
ANALISIS METODE *FUZZY C-MEANS* UNTUK DETEKSI DAN KLASTERISASI PENYAKIT DIABETES DI RSUD PIDIE JAYA

Mutia Arifah¹, Muthmainnah², Dahlan Abdullah³

Universitas Malikussaleh, Aceh

email: ¹mutia.210180064@mhs.unimal.ac.id, ²muthmainnah@unimal.ac.id,
³dahlan@unimal.ac.id

Abstract: *Diabetes is a chronic disease with a steadily increasing number of sufferers every year, including in Pidie Jaya Regency. Early detection and clustering of sufferers are crucial to support effective disease management, personalized care, and more targeted allocation of health resources. This study aims to cluster sub-districts based on the number of diabetes sufferers using the Fuzzy C-Means (FCM) method. This method was chosen because it can flexibly group data by calculating the degree of membership of each data item within each cluster. The data used is secondary data on the number of diabetes sufferers from 12 sub-districts in Pidie Jaya Regency during the 2021-2024 period. The clustering process is carried out through several stages: parameter initialization, random creation of an initial partition matrix, calculation of cluster centers, degree of membership, and iterative evaluation of the objective function until convergence. The results show that Fuzzy C-Means successfully divides the data into two clusters: high and low clusters. Subdistricts such as Trienggadeng and Meureudu are in the high cluster, while Cubo and Jangka Buya are in the low cluster. The objective function value decreased significantly and reached convergence at the 13th iteration. The results of this study are expected to assist local governments and health agencies in developing more targeted policies based on the geographic distribution of diabetes sufferers in the Pidie Jaya region.*

Keyword: *Diabetes, Clustering, Fuzzy C-Means.*

Abstrak: Diabetes merupakan salah satu penyakit kronis yang jumlah penderitanya terus meningkat setiap tahun, termasuk di Kabupaten Pidie Jaya. Deteksi dini dan klasterisasi penderita sangat penting untuk mendukung manajemen penyakit yang efektif, personalisasi perawatan, serta alokasi sumber daya kesehatan yang lebih tepat sasaran. Penelitian ini bertujuan untuk mengklaster wilayah kecamatan berdasarkan jumlah penderita diabetes menggunakan metode *Fuzzy C-Means*. Metode ini dipilih karena mampu mengelompokkan data secara fleksibel dengan memperhitungkan derajat keanggotaan setiap data terhadap masing-masing klaster. Data yang digunakan berupa data sekunder jumlah penderita diabetes dari 12 kecamatan di Kabupaten Pidie Jaya selama periode 2021, 2022, 2023 hingga 2024. Proses klasterisasi dilakukan melalui beberapa tahap, yaitu inisialisasi parameter, pembuatan matriks partisi awal secara acak, perhitungan pusat klaster, derajat keanggotaan, dan evaluasi fungsi objektif secara iteratif hingga konvergensi. Hasil penelitian menunjukkan bahwa *Fuzzy C-Means* berhasil membagi data menjadi dua klaster, yaitu klaster tinggi dan rendah. Kecamatan seperti Trienggadeng dan Meureudu termasuk klaster tinggi, sedangkan Cubo dan Jangka Buya termasuk klaster rendah. Nilai fungsi objektif menurun signifikan dan mencapai konvergensi pada iterasi ke-13. Hasil penelitian ini diharapkan dapat membantu pemerintah daerah dan instansi kesehatan dalam membuat kebijakan yang lebih tepat sasaran berdasarkan distribusi geografis penderita diabetes di wilayah Pidie Jaya.

