

DAFTAR PUSTAKA

- [1] V. Prakruthi, D. Sindhu and D. S. Anupama Kumar, "Real Time Sentiment Analysis Of Twitter Posts," *2018 3rd International Conference on Computational Systems and Information Technology for Sustainable Solutions (CSITSS)*, 2018, pp. 29-34, doi: 10.1109/CSITSS.2018.8768774.
- [2] Kusrini and M. Mashuri, "Sentiment Analysis In Twitter Using Lexicon Based and Polarity Multiplication," *2019 International Conference of Artificial Intelligence and Information Technology (ICAIT)*, 2019, pp. 365-368, doi: 10.1109/ICAIT.2019.8834477.
- [3] L. Mandloi and R. Patel, "Twitter Sentiments Analysis Using Machine Learning Methods," *2020 International Conference for Emerging Technology (INCET)*, 2020, pp. 1-5, doi: 10.1109/INCET49848.2020.9154183.
- [4] F. A. Wenando, R. Hayami, Bakaruddin and A. Y. Novermahakim, "Tweet Sentiment Analysis for 2019 Indonesia Presidential Election Results using Various Classification Algorithms," *2020 1st International Conference on Information Technology, Advanced Mechanical and Electrical Engineering (ICITAMEE)*, 2020, pp. 279-282, doi: 10.1109/ICITAMEE50454.2020.9398513.
- [5] D. A. Damaratih, "Sentiment Analysis Of Online Lecture Opinions On Twitter Social Media Using Naive Bayes Classifier," *2021 International Conference on Computer Science, Information Technology, and Electrical Engineering (ICOMITEE)*, 2021, pp. 24-28, doi: 10.1109/ICOMITEE53461.2021.9650135.
- [6] M. T. H. K. Tusar and M. T. Islam, "A Comparative Study of Sentiment Analysis Using NLP and Different Machine Learning Techniques on US Airline Twitter Data," *2021 International Conference on Electronics, Communications and Information Technology (ICECIT)*, 2021, pp. 1-4, doi: 10.1109/ICECIT54077.2021.9641336.
- [7] P. S. Mung and S. Phyu, "Effective Analytics on Healthcare Big data Using Ensemble Learning," *2020 IEEE Conference on Computer Applications (ICCA)*, 2020, pp. 1-4, doi: 10.1109/ICCA49400.2020.9022853.
- [8] J. L. Fernández-Alemán, J. M. Carrillo-de-Gea, M. Hosni, A. Idri and G. García-Mateos, "Homogeneous and heterogeneous ensemble classification methods in diabetes disease: a review," *2019 41st Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBC)*, 2019, pp. 3956-3959, doi: 10.1109/EMBC.2019.8856341.



- [9] Y. Feng, X. Wang and J. Zhang, "A Heterogeneous Ensemble Learning Method For Neuroblastoma Survival Prediction," in *IEEE Journal of Biomedical and Health Informatics*, vol. 26, no. 4, pp. 1472-1483, April 2022, doi: 10.1109/JBHI.2021.3073056.
- [10] X. LI, T. SHI, P. LI and W. ZHOU, "Application of Bagging Ensemble Classifier based on Genetic Algorithm in the Text Classification of Railway Fault Hazards," *2019 2nd International Conference on Artificial Intelligence and Big Data (ICAIBD)*, 2019, pp. 286-290, doi: 10.1109/ICAIBD.2019.8836988.
- [11] S. Alelyani, "Stable bagging feature selection on medical data", *Journal of Big Data*, vol. 8, no. 11, pp. 1-18, Jan. 2021, doi: <https://doi.org/10.1186/s40537-020-00385-8>.
- [12] L. Li *et al*, "Cluster-based bagging of constrained mixed-effects models for high spatiotemporal resolution nitrogen oxides prediction over large regions", *Environment International*, vol. 128, pp. 310-323, July 2019, doi: <https://doi.org/10.1016/j.envint.2019.04.057>.
- [13] A. Tareq and N. Hewahi, "Sentiment Analysis of Tweets During COVID-19 Pandemic Using BLSTM," *2021 International Conference on Data Analytics for Business and Industry (ICDABI)*, 2021, pp. 245-249, doi: 10.1109/ICDABI53623.2021.9655932.
- [14] A. Poornima and K. S. Priya, "A Comparative Sentiment Analysis Of Sentence Embedding Using Machine Learning Techniques," *2020 6th International Conference on Advanced Computing and Communication Systems (ICACCS)*, 2020, pp. 493-496, doi: 10.1109/ICACCS48705.2020.9074312.
- [15] S. A. El Rahman, F. A. AlOtaibi and W. A. AlShehri, "Sentiment Analysis of Twitter Data," *2019 International Conference on Computer and Information Sciences (ICCIS)*, 2019, pp. 1-4, doi: 10.1109/ICCISci.2019.8716464.
- [16] R. Katarya, G. A. Nath, D. Singhal and A. Shukla, "Analysing the Twitter Sentiments in COVID-19 Using Machine Learning Algorithms," *2022 International Conference on Advances in Computing, Communication and Applied Informatics (ACCAI)*, 2022, pp. 1-8, doi: 10.1109/ACCAI53970.2022.9752511.
- [17] R. Latifah, R. Baddalwan, P. Meilina, A. D. Saputra and Y. Adharani, "Sentiment Analysis of COVID-19 Vaccines from Indonesian Tweets and News Headlines using Various Machine Learning Techniques," *2021 International Conference on Informatics, Multimedia, Cyber and Information System (ICIMCIS)*, 2021, pp. 69-73, doi: 10.1109/ICIMCIS53775.2021.9699187.



