

BAB III

OBJEK DAN METODE PENELITIAN

3.1. Objek Penelitian

Dalam bukunya, Sugiyono (2019a, p. 38) menyebutkan objek penelitian merupakan komponen yang berkaitan dengan sebuah penelitian dan menjadi fokus untuk mendapatkan jawaban atau solusi dari permasalahan yang diteliti. Objek penelitian dapat berupa atribut, sifat, atau nilai dari seseorang. Suatu objek atau kegiatan yang memiliki variasi tertentu yang ditetapkan oleh peneliti untuk dikaji dan kemudian diambil kesimpulannya.

Objek dalam penelitian ini yaitu menganalisis tingkat kemiskinan di Jawa Tengah pada Tahun 2023 dengan mempertimbangkan berbagai dimensi. Adapun subjek dalam penelitian ini adalah Provinsi Jawa Tengah yang terdiri dari 35 kabupaten/kota. Pemilihan Jawa Tengah didasarkan pada beberapa pertimbangan di mana meskipun tingkat persentase kemiskinan Jawa Tengah kedua sepulau Jawa setelah DI Yogyakarta pada tahun 2023, namun luas wilayah serta tingginya indeks kedalaman kemiskinan, Jawa Tengah memiliki nilai yang lebih tinggi dari DI Yogyakarta pada tahun 2023.

3.2. Metode Penelitian

Metode penelitian adalah serangkaian proses sistematis dalam mencari kebenaran dalam sebuah penelitian. Proses ini dimulai dari pemikiran kemudian mengarah pada pembentukan rumusan masalah. Dalam prosesnya, penelitian didukung oleh kajian dan persepsi dari penelitian-penelitian sebelumnya. Data yang

diperoleh kemudian diolah dan dianalisis hingga akhirnya menjadi suatu kesimpulan (Sahir, 2021, p. 1).

3.2.1. Jenis Penelitian

Jenis penelitian yang digunakan pada penelitian ini adalah metode kuantitatif. Metode kuantitatif merupakan pendekatan penelitian yang menggunakan alat statistik untuk mengolah data, sehingga baik data yang dikumpulkan maupun hasil analisisnya berupa angka (Sahir, 2021, p. 13). Metode kuantitatif menurut Sugiyono (2019b, p. 16) diartikan sebagai pendekatan yang berbasis pada filsafat positivisme yang digunakan untuk penelitian terhadap populasi atau sampel tertentu. Dalam metode ini, pengumpulan data dilakukan menggunakan instrumen penelitian dan analisis data dilakukan secara statistik dengan tujuan untuk menguji hipotesis yang telah dirumuskan sebelumnya.

Dalam penelitian kuantitatif komponen masalah dibagi menjadi beberapa variabel, di mana setiap variabel diberikan simbol yang berbeda sesuai dengan kebutuhan penelitian atau permasalahan yang akan dikaji oleh peneliti (Sahir, 2021, p. 14). Hal ini sejalan dengan tujuan dari penelitian ini yaitu untuk menganalisis aspek mana saja yang dapat memengaruhi tingkat kemiskinan di Jawa Tengah pada tahun 2023 serta bagaimana aspek-aspek tersebut dapat dikelompokkan dalam suatu klaster.

3.2.2. Operasionalisasi Variabel

Sugiyono (2019b, p. 68) mendeskripsikan variabel penelitian sebagai atribut, sifat, atau nilai yang berasal dari orang, objek, atau kegiatan yang memiliki variasi tertentu yang telah ditentukan oleh peneliti untuk dikaji lebih lanjut dan

kemudian diambil kesimpulannya. Berdasarkan judul penelitian yang dipilih, adapun operasionalisasi variabel pada penelitian ini diuraikan pada tabel 3.1.