Kata kunci: *Diabetes, Clustering, Fuzzy C-Means.*

PENDAHULUAN

Diabetes adalah suatu kondisi kronis yang mempengaruhi sejumlah besar orang di seluruh dunia. Kondisi kronis medis yang ditandai dengan tingginya tingkat kadar glukosa (gula) dalam darah. Hal ini terjadi ketika tubuh tidak mampu memproduksi insulin untuk menghasilkan jumlah yang cukup atau tidak mampu menggunakan insulin secara efektif. Hormon-hormon insulin diproduksi oleh pankreas dan penting untuk mengendalikan kadar glukosa dalam darah (Satriatama et al., 2023). RSUD Pidie Jaya merupakan rumah sakit yang terletak di Kabupaten Pidie Jaya, Provinsi Aceh. Di Kabupaten Pidie Jaya, jumlah penderita diabetes terus mengalami peningkatan signifikan setiap tahunnya, dengan total kasus mencapai 1.704 orang hingga tahun 2024. Fenomena ini menunjukkan bahwa diabetes merupakan salah satu isu kesehatan yang perlu mendapat perhatian serius, khususnya dalam hal penyebaran dan perencanaan layanan kesehatan (Hana, 2020). Pada penelitian-penelitian sebelumnya, masih sangat terbatas kajian yang secara spesifik menganalisis penyebaran penderita diabetes berdasarkan wilayah, seperti kecamatan. Sebagian besar penelitian hanya berfokus pada aspek klinis, pola konsumsi, atau faktor risiko individu tanpa tanpa mengelompokkan wilayah berdasarkan karakteristik jumlah penderita (Argina, 2020).

Metode *Fuzzy K-Means* merupakan salah satu teknik klusterisasi data yang cocok diterapkan dalam konteks ini. Metode ini memungkinkan satu data memiliki derajat keanggotaan di lebih dari satu cluster, memberikan fleksibilitas dalam pengelompokan dan interpretasi data (Sulistiyawati & Supriyanto, 2021). Oleh karena itu, penelitian ini berupaya mengisi kekosongan dengan menerapkan metode *Fuzzy C-Means* untuk mengelompokkan 12 kecamatan di Kabupaten Pidie Jaya berdasarkan jumlah penderita diabetes selama periode 2021–

2024. Hasil pengelompokan ini diharapkan dapat memberikan kontribusi dalam perencanaan kebijakan kesehatan dan mendukung upaya penanggulangan diabetes secara lebih strategis dan diharapkan dapat membantu pihak rumah sakit dan pemerintah daerah dalam mengambil keputusan strategis terkait pencegahan dan penanganan penyakit diabetes secara lebih efektif. untuk mengantisipasi kepada masyarakat agar lebih berhati-hati pada wilayah dengan tingkat ancaman penderita yang lebih berat nantinya.

METODE

Penelitian ini menggunakan metode analisis deskriptif dengan pendekatan kuantitatif. Metode analisis deskriptif bertujuan untuk menjelaskan dan menggambarkan suatu fenomena melalui penyajian data dalam bentuk yang lebih pendekatan kuantitatif digunakan untuk mengumpulkan dan menganalisis data dalam bentuk angka atau data numerik (Waruwu, 2023).

Proses analisis data yang dilakukan menggunakan metode *Fuzzy C-Means*, *Fuzzy C Means* merupakan metode pengelompokan data yang didasarkan pada derajat keanggotaan, di mana prinsip dasar pembobotannya menggunakan konsep himpunan fuzzy. Setiap data tidak langsung ditempatkan secara mutlak dalam satu klaster, melainkan diberikan nilai kemungkinan untuk menjadi bagian dari setiap klaster yang ada (Hendrastuty, 2024). Dengan kata lain, satu data bisa memiliki potensi menjadi anggota lebih dari satu kelompok. Namun, nilai derajat keanggotaan yang paling tinggi menunjukkan kecenderungan terbesar suatu data untuk tergolong dalam klaster tertentu (Basalamah & Setyadi, 2023). Berikut adalah tahapan awal dalam proses pengelompokan data menggunakan metode *Fuzzy C-Means*:

1. Memasukkan Data: Masukkan data yang akan dikelompokkan X, dalam

- bentuk matriks berukuran $n \times m$ dimana n adalah jumlah data dan m ada.
2. Tentukan Parameter-parameter berikut:
 - a. Jumlah cluster = c
 - b. Pangkat = w
 - c. Maximum Iterasi (Max Iter)
 - d. Error terkecil yang diharapkan = ε
 - e. Fungsi objektif awal : $P_0 = 0$
 - f. Iterasi awal : $t = 1$