- [18] K. Kurnianingsih, N. F. Al Faridi Hadi, E. Dwi Wardihani, N. Kubota and W. H. Chin, "Ensemble Learning Based on Soft Voting for Detecting Methamphetamine in Urine," *2020 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, 2020, pp. 1-6, doi: 10.1109/FUZZ48607.2020.9177622.
- [19] A. J. Nair, V. G and A. Vinayak, "Comparative study of Twitter Sentiment On COVID - 19 Tweets," *2021 5th International Conference on Computing Methodologies and Communication (ICCMC)*, 2021, pp. 1773-1778, doi: 10.1109/ICCMC51019.2021.9418320.
- [20] U. N. Wisesty, R. Rismala, W. Munggana and A. Purwarianti, "Comparative Study of Covid-19 Tweets Sentiment Classification Methods," *2021 9th International Conference on Information and Communication Technology (ICoICT)*, 2021, pp. 588-593, doi: 10.1109/ICoICT52021.2021.9527533.
- [21] P. Puvar et al., "Heart Disease Detection using Ensemble Learning Approach", *International Research Journal of Engineering and Technology (IRJET)*, vol. 8, pp. 1414-1418, May 2021.
- [22] S. Mehta, P. Rana, S. Singh, A. Sharma and P. Agarwal, "Ensemble Learning Approach for Enhanced Stock Prediction," *2019 Twelfth International Conference on Contemporary Computing (IC3)*, 2019, pp. 1-5, doi: 10.1109/IC3.2019.8844891.
- [23] H. R. Sanabila and W. Jatmiko, "Ensemble Learning on Large Scale Financial Imbalanced Data," *2018 International Workshop on Big Data and Information Security (IWBIS)*, 2018, pp. 93-98, doi: 10.1109/IWBIS.2018.8471702.
- [24] J. Tie, X. Lei and Y. Pan, "Metabolite-disease association prediction algorithm combining DeepWalk and random forest," in *Tsinghua Science and Technology*, vol. 27, no. 1, pp. 58-67, Feb. 2022, doi: 10.26599/TST.2021.9010003.
- [25] C. Zhao, Y. Xin, X. Li, Y. Yang, and Y. Chen, "A Heterogeneous Ensemble Learning Framework for Spam Detection in Social Networks with Imbalanced Data," *Applied Sciences*, vol. 10, no. 3, p. 936, Jan. 2020, doi: 10.3390/app10030936.
- [26] Murni, T. Handhika, A. Fahrurrozi, I. Sari, D. P. Lestari and R. I. M. Zen, "Hybrid Method for Sentiment Analysis Using Homogeneous Ensemble Classifier," *2019 2nd International Conference of Computer and Informatics Engineering (IC2IE)*, 2019, pp. 232-236, doi: 10.1109/IC2IE47452.2019.8940896.



- [27] K. S. Gan, K. O. Chin, P. Anthony and S. V. Chang, “Homogeneous Ensemble FeedForward Neural Network in CIMB Stock Price Forecasting,” *2018 IEEE International Conference on Artificial Intelligence in Engineering and Technology (IICAET)*, 2018, pp. 1-6, doi: 10.1109/IICAET.2018.8638452.
- [28] Imamah and F. H. Rachman, “Twitter Sentiment Analysis of Covid-19 Using Term Weighting TF-IDF And Logistic Regresion,” *2020 6th Information Technology International Seminar (ITIS)*, 2020, pp. 238-242, doi: 10.1109/ITIS50118.2020.9320958.
- [29] S. Singh, K. Kumar and B. Kumar, “Sentiment Analysis of Twitter Data Using TF-IDF and Machine Learning Techniques,” *2022 International Conference on Machine Learning, Big Data, Cloud and Parallel Computing (COM-IT-CON)*, 2022, pp. 252-255, doi: 10.1109/COM-IT-CON54601.2022.9850477.
- [30] R. Man and K. Lin, “Sentiment Analysis Algorithm Based on BERT and Convolutional Neural Network,” *2021 IEEE Asia-Pacific Conference on Image Processing, Electronics and Computers (IPEC)*, 2021, pp. 769-772, doi: 10.1109/IPEC51340.2021.9421110.
- [31] A. A. Marouf, J. G. Rokne, and R. Alhaji, “Detecting and Understanding Sentiment Trends and Emotion Patterns of Twitter Users—A Study on the Demise of a Bollywood Celebrity,” *Big Data and Cognitive Computing*, vol. 6, no. 4, p. 129, Oct. 2022, doi: 10.3390/bdcc6040129.
- [32] H. Fang, G. Xu, Y. Long, and W. Tang, “An Effective ELECTRA-Based Pipeline for Sentiment Analysis of Tourist Attraction Reviews,” *Applied Sciences*, vol. 12, no. 21, p. 10881, Oct. 2022, doi: 10.3390/app122110881.
- [33] S. S. Aljameel et al., “A Sentiment Analysis Approach to Predict an Individual’s Awareness of the Precautionary Procedures to Prevent COVID-19 Outbreaks in Saudi Arabia,” *International Journal of Environmental Research and Public Health*, vol. 18, no. 1, p. 218, Dec. 2020, doi: 10.3390/ijerph18010218.
- [34] M. Reusens, M. Reusens, M. Callens, S. vanden Broucke, and B. Baesens, “Comparison of Different Modeling Techniques for Flemish Twitter Sentiment Analysis,” *Analytics*, vol. 1, no. 2, pp. 117–134, Oct. 2022, doi: 10.3390/analytics1020009.
- [35] J. J. E. Macrohon, C. N. Villavicencio, X. A. Inbaraj, and J.-H. Jeng, “A Semi-supervised Approach to Sentiment Analysis of Tweets during the 2022 Philippine Presidential Election,” *Information*, vol. 13, no. 10, p. 484, Oct. 2022, doi: 10.3390/info13100484.



- [36] A. Jenasamanta and S. Mohapatra, “An automated system for the assessment and grading of adolescent delinquency using a machine learning-based soft voting framework,” *Humanities and Social Sciences Communications*, vol. 9, no. 1, Oct. 2022, doi: 10.1057/s41599-022-01407-x.
- [37] A. Taha, “Intelligent Ensemble Learning Approach for Phishing Website Detection Based on Weighted Soft Voting,” *Mathematics*, vol. 9, no. 21, p. 2799, Nov. 2021, doi: 10.3390/math9212799.
- [38] R. Wardoyo, A. Musdholifah, G. Angga Pradipta and I. N. Hariyasa Sanjaya, “Weighted Majority Voting by Statistical Performance Analysis on Ensemble Multiclassifier,” *2020 Fifth International Conference on Informatics and Computing (ICIC)*, Gorontalo, Indonesia, 2020, pp. 1-8, doi: 10.1109/ICIC50835.2020.9288552.
- [39] A. Dogan and D. Birant, “A Weighted Majority Voting Ensemble Approach for Classification,” *2019 4th International Conference on Computer Science and Engineering (UBMK)*, Samsun, Turkey, 2019, pp. 1-6, doi: 10.1109/UBMK.2019.8907028.
- [40] R. Islam, M. I. Sayed, S. Saha, M. J. Hossain, and M. A. Masud, “Android malware classification using Optimum Feature Selection and Ensemble Machine Learning,” *Internet of Things and Cyber-Physical Systems*, vol. 3, pp. 100–111, 2023, doi:10.1016/j.iotcps.2023.03.001.
- [41] P. Agrawal, H. F. Abutarboush, T. Ganesh and A. W. Mohamed, “Metaheuristic Algorithms on Feature Selection: A Survey of One Decade of Research (2009-2019),” in *IEEE Access*, vol. 9, pp. 26766-26791, 2021, doi: 10.1109/ACCESS.2021.3056407.
- [42] T. M. Shami, A. A. El-Saleh, M. Alswaitti, Q. Al-Tashi, M. A. Summakieh and S. Mirjalili, “Particle Swarm Optimization: A Comprehensive Survey,” in *IEEE Access*, vol. 10, pp. 10031-10061, 2022, doi: 10.1109/ACCESS.2022.3142859.
- [43] M. Grandini, Enrico Bagli, and Giorgio Visani, “Metrics for Multi-Class Classification: an Overview,” *arXiv (Cornell University)*, Aug. 2020, doi: <https://doi.org/10.48550/arxiv.2008.05756>.