Tabel 3.1 Operasionalisasi Variabel

Dimensi	Variabel	Definisi Variabel	Simbol	Satuan
(1)	(2)	(3)	(4)	(5)
Sosial - Ekonomi	Persentase Penduduk Miskin	Seberapa persen penduduk yang hidup di bawah garis kemiskinan	X1	Persen
	Indeks Kedalaman Kemiskinan	Seberapa jauh jarak rata-rata pengeluaran penduduk miskin dibandingkan dengan garis kemiskinan	X2	Rasio
	Indeks Keparahan Kemiskinan	Seberapa jauh distribusi pengeluaran di kalangan masyarakat miskin	X3	Rasio
	Persentase Pengeluaran per kapita Sebulan untuk Makanan	Seberapa besar porsi pengeluaran untuk konsumsi makanan sebulan dibandingkan dengan total pengeluaran per kapita sebulan (makanan dan non-makanan)	X4	Persen
	Rata-rata Upah/Gaji Bersih Sebulan Pekerja Formal berdasarkan Lapangan Pekerjaan Utama	Rata-rata penghasilan bersih dalam sebulan penduduk yang memiliki status pekerjaan formal	X5	Rp
	Rata-rata Upah/Gaji Bersih Sebulan Pekerja Informal berdasarkan Lapangan Pekerjaan Utama	Rata-rata penghasilan bersih dalam sebulan penduduk yang memiliki status pekerjaan informal	X6	Rp
Pendi- dikan	Rata-Rata Lama Sekolah	Rata-rata jumlah tahun penduduk usia 25 tahun ke atas yang telah	X7	Tahun

(1)	(2)	(3)	(4)	(5)
		menyelesaikan pendidikan formal		
	Harapan Lama Sekolah	Rata-rata jumlah tahun sekolah formal yang diharapkan akan ditempuh oleh penduduk	X8	Tahun
Kesehatan	Usia Harapan Hidup	Rata-rata jumlah tahun yang diperkirakan akan dijalani oleh seseorang	X9	Tahun
	Persentase Rumah Tangga Terhadap Akses Sumber Air Minum Bersih	Seberapa persen rumah tangga yang mempunyai akses terhadap sumber air minum bersih	X10	Persen
Infra- struktur/ Standar Hidup Layak	Persentase Rumah Tangga yang Memiliki Akses Terhadap Sanitasi Layak	Seberapa persen rumah tangga yang mempunyai akses dengan sanitasi layak	X11	Persen
	Persentase Rumah Tangga dengan Sumber Penerangan Listrik PLN	Seberapa persen rumah tangga yang sumber penerangannya menggunakan listrik PLN	X12	Persen
	Persentase Rumah Tangga dengan Jenis Lantai Terluas Bukan Tanah	Seberapa persen rumah tangga yang menggunakan jenis lantai terluasnya bukan tanah	X13	Persen

3.2.3. Teknik Pengumpulan Data

3.2.3.1. Jenis dan Sumber Data

Jenis dan sumber data yang digunakan dalam penelitian ini yaitu menggunakan data sekunder. Data sekunder adalah data yang diperoleh secara tidak langsung melalui proses membaca, mempelajari, dan memahami berbagai media, seperti literatur, buku-buku, dan dokumen. Data sekunder merupakan data yang

didapatkan oleh peneliti secara tidak langsung dari subjek penelitian, melainkan diperoleh melalui perantara atau sumber lain (Sugiyono, 2019b, p. 194).

Data dalam penelitian ini bersumber dari Badan Pusat Statistik Provinsi Jawa Tengah. Data yang digunakan merupakan data *cross-section* yang mencakup 35 kabupaten/kota di Jawa Tengah pada tahun 2023.

