PJCS.2286-FIG-14.
- [9] R. Guatriansyah, N. Suhandi, and F. Antony, "Clustering optimization in RFM analysis Based on k-Means," *IJECS*, vol. 18, no. 1, pp. 470–477, Apr. 2020, doi: 10.11591/ijeecs.v18.i1.pp470-477.
- [10] R. Richard, H. Cao, and M. Wachowicz, "An Automated Clustering Process for Helping Practitioners to Identify Similar EV Charging Patterns across Multiple Temporal Granularities," *International Conference on Smart Cities and Green ICT Systems*, pp. 67–77, 2021, doi: 10.5220/010485000670077.
- [11] M. T. Guerreiro et al., "Anomaly Detection in Automotive Industry Using Clustering Methods—A Case Study," *Applied Sciences* 2021, Vol. 11, Page 9868, vol. 11, no. 21, p. 9868, Oct. 2021, doi: 10.3390/AP11219868.
- [12] M. Jain, V. Satishpal, N. Singh, and S. B. Singh, "An Overview of Variants and Advancements of PSO Algorithm," *MDPI Applied Sciences*, 2022, doi: 10.3390/ap12178392.
- [13] M. C. Thrun and A. Ultsch, "Clustering benchmark datasets exploiting the fundamental clustering problems," *Data Brief*, vol. 30, p. 105501, Jun. 2020, doi: 10.1016/J.DIB.2020.105501.
- [14] A. Ultsch, "Strategies for an Artificial Life System to cluster high dimensional Data," 2004, Accessed: Apr. 14,

2025. [Online]. Available: <https://www.researchgate.net/publication/228932819>
- [15] A. Ullsch, G. Guimaraes, D. Korus, and H. Li, "Knowledge Extraction from Artificial Neural Networks and Applications," *Parallele Datenverarbeitung mit dem Transputer*, pp. 148–162, 1994, doi: 10.1007/978-3-642-78901-4_11.
- [16] P. Mangiameli, S. K. Chen, and D. West, "A comparison of SOM neural network and hierarchical clustering methods," *Eur J Oper Res*, vol. 93, no. 2, pp. 402–417, Sep. 1996, doi: 10.1016/0377-2217(96)00038-0.
- [17] S. P. Chatzis and D. I. Kosmopoulos, "A variational Bayesian methodology for hidden Markov models utilizing Student's-t mixtures," *Pattern Recognit*, vol. 44, no. 2, pp. 295–306, Feb. 2011, doi: 10.1016/j.patcog.2010.09.001.
- [18] J. Poelmans, M. M. Van Hulle, S. Vaece, P. Elzinga, and G. Dedene, "Text mining with emergent self organizing maps and multi-dimensional scaling: A comparative study on domestic violence," *Appl Soft Comput*, vol. 11, no. 4, pp. 3870–3876, Jun. 2011, doi: 10.1016/j.asoc.2011.02.026.
- [19] A. Ullsch, "U*-Matrix: a Tool to visualize Clusters in high dimensional Data," 2004.
- [20] A. Ullsch, "Density Estimation and Visualization for Data Containing Clusters of Unknown Structure," *Studies in Classification, Data Analysis, and Knowledge Organization*, pp. 232–239, 2005, doi: 10.1007/3-540-28084-7_25.
- [21] B. Khusul Khotimah, F. Irfhami, and T. Sundarwati, "A GENETIC ALGORITHM FOR OPTIMIZED INITIAL CENTERS K-MEANS CLUSTERING IN SMEs," *J Theor Appl Inf Technol*, vol. 15, no. 1, 2016, Accessed: Jun. 13, 2025. [Online]. Available: www.jatit.org
- [22] H. He, B. Sun, Y. Yang, and S. Liu, "An improved K-Means clustering based on differential evolution," *J Phys Conf Ser*, vol. 2595, no. 1, p. 012010, Sep. 2023, doi: 10.1088/1742-6596/2595/1/012010.
- [23] A. Hassanat, K. Almohammadi, E. Alkafaween, E. Abunawas, A. Hammouri, and V. B. S. Prasath, "Choosing Mutation and Crossover Ratios for Genetic Algorithms—A Review with a New Dynamic Approach," *Information 2019, Vol. 10, Page 390*, vol. 10, no. 12, p. 390, Dec. 2019, doi: 10.3390/INFO10120390.
- [24] M. Jain, V. Saihpal, N. Singh, and S. B. Singh, "An Overview of Variants and Advancements of PSO Algorithm," *Applied Sciences 2022, Vol. 12, Page 8392*, vol. 12, no. 17, p. 8392, Aug. 2022, doi: 10.3390/AP12178392.

A Multi-Objective Particle Swarm Optimization Approach for Optimizing K-Means Clustering Centroids

ORIGINALITY REPORT

2%

SIMILARITY INDEX

2%

INTERNET SOURCES

2%

PUBLICATIONS

1%

STUDENT PAPERS

PRIMARY SOURCES

1

Submitted to Universitas Muhammadiyah
Surakarta

Student Paper

1%

2

kfh.libraryservices.nhs.uk

Internet Source

<1%

3

Mohammad Siami, Mohsen Naderpour, Jie Lu.
"A Mobile Telematics Pattern Recognition
Framework for Driving Behavior Extraction",
IEEE Transactions on Intelligent
Transportation Systems, 2020

Publication

<1%

Exclude quotes On

Exclude matches < 10 words

Exclude bibliography On