LAMPIRAN



Lampiran 1. Script Program untuk Prapemrosesan Data

```

1 import pandas as pd
2 import re
3 import nltk
4 from sklearn.pipeline import Pipeline
5 from textblob import TextBlob
6 nltk.download('stopwords')
7 from nltk.corpus import stopwords
8 from Sastrawi.Stemmer.StemmerFactory import StemmerFactory
9
10 def clean(data):
11
12     def casefolding(tweet):
13         tweet = tweet.lower()
14         #tweet = tweet.strip(" ")
15         tweet = re.sub(r'[\r\n]+', ' ', tweet)
16         tweet = re.sub(r'https?://\S+', ' ', tweet)
17         tweet = re.sub(r'@\w+', ' ', tweet)
18         # tweet = re.sub(r'[A-Za-z0-9\s]+', ' ', tweet)
19         tweet = re.sub(r'[A-Za-z\s]+', ' ', tweet)
20         tweet = re.sub(r'\W+', ' ', tweet)
21         return tweet
22     data['tweet'] = data['tweet'].apply(casefolding)
23
24     #Tokenizing
25     def token(tweet):
26         nstr = tweet.split(' ')
27         dat = []
28         for hu in nstr:
29             if (hu != ''):
30                 dat.append(hu)
31         return dat
32     data['tweet'] = data['tweet'].apply(token)
33
34     dictionary = {
35         "knp": "kenapa", "bru": "baru", "lg": "lagi", "dr": "dari", "sdh": "sudah",
36         "dg": "dengan", "tp": "tapi", "sblm": "sebelum", "jg": "juga", "sy": "saya",
37         "gk": "gak", "bs": "bisa", "gpp": "gapapa", "tpi": "tapi", "sm": "sama",
38         "bgt": "banget", "gmn": "bagaimana", "dgn": "dengan", "bsa": "bisa",
39         "jd": "jadi", "tdk": "tidak", "spt": "seperti", "sdng": "sedang",
40         "gini": "begini", "kmu": "kamu", "krn": "karena", "dlm": "dalam",
41         "bkn": "bukan", "mrnt": "menurut", "trs": "terus", "sdikit": "sedikit", "sdkt": "sedikit",
42         "brp": "berapa", "gtu": "begitu", "w": "aku", "kmn": "kemana", "lm": "lama",
43         "smg": "semoga", "bnyk": "banyak", "mn": "mana", "slm": "selama", "gk": "gak",
44         "btl": "betul", "bnr": "benar", "pd": "pada", "yg": "yang", "lw": "kalau",
45         "mlm": "malam", "bsr": "besar", "org": "orang", "ntr": "nant", "dlu": "dulu",
46         "bwh": "bawah", "tmn": "teman", "syg": "sayang", "trus": "terus", "tp": "tapi",
47         "thn": "tahun", "cma": "cuma", "smpe": "sampai", "pdhl": "padahal", "lh": "lah",
48         "ny": "nya", "bbrrp": "beberapa", "lmyn": "lumayan", "n": "dan",
49         "pa": "apa", "tar": "sementar", "bca": "baca", "lm": "lama", "ps": "pas",
50         "lsg": "langsung", "msh": "masih", "bnyk": "banyak", "mrnt": "menurut",
51         "hnya": "hanya", "msh": "masalah", "sih": "sih", "hri": "hari",
52         "dri": "dari", "dkt": "dekat", "skr": "sekarang", "jm": "jam", "y": "iya",
53         "k": "oke", "byk": "banyak", "ntar": "nant", "msk": "masuk", "ktnya": "katanya",
54         "mksh": "makasih", "br": "baru", "mgkn": "mungkin", "nanya": "bertanya",
55         "mau": "ingin", "u": "kamu", "tngh": "tengah", "blm": "belum",
56         "q": "aku", "nya": "dia", "s": "sudah", "btw": "by the way", "g": "ga",
57         "w": "gw", "tw": "tahu", "kt": "kita", "bgs": "bagus", "tpi": "tapi", "sm": "sama",
58         "jg": "juga", "pd": "pada", "sdh": "sudah", "dn": "dan", "bw": "bawa",
59         "sndr": "sendiri", "jln": "jalan", "ngr": "negara", "gt": "gitu", "bhs": "bahasa",
60         "klo": "kalau", "lh": "lah", "bgt": "banget", "sdikit": "sedikit",
61         "tig": "tolong", "plg": "pulang", "dl": "dulu", "sms": "pesan",
62         "mksd": "maksud", "ngmng": "ngomong", "mlu": "malu", "thnks": "thanks",
63         "bgtu": "begitu", "ktny": "katanya", "mksdny": "maksudnya",
64         "mkasih": "makasih", "blg": "bilang", "tmn": "teman", "kbr": "kabar",
65         "ksh": "kasih", "kmrn": "kemarin", "kbtulan": "kebetulan",
66         "mndng": "mending", "brp": "berapa", "kluar": "keluar", "stlh": "setelah",
67         "shbt": "sahabat", "pny": "punya", "mgkn": "mungkin", "mgk": "mungkin",
68         "jwb": "jawab", "ntn": "nonton", "sblm": "sebelum", "byr": "bayar"
69     }
70
71     def ubah_singkatan(tweet):
72         ubah = []
73         for i in tweet:
74             p = 0
75             for j in dictionary:
76                 if i == j:
77                     p = p + 1
78                     ubah.append(dictionary[j])
79                     break
80             if p == 0:
81                 ubah.append(i)
82         return ubah
83     data['tweet'] = data['tweet'].apply(ubah_singkatan)