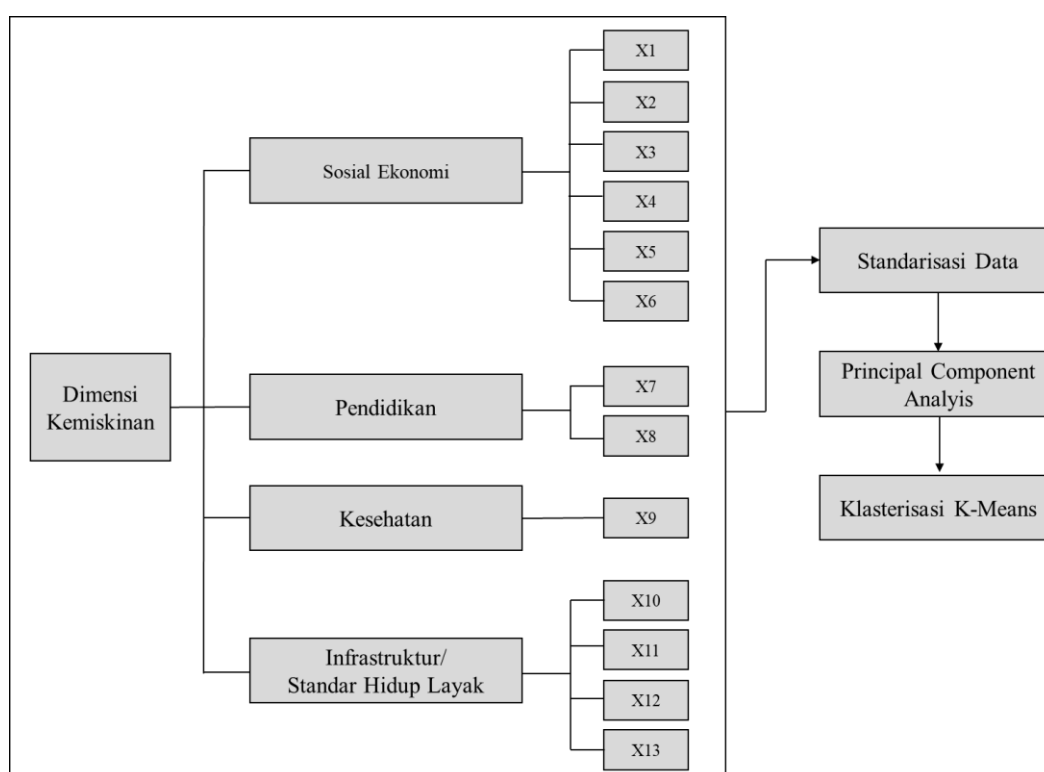
3.2.3.2. Populasi dan Sampel Penelitian

Populasi merupakan suatu wilayah umum yang mencakup objek atau subjek dengan jumlah dan karakteristik tertentu, yang telah ditetapkan oleh peneliti untuk dikaji dan diambil kesimpulannya (Sugiyono, 2019b, p. 126). Populasi tidak hanya terdiri dari manusia saja, tetapi juga mencakup objek dan benda-benda alam lainnya. Populasi bukan sekedar jumlah yang ada pada objek/subjek yang diteliti, tetapi mencakup seluruh ciri, karakteristik dan sifat yang dimiliki oleh subjek atau objek tersebut. Adapun populasi dalam penelitian ini yaitu kabupaten/kota di Provinsi Jawa Tengah yang berjumlah 35 wilayah dengan terdiri dari 29 kabupaten dan 6 kota.

Sampel merupakan bagian dari populasi yang dipilih oleh peneliti untuk dikaji dan diambil kesimpulannya (Sugiyono, 2019b, p. 127). Dalam penelitian ini penulis menggunakan teknik sampling jenuh. Sampel jenuh merupakan kondisi di mana penambahan jumlah sampel tidak lagi memberikan kontribusi signifikan terhadap penelitian (Sugiyono, 2019b, p. 133). Dalam penelitian ini seluruh populasi yaitu Jawa Tengah menjadi subjek penelitian. Hal ini dipilih untuk mendapatkan gambaran komprehensif tentang kondisi kemiskinan di seluruh wilayah Jawa Tengah pada tahun 2023.