Langkah-Langkah ini merupakan persiapan sebelum memulai iterasi dalam metode *Fuzzy C-Means* untuk membentuk cluster yang optimal.
 3. Membangkitkan bilangan random μ_{ik} dengan $i = 1, 2, \dots, n; k = 1, 2, \dots, c$, yang menjadi elemen-elemen dalam matriks partisi awal U . Selanjutnya, hitung total nilai pada setiap kolom matriks tersebut:

$$Q_i = \sum_k^c \mu_{ik} \text{ dengan } j = 1, 2, \dots, n \quad (1)$$
 4. Menghitung pusat cluster ke- k : V_{kj}

$$V_{kj} = \frac{\sum_{i=1}^n (\mu_{ik})^w x_{ij}}{\sum_{i=1}^n (\mu_{ik})^w} \quad (2)$$
 5. Menghitung fungsi objektif pada iterasi ke- t :

$$P_t = \sum_{i=1}^n \sum_{k=1}^c ([\sum_{j=1}^m (x_{ij} - V_{kj})^2] (\mu_{ik})^w) \quad (3)$$
 6. Menghitung perubahan matriks partisi:

$$V_{kj} = \frac{[\sum_{i=1}^n (x_{ij} - V_{kj})^2] (\mu_{ik})^w}{[\sum_{k=1}^c \sum_{j=1}^m (x_{ij} - V_{kj})^2]^{\frac{w}{w-1}}} \quad (4)$$
 7. Periksa Kondisi berhenti:
 - a. Jika nilai $|P_t - P_{t-1}| \leq |MaxIter|$ telah mencapai batas maximum maka proses dihentikan.
 - b. Jika tidak, tingkatkan nilai iterasi $t = t + 1$ dan ulangi langkah ke-4 hasil akhir dari perhitungan menggunakan metode *Fuzzy C-Means* berupa kumpulan pusat cluster serta tingkat keanggotaan masing-masing data terhadap setiap cluster (Kurniawan et al., 2023).

HASIL DAN PEMBAHASAN

Pada penelitian ini menggunakan data jumlah penderita Diabetes dari 12 Kecamatan yang ada di Kabupaten Pidie Jaya selama periode 2021 hingga 2024. Data yang dianalisis berupa jumlah kasus penderita tiap tahun dari masing-masing kecamatan yang kemudian dilakukan proses klusterisasi menggunakan metode *Fuzzy C Means* dengan jumlah cluster sebanyak dua, yaitu kluster penderita diabetes tertinggi dan terendah. Berikut langkah-langkah proses perhitungan *Clustering* menggunakan metode *Fuzzy C-Means* untuk iterasi awal:

Analisis Data

1. Memasukkan data yang akan digunakan:

Tabel 1 Dataset Jumlah Penderita Diabetes

No	Kecamatan	2021	2022	2023	2024
1.	Bandar Baru	52	63	74	290
2.	Nyong	51	58	65	95
3.	Cubo	47	57	67	24
4.	Panteraja	65	75	91	73
5.	Trienggadeng	62	93	114	392
6.	Meureudu	112	144	166	258
7.	Meurah Dua	68	88	113	104
8.	Ulim	90	110	132	235

9.	Jangka Buya	56	72	97	35
10.	Bandar Dua	66	76	92	84
11.	Kuta Krueng	46	54	67	96
12.	Blang Kuta	60	62	64	81
13.	Jumlah	785	952	1142	1704

2. Menentukan parameter awal:

Jumlah Cluster :

2

Pangkat :

2

Maximum Iterasi : 100

Error terkecil yang diharapkan:
0.0001

Fungsi objektif awal : 0

3. Membangkitkan nilai bilangan random sebagai matriks partisi awal U: Jumlah derajat keanggotaan nilai bilangan random μ_{ik} untuk setiap data harus bernilai 1, Untuk menghitung pusat klaster nilai matriks partisi ini dipangkatkan, sehingga data dengan keanggotaan lebih tinggi memiliki pengaruh lebih besar dalam pembentukan pusat klaster.

C2	67.4	76.8024	93.0170	104.6
----	------	---------	---------	-------

5. Menghitung fungsi objektif menggunakan persamaan (3), nilai yang didapatkan adalah 47929.67. jika perubahan nilai fungsi objektif memenuhi batas error terkecil yang telah ditentukan maka proses iterasi dapat dihentikan namun jika belum mencapai error terkecil maka iterasi akan terus dilanjutkan hingga mencapai error terkecil yang diharapkan.
6. Selanjutnya adalah menghitung perubahan matriks partisi menggunakan persamaan (4), nilai yang didapatkan dalam perubahan matriks partisi ini adalah nilai yang akan digunakan sebagai pusat klaster pada iterasi berikutnya.