```



```

83 #Filtering/Stopwords Removal
84 def stopword_removal(tweet):
85     pkg_stopwords = stopwords.words('Indonesian','english')
86     more_stopwords = [
87         'yg', 'utk', 'cuman', 'deh', 'btw', 'gua', 'gue', 'gw', 'lo', 'lu',
88         'kalo', 'trs', 'jd', 'nih', 'ntar', 'nya', 'lg', 'gk', 'ecusli', 'dpt',
89         'dr', 'kpn', 'kok', 'kyk', 'donk', 'dong', 'yah', 'u', 'ya', 'km', 'eh',
90         'sih', 'eh', 'bang', 'br', 'rp', 'jt', 'kan', 'gpp', 'sm', 'usah',
91         'mas', 'sob', 'thx', 'ato', 'jg', 'wkwk', 'mak', 'haha', 'ly', 'k',
92         'tp', 'hahaha', 'dg', 'dri', 'duh', 'ye', 'wkwk', 'syg', 'btw',
93         'gaes', 'guys', 'moga', 'kmrn', 'nemu', 'yuk', 'yukk', 'yukkk',
94         'klas', 'iw', 'ew', 'lho', 'loh', 'sbny', 'org', 'gtu', 'bwt',
95         'klrga', 'clau', 'lbh', 'cpet', 'ku', 'wke', 'mba', 'mas', 'sdh', 'kmrn',
96         'oi', 'spt', 'dlm', 'bs', 'krn', 'sapa', 'spt', 'sh', 'wakakaka', 'wakakak',
97         'sihh', 'sihhh', 'hehe', 'hehehe', 'lh', 'dgn', 'la', 'kl', 'ttg', 'mana', 'kmna', 'knn',
98         'tuh', 'dah', 'kek', 'ko', 'pls', 'bbrp', 'pd', 'mah', 'dhhh', 'dhh', 'dh',
99         'kpd', 'tuh', 'kzl', 'sbl', 'byr', 'byar', 'sli', 'si', 'cm', 'cmn', 'sy', 'hahahaha', 'weh',
100        'dlu', 'tuh', 'nah', 'naahhh', 'nahh'
101    ]
102    filtering = more_stopwords + pkg_stopwords
103    x = []
104    data = []
105    def myFunc(x):
106        if x in filtering:
107            return False
108        else:
109            return True
110    fit = filter(myFunc, tweet)
111    for x in fit:
112        data.append(x)
113    return data
114    data['tweet'] = data['tweet'].apply(stopword_removal)
115
116 #Stemming
117 def stemming(tweet):
118     factory = StemmerFactory()
119     stemmer = factory.create_stemmer()
120     do=[]
121
122     for w in tweet:
123         dt = stemmer.stem(w)
124         do.append(dt)
125
126     d_clean=[]
127     d_clean= " ".join(do)
128     print(d_clean)
129     return d_clean
130
131    data['tweet'] = data['tweet'].apply(stemming)
132
133    return data

```



Lampiran 2. Script Program untuk Vektorisasi Data menggunakan TF-IDF

```

1 import math
2 import pandas as pd
3 import numpy as np
4
5 def tf(d_list, d_semua):
6     data_dict = {}
7     data_tf = []
8     for kalimat in d_list:
9         for kata in d_semua:
10             data_dict[kata] = kalimat.split().count(kata)/len(kalimat.split())
11             data_tf.append(data_dict)
12     data_dict = {}
13     return data_tf
14
15 def df(d_list, d_semua):
16     df_dict = {}
17     for kalimat in d_list:
18         for term in d_semua:
19             if term in kalimat.split():
20                 if term in df_dict:
21                     df_dict[term] += 1
22                 else:
23                     df_dict[term] = 1
24             else:
25                 if term not in df_dict:
26                     df_dict[term] = 0
27     return df_dict
28
29 def idf(df_dict, d_list):
30     data_idf = {}
31     for kata in df_dict:
32         hitung = math.log(len(d_list)/(df_dict[kata]+1))
33         data_idf[kata] = hitung
34     return data_idf
35
36 def tfidf(data_tf, data_idf):
37     hasil = []
38     data_sementara = {}
39     for index in range(len(data_tf)):
40         for i in data_tf[index]:
41             kalkulasi = data_tf[index][i] * data_idf[i]
42             data_sementara[i] = round(kalkulasi, 4)
43         hasil.append(data_sementara)
44         data_sementara = {}
45     return hasil
46
47 def tfidf_proses(data, d_semua=[], idf_train={}):
48     n=1
49     if d_semua == []:
50         n = 0
51     d_list = []
52     for i in data['tweet']:
53         d_list.append(i)
54
55     def token(tweet):
56         nstr = tweet.split(' ')
57         dat = []
58         for hu in nstr:
59             if {hu != ''}:
60                 dat.append(hu)
61         return dat
62     data['tweet'] = data['tweet'].apply(token)
63
64     if n == 0:
65         d_semua = []
66         for i in data['tweet']:
67             for j in i:
68                 if j not in d_semua:
69                     d_semua.append(j)
70
71     #TF
72     data_tf = tf(d_list, d_semua)
73
74     #DF
75     df_dict = df(d_list, d_semua)
76
77     #IDF
78     if idf_train == {}:
79         data_idf = idf(df_dict, d_list)
80     else:
81         data_idf = idf_train
82
83     #TFIDF
84     hasil = tfidf(data_tf, data_idf)
85     tfidf_data = pd.DataFrame(hasil)
86
87     if n==0:
88         all_data=tfidf_data.join(data['label'])
89         all_data.head()
90
91     return tfidf_data, all_data, d_semua, data_idf
92
93 else:
94     return tfidf_data
95

```



Lampiran 3. Script Program Algoritma Naïve Bayes

```

1 import numpy as np
2
3 class MultiNB:
4     def __init__(self, alpha=1):
5         self.alpha = alpha
6
7     def _prior(self): # CHECKED
8         P = np.zeros((self.n_classes_))
9         _, self.dist = np.unique(self.y, return_counts=True)
10        for i in range(self.classes_.shape[0]):
11            P[i] = self.dist[i] / self.n_samples
12        return P
13
14    def fit(self, X, y): # CHECKED, matches with sklearn
15        self.y = y
16        self.n_samples, self.n_features = X.shape
17        self.classes_ = np.unique(y)
18        self.n_classes_ = self.classes_.shape[0]
19        self.class_priors_ = self._prior()
20        # distinct values in each features
21        self.uniques = []
22        for i in range(self.n_features):
23            tmp = np.unique(X[:,i])
24            self.uniques.append( tmp )
25        self.N_yi = np.zeros((self.n_classes_, self.n_features)) # feature count
26        self.N_y = np.zeros((self.n_classes_)) # total count
27        for i in self.classes_: # x axis
28            indices = np.argwhere(self.y==i).flatten()
29            columnwise_sum = []
30            for j in range(self.n_features): # y axis
31                columnwise_sum.append(np.sum(X[indices,j]))
32            self.N_yi[i] = columnwise_sum # 2d
33            self.N_y[i] = np.sum(columnwise_sum) # 1d
34        return self
35
36    def _theta(self, x_i, i, h):
37        Nyi = self.N_yi[h,i]
38        Ny = self.N_y[h]
39        numerator = Nyi + self.alpha
40        denominator = Ny + (self.alpha * self.n_features)
41        return (numerator / denominator)**x_i
42
43    def _likelihood(self, x, h):
44        tmp = []
45        for i in range(x.shape[0]):
46            tmp.append(self._theta(x[i], i, h))
47        return np.prod(tmp)
48
49    def predict(self, X):
50        samples, features = X.shape
51        self.predict_proba = np.zeros((samples, self.n_classes_))
52        for i in range(X.shape[0]):
53            joint_likelihood = np.zeros((self.n_classes_))
54
55            for h in range(self.n_classes_):
56                joint_likelihood[h] = self.class_priors_[h] * self._likelihood(X[i], h) # P(y) P(X|y)
57            denominator = np.sum(joint_likelihood)
58            for h in range(self.n_classes_):
59                numerator = joint_likelihood[h]
60                self.predict_proba[i,h] = (numerator / denominator)
61        indices = np.argmax(self.predict_proba, axis=1)
62        return self.classes_[indices], self.predict_proba