3.2.4. Model Penelitian

Model penelitian adalah suatu model yang menggambarkan keterkaitan antar variabel dari kerangka berpikir yang didasarkan pada teori tertentu. Kerangka berpikir tersebut tidak hanya menunjukkan bagaimana variabel-variabel penelitian saling berhubungan, tetapi juga mencerminkan permasalahan yang ingin dipecahkan dalam penelitian. (Sugiyono, 2019a, p. 61). Oleh karena itu, model penelitian yang digunakan dalam judul ini digambarkan pada gambar 3.1.



Gambar 3.1 Model Penelitian

3.2.5. Teknis Analisis Data

Tahap pertama dalam penelitian ini yaitu melakukan pengolahan data dengan melakukan standarisasi data. Tahap pertama sebelum penentuan aspek atau komponen utama dan Klasterisasi adalah melakukan standarisasi data menggunakan *Z-Score Standardization*. Standarisasi *Z-score* mengubah data agar

berada pada rentang nilai yang sama (Yafi et al., 2023). Proses ini penting dilakukan karena variabel-variabel yang digunakan dalam penelitian ini memiliki satuan yang berbeda (Aldawiyah et al., 2024). Standarisasi data membantu menghilangkan pengaruh perbedaan satuan pengukuran dan membuat semua variabel memiliki kontribusi yang setara dalam analisis (atlan, 2023).

3.2.5.1. *Principal Component Analysis*

Principal Component Analysis (PCA) merupakan metode analisis yang bertujuan untuk mereduksi data dengan cara menerangkan struktur varians dan kovarians melalui kombinasi linier sesedikit mungkin dari variabel aslinya (Rusli & Setyawan, 2022). PCA mencoba mengestrak struktur dari suatu set data berukuran besar menjadi dimensi yang lebih kecil tanpa kehilangan informasi penting yang dapat menggambarkan keseluruhan data (Yafi et al., 2023).

Dalam penelitian ini, PCA digunakan untuk menemukan hubungan antar variabel yang independent, lalu mengelompokkannya menjadi faktor-faktor atau komponen-komponen dengan jumlah yang lebih sedikit dari jumlah variabel awal. Dengan demikian PCA dapat membantu menyederhanakan data multivariat menjadi dimensi yang lebih ringkas namun tetap mampu menjelaskan keragaman data secara komprehensif (Syahrani et al., 2021). Dalam bentuk matematis, katakanlah Y merupakan kombinasi linier dari variabel-variabel $X_1, X_2, X_3, \dots$ yang dapat dinyatakan sebagai berikut:

$$Y_p = W_{1p}X_1 + W_{2p}X_2 + W_{3p}X_3 + \dots + \beta_{pp}X_p$$

Keterangan:

Y = perkiraan faktor/komponen (kombinasi linier dari variabel X)

W = bobot atau koefisien nilai variabel

p = banyaknya variabel

Adapun dalam penelitian ini proses analisis PCA dilakukan dengan bantuan *software* SPSS di mana dalam tahapannya meliputi:

1. Pembentukan Matriks Korelasi
 - a. Uji *Kaiser-Meyer-Olkin* (KMO) dan *Bartlett's Test of Sphericity*

Uji *Kaiser-Meyer-Olkin* (KMO) adalah pengujian untuk menentukan ukuran kecocokan data dalam analisis faktor. Uji ini mengukur kecukupan pengambilan sampel untuk setiap variabel dalam model. Uji bertujuan untuk mengidentifikasi hubungan antar variabel. Pada tahap ini dilakukan pengujian kelayakan data menggunakan *Kaiser-Meyer-Olkin* (KMO) dan *Bartlett's Test of Sphericity*. Menurut Syahrani (2021) KMO adalah alat yang digunakan untuk mengukur seberapa kuat hubungan atau korelasi antar variabel. Nilai KMO ini dapat menjadi indikator apakah data tersebut layak atau tidak untuk dilakukan analisis faktor. Di samping itu, uji *Bartlett's Test of Sphericity* merupakan uji statistik yang bertujuan untuk mendeteksi ada atau tidaknya korelasi antar variabel dalam data. Uji ini menjadi pelengkap dari pengujian KMO. Adapun kriteria kelayakan analisis faktor berdasarkan hasil uji KMO dan *Bartlett's Test of Sphericity* adalah nilai KMO harus lebih dari 0,5 dan nilai signifikansi uji *Bartlett's Test of Sphericity* harus kurang dari 0,05. Jika kedua kriteria ini terpenuhi, maka data dinyatakan memenuhi syarat untuk dilakukan analisis faktor.