Tabel 2 Matrisk Partisi Awal

C1	C2
0.5	0.5
0.6	0.4
0.7	0.3
0.2	0.8
0.9	0.1
0.4	0.6
0.3	0.7
0.8	0.2
0.1	0.9
0.5	0.5
0.4	0.6
0.2	0.8

4. Menghitung pusat cluster dengan persamaan (2), maka didapatkan hasil sebagai berikut:

Tabel 3 Pusat Clusster 1

PUSAT CLUSTER				
C1	65.8393	83.3606	99.7484	190.563

Tabel 4 Perubahan Matriks Partisi

C1	C2
0.75874	0.24126
0.112374	0.887626
0.208738	0.791262
0.06718	0.93282
0.723913	0.276087
0.699226	0.300774
0.065114	0.934886
0.824688	0.175312
0.16871	0.83129
0.03595	0.96405
0.125708	0.874292
0.106416	0.893584

7. Melakukan pemeriksaan kondisi berhenti dengan menghitung selisih nilai fungsi objektif antara iterasi berjalan dengan iterasi pertama. Apabila selisih nilai tersebut masih lebih besar dari 0,0001 maka proses iterasi harus dilanjutkan namun

apabila selisih nilai sudah kurang dari atau sama dengan 0,0001 maka proses iterasi dihentikan. Berdasarkan hasil perhitungan, selisih nilai fungsi objektif masih lebih besar dari 0,0001 sehingga iterasi harus terus berlanjut.

Pada proses klasterisasi jumlah penderita diabetes di kabupaten Pidie Jaya untuk menentukan kecamatan dengan penderita diabetes tertinggi dan terendah berdasarkan hasil perhitungannya proses iterasi berhenti pada iterasi ke 13 dengan nilai error minimum 0,0001. Karena selisih fungsi objective iterasi ini dengan iterasi sebelumnya lebih kecil dari error terkecil yang diharapkan ($0,0001 < 0,0001$), maka iterasi dihentikan.

Tabel 5 Pusat Cluster pada Iterasi 13

PUSAT CLUSTER				
C1	80.805	102.032	121.113	278.321
C2	57.481	67.726	81.830	74.641

Tabel 6 Hasil Klasterisasi

Kecamatan	C1	C2	Cluster
Bandar Baru	0.000212	0.000022	1
Nyong	0.000025	0.001199	2
Cubo	0.000015	0.000332	2
Panteraja	0.000023	0.005096	2
Trieng gadeng	0.000327	0.000015	1
Meureudu	0.000170	0.000020	1
Meurah Dua	0.000032	0.000425	2
Ulim	0.000467	0.000032	1
Jangka Buya	0.000016	0.000032	2
Bandar Dua	0.000025	0.003012	2
Kuta Krueng	0.000025	0.001004	2
Blang Kuta	0.000023	0.002516	2

Dari hasil klasterisasi terhadap 12 kecamatan menunjukkan sebanyak 4 kecamatan termasuk dalam Cluster 1 dengan jumlah penderita diabetes tertinggi, sedangkan 8 kecamatan lainnya termasuk dalam Cluster 2 dengan jumlah penderita diabetes terendah. Hasil ini dapat digunakan untuk penanganan kesehatan

yang lebih terfokus di Kecamatan penderita tertinggi.

Analisa Python Fuzzy C-Means

Instalasi Library Fuzzy

Langkah awal yang dilakukan pada proses Analisa Python adalah menginstal Library Fuzzy yang dapat dilakukan dengan perintah:



Gambar 1 Instalasi Library Fuzzy

Menampilkan Library yang dibutuhkan Langkah awal dalam proses analisis data adalah mengimpor library yang diperlukan untuk membaca data set yang akan digunakan.



Gambar 2 Menampilkan Library

Menampilkan Data yang dibutuhkan Perintah pada gambar 3 dibawah akan menampilkan 12 baris dataset yang telah dimuat tujuannya untuk memastikan bahwa data telah berhasil di muat.

```
Out[2]:
```

	KECAMATAN	2021	2022	2023	2024
0	Bandar Baru	52	53	74	290
1	Nyong	51	58	65	95
2	Cubo	47	57	67	24
3	Panteraja	65	75	91	73
4	Trieng gadeng	62	93	114	329
5	Meureudu	122	144	166	258
6	Meurah Dua	68	88	113	104
7	Ulim	90	110	132	235
8	Jangka Buya	56	72	97	35
9	Bandar Dua	66	78	92	101
10	Kuta Krueng	46	54	67	96
11	Blang Kuta	60	62	64	81