```



Lampiran 4. Script Program Algoritma SVM

```

1 import numpy as np
2 from sklearn.base import BaseEstimator, ClassifierMixin
3 from sklearn.utils import check_random_state
4 from sklearn.preprocessing import LabelEncoder
5
6 def projection_simplex(v, z=1):
7     n_features = v.shape[0]
8     u = np.sort(v)[::-1]
9     cssv = np.cumsum(u) - z
10    ind = np.arange(n_features) + 1
11    cond = u - cssv / ind > 0
12    rho = ind[cond][-1]
13    theta = cssv[cond][-1] / float(rho)
14    w = np.maximum(v - theta, 0)
15    return w
16
17 class MulticlassSVM(BaseEstimator, ClassifierMixin):
18
19     def __init__(self, C=1, max_iter=50, tol=0.05,
20                  random_state=None, verbose=0):
21         self.C = C
22         self.max_iter = max_iter
23         self.tol = tol
24         self.random_state = random_state
25         self.verbose = verbose
26
27     def _partial_gradient(self, X, y, i):
28         g = np.dot(X[i], self.coef_.T) + 1
29         g[y[i]] -= 1
30         return g
31
32     def _violation(self, g, y, i):
33         smallest = np.inf
34         for k in range(g.shape[0]):
35             if k == y[i] and self.dual_coef_[k, i] >= self.C:
36                 continue
37             elif k != y[i] and self.dual_coef_[k, i] >= 0:
38                 continue
39             smallest = min(smallest, g[k])
40         return g.max() - smallest
41
42     def _solve_subproblem(self, g, y, norms, i):
43         Ci = np.zeros(g.shape[0])
44         Ci[y[i]] = self.C
45         beta_hat = norms[i] * (Ci - self.dual_coef_[i, i]) + g / norms[i]
46         z = self.C * norms[i]
47         beta = projection_simplex(beta_hat, z)
48         return Ci - self.dual_coef_[i, i] - beta / norms[i]
49
50     def fit(self, X, y):
51         n_samples, n_features = X.shape
52         # Normalize labels.
53         self._label_encoder = LabelEncoder()
54         y = self._label_encoder.fit_transform(y)
55         # Initialize primal and dual coefficients.
56         n_classes = len(self._label_encoder.classes_)
57         self.dual_coef_ = np.zeros((n_classes, n_samples), dtype=np.float64)
58         self.coef_ = np.zeros((n_classes, n_features))
59         # Pre-compute norms.
60         norms = np.sqrt(np.sum(X ** 2, axis=1))
61         # Shuffle sample indices.
62         rs = check_random_state(self.random_state)
63         ind = np.arange(n_samples)
64         rs.shuffle(ind)
65
66         violation_init = None
67         for it in range(self.max_iter):
68             violation_sum = 0
69             for ii in range(n_samples):
70                 i = ind[ii]
71                 # All-zero samples can be safely ignored.
72                 if norms[i] == 0:
73                     continue
74                 g = self._partial_gradient(X, y, i)
75                 v = self._violation(g, y, i)
76                 violation_sum += v
77                 if v < 1e-12:
78                     continue
79                 # Solve subproblem for the ith sample.
80                 delta = self._solve_subproblem(g, y, norms, i)
81                 # Update primal and dual coefficients.
82                 self.coef_ += (delta * X[i][:, np.newaxis]).T
83                 self.dual_coef_[i, i] += delta
84                 if it == 0:
85                     violation_init = violation_sum
86                 vratio = violation_sum / violation_init
87                 if vratio < self.tol:
88                     if self.verbose >= 1:
89                         print("Converged")
90                     break
91             return self
92
93     def predict_proba(self, X):
94         decision = np.dot(X, self.coef_.T)
95         return decision
96
97     def predict(self, X):
98         decision = self.predict_proba(X)
99         pred = decision.argmax(axis=1)
100        return self._label_encoder.inverse_transform(pred)

```



Lampiran 5. Script Program Algoritma Logistic Regression

```

1 import numpy as np
2
3 class SoftmaxRegression:
4     def __init__(self, label, fitur):
5         seed_value = 50
6         np.random.seed(seed_value)
7         self.weights = np.random.uniform(0, 1, (label, fitur))
8         self.label = label
9
10    def softmax(self, x):
11        return np.exp(x) / np.sum(np.exp(x))
12
13    def train_sgd(self, x_train, y_train, max_itr=1, alpha=0.1, lamda=0):
14        for itr in range(max_itr):
15            for i in range(len(x_train)):
16                x = x_train[i].reshape(-1, 1)
17                y = np.zeros((self.label, 1))
18                y[y_train[i]] = 1
19                h_y = self.softmax(np.dot(self.weights, x))
20                self.weights = self.weights - alpha * (np.dot((h_y - y), x.T)) + alpha * lamda * self.weights
21            return self.weights
22
23    def train_bgd(self, x_train, y_train, max_itr=100, alpha=0.1, lamda=0): # batch gradient decent
24        print("-----Training-----")
25        for itr in range(max_itr):
26            print('iteration:%d' % itr)
27            err = np.zeros((5, 14942))
28            for i in range(len(x_train)):
29                x = x_train[i].reshape(-1, 1)
30                y = np.zeros((5, 1))
31                y[y_train.iloc[i]] = 1
32                h_y = self.softmax(np.dot(self.weights, x))
33                singleErr = np.dot((h_y - y), x.T)
34                err += singleErr
35            err = err / len(x_train)
36            regular = (alpha * lamda * self.weights) / len(x_train)
37            self.weights = self.weights - alpha * err + regular
38            print('-----Train Finished-----')
39            return self.weights
40
41    def predict_proba(self, x_test):
42        y_proba = []
43        for i in range(len(x_test)):
44            x = x_test[i]
45            y_proba.append(np.dot(self.weights, x))
46        return y_proba
47
48    def predict(self, x_test):
49        y_predict = []
50        for i in range(len(x_test)):
51            x = x_test[i]
52            y = np.argmax(np.dot(self.weights, x))
53            y_predict.append(y)
54        y_predict = np.array(y_predict)
55        return y_predict
56
57    def test(self, y_test, y_predict):
58        print('-----Testing-----')
59        accuracy = 0
60        for i in range(len(y_test)):
61            if y_predict[i] == y_test.iloc[i]:
62                accuracy += 1
63        accuracy = accuracy / len(y_test)
64        print('-----Test Finished-----')
65        return accuracy

```



Lampiran 6. Script Program Confusion Matrix untuk Menghitung Performa

```

1 def performance_test(y_test, predict):
2     conf_matrix = confusion_matrix(y_test, predict)
3
4     if len(conf_matrix) == 3:
5         accuracy = (conf_matrix[0,0] + conf_matrix[1,1] + conf_matrix[2,2])/len(y_test)
6         precision = ((conf_matrix[0,0]/(conf_matrix[0,0]+conf_matrix[1,0]+conf_matrix[2,0])) +
7             (conf_matrix[1,1]/(conf_matrix[1,1]+conf_matrix[0,1]+conf_matrix[2,1])) + (conf_matrix[2,2]/(conf_matrix[2,2]+
8             conf_matrix[0,2]+conf_matrix[1,2]))) / 3
9         recall = ((conf_matrix[0,0]/(conf_matrix[0,0]+conf_matrix[0,1]+conf_matrix[0,2])) +
10             (conf_matrix[1,1]/(conf_matrix[1,1]+conf_matrix[1,0]+conf_matrix[1,2])) + (conf_matrix[2,2]/(conf_matrix[2,2]+
11             conf_matrix[2,0]+conf_matrix[2,1]))) / 3
12         f1 = (2 * (precision * recall)) / (precision + recall)
13     else:
14         accuracy = (conf_matrix[0,0] + conf_matrix[1,1]) / len(y_test)
15         precision = conf_matrix[1,1] / (conf_matrix[1,1]+conf_matrix[0,1])
16         recall = conf_matrix[1,1]/(conf_matrix[1,1]+conf_matrix[1,0])
17         f1 = (2 * (precision * recall)) / (precision + recall)
18
19     return conf_matrix, accuracy, precision, recall, f1

```



Lampiran 7. Script Program Heterogeneous Ensemble Learning

```

1 import numpy as np
2 import pandas as pd
3 import math
4 import seaborn as sns
5 import matplotlib.pyplot as plt
6 from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score
7 from sklearn.metrics import confusion_matrix
8 from scipy.sparse import csr_matrix
9 from sklearn.preprocessing import LabelEncoder
10 from sklearn.feature_selection import SelectKBest, chi2
11 from MLAlgorithm.NaiveBayes import MultiNB
12 from MLAlgorithm.LogisticRegression import SoftmaxRegression
13 from MLAlgorithm.SVM import MulticlassSVM
14 import pickle
15
16 training_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_training_data_new.csv", encoding='utf-8')
17 X_train = csr_matrix(training_data.iloc[:, :-1]).toarray()
18 label_encoding = LabelEncoder()
19 y_train = label_encoding.fit_transform(training_data['label'])
20
21 #PROSES TRAINING
22 #Naive Bayes
23 NBModel = MultiNB()
24 NBModel.fit(X_train, y_train)
25 #SVM
26 SVMModel = MulticlassSVM(C=1, tol=0.05, max_iter=1000, random_state=0, verbose=1)
27 SVMModel.fit(X_train, y_train)
28 #Logistic Regression
29 LRModel = SoftmaxRegression(len(set(y_train)), len(X_train[0]))
30 LRModel.train_sgd(X_train, y_train)
31
32 testing_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_testing_data_new.csv", encoding='utf-8')
33 X_test = csr_matrix(testing_data.iloc[:, :-1]).toarray()
34 y_test = label_encoding.fit_transform(testing_data['label'])
35
36 #PROSES TESTING
37 #Naive Bayes
38 NBPredicted_prob = NBModel.predict(X_test)
39 #SVM
40 SVMPredicted_prob = SVMModel.predict_proba(X_test)
41 #Logistic Regression
42 LRPredicted_prob = LRModel.predict_proba(X_test)
43
44 #PROSES VOTING DENGAN WEIGHTED SOFT VOTING
45 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/optimal_weight/heterogeneous_optimal_weights_new.pkl', 'rb') as f:
46     optimal_weights = pickle.load(f)
47     predicted_hetero = []
48     for baris in range(len(NBPredicted_prob)):
49         jml = ((np.array(NBPredicted_prob[baris]) * optimal_weights[0]) + (np.array(SVMPredicted_prob[baris]) *
         optimal_weights[1]) + (np.array(LRPredicted_prob[baris]) * optimal_weights[2])) / 3
50         if jml.argmax() == 0:
51             predicted_hetero.append(0)
52         elif jml.argmax() == 1:
53             predicted_hetero.append(1)
54         elif jml.argmax() == 2:
55             predicted_hetero.append(2)
56     predicted_hetero = np.array(predicted_hetero)
57
58 conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_hetero)

```



Lampiran 8. Script Program Homogeneous Ensemble Learning Naïve Bayes

```

1 import numpy as np
2 import pandas as pd
3 import math
4 import seaborn as sns
5 import matplotlib.pyplot as plt
6 from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score
7 from sklearn.metrics import confusion_matrix
8 from scipy.sparse import csr_matrix
9 from sklearn.preprocessing import LabelEncoder
10 from MLAlgorithm.NaiveBayes import MultiNB
11 import pickle
12
13 training_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_training_data_new.csv", encoding='utf-8')
14
15 #PROSES TRAINING
16 n_sample = 10
17 models = []
18 for i in range(n_sample):
19     sample_data = training_data.sample(frac=0.7, replace=False, random_state=i)
20
21     X_train = csr_matrix(sample_data.iloc[:, :-1]).toarray()
22     label_encoding = LabelEncoder()
23     y_train = label_encoding.fit_transform(sample_data['label'])
24
25     NBModel = MultiNB()
26     NBModel.fit(X_train, y_train)
27     models.append(NBModel)
28
29 testing_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_testing_data_new.csv", encoding='utf-8')
30 X_test = csr_matrix(testing_data.iloc[:, :-1]).toarray()
31 y_test = label_encoding.fit_transform(testing_data['label'])
32
33 #PROSES TESTING
34 NBPredicted_prob = []
35 for model in models:
36     predicted, prob = model.predict(X_test)
37     NBPredicted_prob.append(prob)
38
39 #PROSES VOTING DENGAN WEIGHTED SOFT VOTING
40 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/optimal_weight/homogeneousNB_optimal_weights_new.pkl', 'rb') as f:
41     optimal_weights = pickle.load(f)
42
43 predicted_homogen = []
44 for baris in range(len(NBPredicted_prob[0])):
45     jml=0
46     for model in range(len(NBPredicted_prob)):
47         jml = jml + (np.array(NBPredicted_prob[model])[baris]) * optimal_weights[model]
48     jml = jml/len(NBPredicted_prob)
49     if jml.argmax() == 0:
50         predicted_homogen.append(0)
51     elif jml.argmax() == 1:
52         predicted_homogen.append(1)
53     elif jml.argmax() == 2:
54         predicted_homogen.append(2)
55 predicted_homogen = np.array(predicted_homogen)
56
57 conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_homogen)