Gambar 3 Menampilkan Dataset

Penerapan Metode Fuzzy C-Means

Langkah awal dalam proses klasterisasi adalah melakukan inisialisasi pembacaan dan ekstraksi data menggunakan perintah pada gambar 4.

```
In [8]: import skfuzzy as fuzz
```

```
r = fuzz.cluster.cmeans(
    data=X.T,
    c=2,
    m=2.0,
    error=0.0001,
    maxiter=100,
    init=None
)
```

Gambar 4 Inisialisasi Cluster

Hasil *Clustering*

Setelah proses inisialisasi selesai maka hasil yang diperoleh disimpan dalam variable *r*. untuk mengakses hasil tertentu dari proses klasterisasi, cukup memanggil elemen array tersebut menggunakan indeks *r*. berikut adalah contoh bentuk hasil yang ditampilkan

[illegible]

Gambar 5 Menampilkan Array Klasterisasi

Menampilkan Cluster setiap Data

Langkah selanjutnya setelah mendapatkan hasil perhitungan adalah menentukan setiap baris data dapat diketahui termasuk ke dalam Cluster mana berdasarkan hasil. Hal ini dilakukan menggunakan perintah pada gambar 6. Dimana berdasarkan gambar setelah proses klasterisasi selesai dilakukan, seluruh kecamatan di Kabupaten Pidie Jaya terbagi menjadi dua klaster, yaitu Cluster 0 yang terdiri dari kecamatan dengan jumlah penderita Diabetes terendah dan Cluster 1 mencakup kecamatan dengan jumlah penderita Diabetes tertinggi.

```

In [28]: cluster_membership = np.argmax(u, axis=0)

In [29]: base['Cluster'] = cluster_membership

In [30]: print(base)

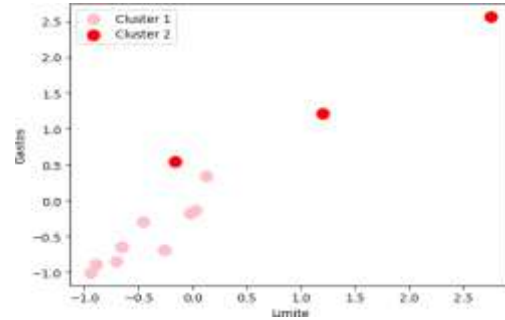
```

	KECAMATAN	2021	2022	2023	2024	Cluster
0	Bandar Baru	52	63	74	298	1
1	Nyong	11	58	65	95	1
2	Cube	47	57	67	24	1
3	Panteraja	65	73	91	72	1
4	Trianggadeng	53	93	114	129	0
5	Necudu	122	144	166	250	0
6	Muarah Dua	68	88	113	204	0
7	Ulin	98	118	152	235	0
8	Jangka Buaya	50	72	97	35	1
9	Bandar Dua	66	76	92	88	1
10	Kuta Trueng	46	54	67	90	1
11	Blang Kuta	60	62	64	81	1

Gambar 6 Hasil Setiap Cluster

Menampilkan Visualisasi Scatter Plot Klasterisasi

Langkah selanjutnya setelah proses klusterisasi selesai adalah menampilkan Visualisasi hasil kluster menggunakan Scatter Plot untuk menganalisis secara Visual.

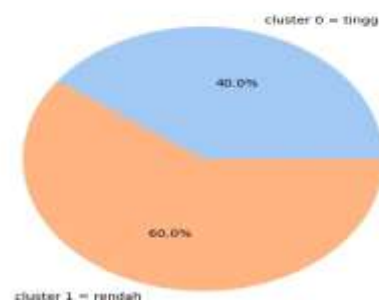


Gambar 7 Visualisasi Scatter Plot

Titik-titik yang berwarna pink mewakili wilayah yang tergolong kedalam Cluster 1 yang merupakan daerah dengan jumlah penderita Diabetes terendah sedangkan titik berwarna merah merupakan wilayah kecamatan yang masuk ke Cluster 2 merupakan daerah dengan jumlah penderita Diabetes tertinggi. Visualisasi ini membantu dalam analisis spasial dan pengambilan keputusan, wilayah berwarna pink perlu mendapatkan perhatian lebih karena memiliki konsentrasi penderita Diabetes yang lebih tinggi.