```



Lampiran 9. Script Program Homogeneous Ensemble Learning SVM

```

1 import numpy as np
2 import pandas as pd
3 import math
4 import seaborn as sns
5 import matplotlib.pyplot as plt
6 from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score
7 from sklearn.metrics import confusion_matrix
8 from scipy.sparse import csr_matrix
9 from sklearn.preprocessing import LabelEncoder
10 from MLAlgorithm.SVM import MulticlassSVM
11 import pickle
12
13 training_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_training_data_new.csv", encoding='utf-8')
14
15 #PROSES TRAINING
16 n_sample = 10
17 models = []
18 for i in range(n_sample):
19     sample_data = training_data.sample(frac=0.7, replace=False, random_state=i)
20
21     X_train = csr_matrix(sample_data.iloc[:, :-1]).toarray()
22     label_encoding = LabelEncoder()
23     y_train = label_encoding.fit_transform(sample_data['label'])
24
25     SVMModel = MulticlassSVM(C=1, tol=0.05, max_iter=1000, random_state=0, verbose=1)
26     SVMModel.fit(X_train, y_train)
27     models.append(SVMModel)
28
29 testing_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_testing_data_new.csv", encoding='utf-8')
30 X_test = csr_matrix(testing_data.iloc[:, :-1]).toarray()
31 y_test = label_encoding.fit_transform(testing_data['label'])
32
33 #PROSES TESTING
34 SVMPredicted_prob = []
35 for model in models:
36     prob = model.predict_proba(X_test)
37     SVMPredicted_prob.append(prob)
38
39 #PROSES VOTING DENGAN WEIGHTED SOFT VOTING
40 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/optimal_weight/homogeneousSVM_optimal_weights_new.pkl', 'rb') as f:
41     optimal_weights = pickle.load(f)
42
43 predicted_homogen = []
44 for baris in range(len(SVMPredicted_prob[0])):
45     jml=0
46     for model in range(len(SVMPredicted_prob)):
47         jml = jml + (np.array(SVMPredicted_prob[model])[baris]) * optimal_weights[model]
48     jml = jml/len(SVMPredicted_prob)
49     if jml.argmax() == 0:
50         predicted_homogen.append(0)
51     elif jml.argmax() == 1:
52         predicted_homogen.append(1)
53     elif jml.argmax() == 2:
54         predicted_homogen.append(2)
55 predicted_homogen = np.array(predicted_homogen)
56
57 conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_homogen)

```



Lampiran 10. Script Program Homogeneous Ensemble Learning LR

```

1 import numpy as np
2 import pandas as pd
3 import math
4 import seaborn as sns
5 import matplotlib.pyplot as plt
6 from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score
7 from sklearn.metrics import confusion_matrix
8 from scipy.sparse import csr_matrix
9 from sklearn.preprocessing import LabelEncoder
10 from MLAlgorithm.LogisticRegression import SoftmaxRegression
11 import pickle
12
13 training_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_training_data_new.csv", encoding='utf-8')
14
15 #PROSES TRAINING
16 n_sample = 10
17 models = []
18 for i in range(n_sample):
19     sample_data = training_data.sample(frac=0.7, replace=False, random_state=i)
20
21     X_train = csr_matrix(sample_data.iloc[:, :-1]).toarray()
22     label_encoding = LabelEncoder()
23     y_train = label_encoding.fit_transform(sample_data['label'])
24
25     LRModel = SoftmaxRegression(len(set(y_train)), len(X_train[0]))
26     LRModel.train_sgd(X_train, y_train)
27     models.append(LRModel)
28
29 testing_data = pd.read_csv("E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/vector_save/multiclass_testing_data_new.csv", encoding='utf-8')
30 X_test = csr_matrix(testing_data.iloc[:, :-1]).toarray()
31 y_test = label_encoding.fit_transform(testing_data['label'])
32
33 #PROSES TESTING
34 LRPredicted_prob = []
35 for model in models:
36     prob = model.predict_proba(X_test)
37     LRPredicted_prob.append(prob)
38
39 #PROSES VOTING DENGAN WEIGHTED SOFT VOTING
40 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
New/optimal_weight/homogeneousLR_optimal_weights_new.pkl', 'rb') as f:
41     optimal_weights = pickle.load(f)
42
43 predicted_homogen = []
44 for baris in range(len(LRPredicted_prob[0])):
45     jml=0
46     for model in range(len(LRPredicted_prob)):
47         jml = jml + (np.array(LRPredicted_prob[model])[baris]) * optimal_weights[model]
48     jml = jml/len(LRPredicted_prob)
49     if jml.argmax() == 0:
50         predicted_homogen.append(0)
51     elif jml.argmax() == 1:
52         predicted_homogen.append(1)
53     elif jml.argmax() == 2:
54         predicted_homogen.append(2)
55
56 predicted_homogen = np.array(predicted_homogen)
57
58 conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_homogen)

```



Lampiran 11. Potongan Script Algoritma PSO untuk Model *Heterogeneous*

```

1 from pyswarm import pso
2
3 def objective_function(weights):
4     # Split bobot untuk masing-masing model
5     nb_weight, svm_weight, lr_weight = weights
6     #Naive Bayes
7     NBModel = MultiNB()
8     NBModel.fit(X_train, y_train)
9     #SVM
10    SVMModel = MulticlassSVM(C=1, tol=0.05, max_iter=1000, random_state=0, verbose=1)
11    SVMModel.fit(X_train, y_train)
12    #Logistic Regression
13    LRModel = SoftmaxRegression(len(set(y_train)), len(X_train[0]))
14    LRModel.train_sgd(X_train, y_train)
15
16    #Naive Bayes
17    NBPredicted, NBPredicted_prob = NBModel.predict(X_test)
18    #SVM
19    SVMPredicted_prob = SVMModel.predict_proba(X_test)
20    #Logistic Regression
21    LRPredicted_prob = LRModel.predict_proba(X_test)
22
23    predicted_hetero = []
24    for baris in range(len(NBPredicted_prob)):
25        jml = ((np.array(NBPredicted_prob[baris]) * nb_weight) + (np.array(SVMPredicted_prob[baris]) * svm_weight)
26        + (np.array(LRPredicted_prob[baris]) * lr_weight)) / 3
27        if jml.argmax() == 0:
28            predicted_hetero.append(0)
29        elif jml.argmax() == 1:
30            predicted_hetero.append(1)
31        elif jml.argmax() == 2:
32            predicted_hetero.append(2)
33    predicted_hetero = np.array(predicted_hetero)
34
35    conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_hetero)
36
37    return -accuracy
38
39 # Batasan bobot (dalam rentang [0, 1])
40 lb = [0, 0, 0]
41 ub = [1, 1, 1]
42
43 # Jalankan algoritma PSO untuk mencari bobot optimal
44 best_weights, _ = pso(objective_function, lb, ub, swarmsize=10, maxiter=30)
45 optimal_nb_weight, optimal_svm_weight, optimal_lr_weight = best_weights
46 optimal_weights = [optimal_nb_weight, optimal_svm_weight, optimal_lr_weight]
47
48 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
    New/optimal_weight/heterogeneous_optimal_weights_new.pkl', 'wb') as f:
49     pickle.dump(optimal_weights, f)

```



Lampiran 12. Potongan Script Algoritma PSO Model *Homogeneous* NB

```

1 from pyswarm import pso
2
3 def objective_function(weights):
4     n_sample = 10
5     models = []
6     for i in range(n_sample):
7         sample_data = training_data.sample(frac=0.7, replace=False, random_state=i)
8
9         X_train = csr_matrix(sample_data.iloc[:, :-1]).toarray()
10        label_encoding = LabelEncoder()
11        y_train = label_encoding.fit_transform(sample_data['label'])
12
13        NBModel = MultinNB()
14        NBModel.fit(X_train, y_train)
15        models.append(NBModel)
16
17    NBPredicted_prob = []
18    NBAccuracy = []
19    for model in models:
20        predicted, prob = model.predict(X_test)
21        NBPredicted_prob.append(prob)
22
23    predicted_homogen = []
24    for baris in range(len(NBPredicted_prob[0])):
25        jml=0
26        for model in range(len(NBPredicted_prob)):
27            jml = jml + (np.array(NBPredicted_prob[model][baris]) * weights[model])
28        jml = jml/len(NBPredicted_prob)
29        if jml.argmax() == 0:
30            predicted_homogen.append(0)
31        elif jml.argmax() == 1:
32            predicted_homogen.append(1)
33        elif jml.argmax() == 2:
34            predicted_homogen.append(2)
35
36    predicted_homogen = np.array(predicted_homogen)
37
38    conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_homogen)
39
40    return -accuracy
41
42 # Batasan bobot (dalam rentang [0, 1])
43 jml = 10
44 lb = [0 for _ in range(10)]
45 ub = [1 for _ in range(10)]
46
47 # Jalankan algoritma PSO untuk mencari bobot optimal
48 best_weights, _ = pso(objective_function, lb, ub, swarmsize=10, maxiter=5)
49 optimal_weights = best_weights
50
51 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
52    New/optimal_weight/homogeneousNB_optimal_weights_new.pkl', 'wb') as f:
53     pickle.dump(optimal_weights, f)

```



Lampiran 13. Potongan Script Algoritma PSO Model *Homogeneous SVM*

```

1 from pyswarm import pso
2
3 def objective_function(weights):
4     n_sample = 10
5     models = []
6     for i in range(n_sample):
7         sample_data = training_data.sample(frac=0.7, replace=False, random_state=i)
8
9         X_train = csr_matrix(sample_data.iloc[:, :-1]).toarray()
10        label_encoding = LabelEncoder()
11        y_train = label_encoding.fit_transform(sample_data['label'])
12
13        SVMModel = MulticlassSVM(C=1, tol=0.05, max_iter=1000, random_state=0, verbose=1)
14        SVMModel.fit(X_train, y_train)
15        models.append(SVMModel)
16
17    SVMPredicted_prob = []
18    SVMAccuracy = []
19    for model in models:
20        prob = model.predict_proba(X_test)
21        SVMPredicted_prob.append(prob)
22
23    predicted_homogen = []
24    for baris in range(len(SVMPredicted_prob[0])):
25        jml=0
26        for model in range(len(SVMPredicted_prob)):
27            jml = jml + (np.array(SVMPredicted_prob[model])[baris] * weights[model])
28        jml = jml/len(SVMPredicted_prob)
29        if jml.argmax() == 0:
30            predicted_homogen.append(0)
31        elif jml.argmax() == 1:
32            predicted_homogen.append(1)
33        elif jml.argmax() == 2:
34            predicted_homogen.append(2)
35
36    predicted_homogen = np.array(predicted_homogen)
37
38    conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_homogen)
39
40    return -accuracy
41
42 # Batasan bobot (dalam rentang [0, 1])
43 jml = 10
44 lb = [0 for _ in range(10)]
45 ub = [1 for _ in range(10)]
46
47 # Jalankan algoritma PSO untuk mencari bobot optimal
48 best_weights, _ = pso(objective_function, lb, ub, swarmsize=10, maxiter=30)
49 optimal_weights = best_weights
50
51 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
52         New/optimal_weight/homogeneousSVM_optimal_weights_new.pkl', 'wb') as f:
53     pickle.dump(optimal_weights, f)

```



Lampiran 14. Potongan Script Algoritma PSO Model *Homogeneous* LR

```

1 from pyswarm import pso
2
3 def objective_function(weights):
4     n_sample = 10
5     models = []
6     for i in range(n_sample):
7         sample_data = training_data.sample(frac=0.7, replace=False, random_state=i)
8
9         X_train = csr_matrix(sample_data.iloc[:, :-1]).toarray()
10        label_encoding = LabelEncoder()
11        y_train = label_encoding.fit_transform(sample_data['label'])
12
13        LRModel = SoftmaxRegression(len(set(y_train)), len(X_train[0]))
14        LRModel.train_sgd(X_train, y_train)
15        models.append(LRModel)
16
17    LRPredicted_prob = []
18    LRAccuracy = []
19    for model in models:
20        prob = model.predict_proba(X_test)
21        LRPredicted_prob.append(prob)
22
23    predicted_homogen = []
24    for baris in range(len(LRPredicted_prob[0])):
25        jml=0
26        for model in range(len(LRPredicted_prob)):
27            jml = jml + (np.array(LRPredicted_prob[model])[baris]) * weights[model]
28        jml = jml/len(LRPredicted_prob)
29        if jml.argmax() == 0:
30            predicted_homogen.append(0)
31        elif jml.argmax() == 1:
32            predicted_homogen.append(1)
33        elif jml.argmax() == 2:
34            predicted_homogen.append(2)
35
36    predicted_homogen = np.array(predicted_homogen)
37
38    conf_matrix, accuracy, precision, recall, f1 = performance_test(y_test, predicted_homogen)
39
40    return -accuracy
41
42 # Batasan bobot (dalam rentang [0, 1])
43 jml = 10
44 lb = [0 for _ in range(10)]
45 ub = [1 for _ in range(10)]
46
47 # Jalankan algoritma PSO untuk mencari bobot optimal
48 best_weights, _ = pso(objective_function, lb, ub, swarmsize=10, maxiter=30)
49 optimal_weights = best_weights
50
51 with open('E:/Python/Penelitian Tesis/Sistem/Ready Sistem
52         New/optimal_weight/homogeneousLR_optimal_weights_new.pkl', 'wb') as f:
53     pickle.dump(optimal_weights, f)

```