Visualisasi Presentase Klasterisasi Penderita Diabetes

Langkah yang selanjutnya adalah menampilkan Visualisasi Persentase dari hasil klasterisasi dalam bentuk diagram pie yang bertujuan untuk memberikan gambaran proporsi jumlah penderita diabetes dalam masing-masing klaster.



Gambar 8 Visualisasi Presentase *Clustering*

Berdasarkan Diagram Pie pada gambar 8 menunjukkan hasil visualisasi yang mengelompokkan penderita diabetes kedalam 2 cluster. 60% Kecamatan termasuk kedalam Cluster terendah dan 40% Kecamatan termasuk kedalam Cluster jumlah penderita Diabetes tertinggi.

SIMPULAN

Penerapan metode *Fuzzy C-Means* dalam deteksi dan klasterisasi penderita diabetes di Pidie Jaya berhasil mengelompokkan 12 kecamatan menjadi dua klaster berdasarkan jumlah kasus dari tahun 2021 hingga 2024. Setiap kecamatan memperoleh derajat keanggotaan yang memungkinkan analisis lebih fleksibel. Metode ini terbukti efektif dalam mengidentifikasi pola distribusi penderita, ditunjukkan oleh penurunan signifikan nilai fungsi objektif hingga konvergen pada iterasi ke-13. Hasil klasterisasi memberikan informasi strategis bagi pemerintah daerah dan rumah sakit dalam menetapkan prioritas penanganan dan perencanaan layanan kesehatan, dengan kecamatan Trienggadeng dan Meureudu sebagai wilayah dengan tingkat kasus tertinggi.

DAFTAR PUSTAKA

- Argina, A. M. (2020). Penerapan Metode Klasifikasi K-Nearest Neighbor pada Dataset Penderita Penyakit Diabetes. *Indonesian Journal of Data and Science*, 1(2), 29–33. <https://doi.org/10.33096/ijodas.v1i2.11>
- Basalamah, A. T., & Setyadi, R. (2023). Penerapan Algoritma K-Means Clustering Pada Tingkat Penyelesaian Pendidikan Di Provinsi Indonesia. *Jurnal Informatika Dan Teknologi Komputer*, 4(2), 114–121. <https://ejurnalunsam.id/index.php/jic om/>
- Hana, F. M. (2020). Klasifikasi Penderita Penyakit Diabetes Menggunakan Algoritma Decision Tree C4.5. *Jurnal SISKOM-KB (Sistem Komputer Dan Kecerdasan Buatan)*.
- Hendrastuty, N. (2024). Penerapan Data Mining Menggunakan Algoritma K-Means Clustering Dalam Evaluasi Hasil Pembelajaran Siswa. *Jurnal Ilmiah Informatika Dan Ilmu Komputer (Jima-Ilkom)*, 3(1), 46–56. <https://doi.org/10.58602/jima-ilkom.v3i1.26>
- Kurniawan, S., Siregar, A. M., & Novita, H. Y. (2023). Penerapan Algoritma K-Means dan Fuzzy C-Means Dalam Mengelompokkan Prestasi Siswa Berdasarkan Nilai Akademik. *Scientific Student Journal for Information, Technology and Science*, IV(1), 73–81.
- Satriatama, A. E., Wibowo, A. P., Arnold, I. G. N., Pratama, R. B., Masyhuda, T. A., Agusti, Y. A., Purwanti, E., & Werdiningsih, I. (2023). Analisis Klaster Data Pasien Diabetes untuk Identifikasi Pola dan Karakteristik Pasien. *Jurnal Teknologi Dan Sistem Informasi Bisnis*, 5(3), 172–182. <https://doi.org/10.47233/jteksis.v5i3.828>
- Sulistiyawati, A., & Supriyanto, E. (2021). Implementasi Algoritma K-means Clustering dalam Penentuan Siswa Kelas Unggulan. *Jurnal Tekno Kompak*, 15(2), 25. <https://doi.org/10.33365/jtk.v15i2.1162>
- Waruwu, M. (2023). Pendekatan Penelitian Pendidikan: Metode Penelitian Kualitatif, Metode Penelitian Kuantitatif dan Metode Penelitian Kombinasi (Mixed Method). *Jurnal Pendidikan Tambusai*, 7(1), 2896–2910.