Adapun kriteria KMO dalam penelitiannya Kaiser (1974) adalah:

0,9 – 1.00 : data sangat baik

- 0,8 – 0.89 : data baik
- 0.7 – 0.79 : data cukup baik
- 0,6 – 0.69 : data kurang baik/biasa
- 0.5 – 0.59 : data sangat kurang
- 0.00 – 0.49 : data tidak layak/tidak dapat diterima

b. Uji *Measure of Sampling Adequacy* (MSA)

Selain menggunakan nilai *Kaiser-Meyer-Olkin* (KMO) dan *Bartlett's Test of Sphericity* sebagai uji kelayakan suatu data, dilakukan juga uji *Measure of Sampling Adequacy* (MSA) untuk melihat layak tidaknya suatu variabel dalam analisis faktor. Adapun kriteria yang digunakan dalam MSA adalah kriteria yang sama digunakan dalam KMO. Apabila nilai MSA suatu variabel kurang dari 0,5 maka variabel tersebut dikeluarkan dan dilakukan pengujian ulang dari tahap awal (Aldawiyah et al., 2024).

2. Ekstraksi Komponen

a. Uji *Communalities*

Setelah melihat layak tidaknya suatu variabel terhadap penelitian yang dicerminkan melalui nilai MSA. Maka langkah selanjutnya adalah dengan melihat nilai *communalities*. Uji *communalities* merupakan pengukuran untuk mengetahui jumlah varians atau keragaman dari suatu variabel yang dapat dijelaskan atau diterangkan oleh faktor-faktor yang terbentuk dalam analisis. Nilai *communalities* lebih besar dari 0,5 perlu dipenuhi agar analisis dapat dilanjutkan dan dilakukan dengan baik (Syahrani et al., 2021).

b. Menentukan banyaknya nilai komponen utama

Terdapat tiga pendekatan yang dapat digunakan untuk menentukan berapa banyak faktor atau komponen yang terbentuk. Rusli dan Setyawan (2022) menyebutkan dengan melihat nilai *eigen* dari masing-masing komponen. Hanya komponen dengan nilai *eigen* di atas 1 yang akan dipertahankan, karena dianggap mampu menjelaskan varians data secara memadai. Selanjutnya jumlah faktor yang terbentuk juga harus dapat menjelaskan setidaknya 66% dari total varians data secara keseluruhan. Pendekatan yang terakhir adalah dengan melihat grafik *scree plot* yang menampilkan nilai *eigen* untuk masing-masing komponen.

3. Rotasi Komponen

Tahap terakhir dalam PCA yaitu melakukan rotasi komponen untuk mengelompokkan variabel-variabel yang saling berkorelasi kuat ke dalam faktor atau komponen yang sama. Hal ini dilakukan dengan melihat nilai *loading factor* dari masing-masing variabel terhadap komponen-komponen yang terbentuk. Jika suatu variabel memiliki nilai *loading factor* yang lebih besar atau sama dengan 0,5 maka dapat dikatakan bahwa variabel tersebut memiliki korelasi yang tinggi terhadap komponen yang dibentuknya (Rusli & Setyawan, 2022)

Untuk mengidentifikasi keterkaitan antara variabel dengan komponen baru yang terbentuk digunakan rotasi *Varimax*. Rotasi *varimax* merupakan salah satu metode yang digunakan dalam analisis faktor atau analisis komponen utama. Interpretasi komponen dilakukan berdasarkan variabel-variabel yang memiliki *loading factor* tinggi pada masing-masing komponen (Aldawiyah et al., 2024).

3.2.5.2. Klasterisasi K-Means

Setelah dimensi data direduksi melalui PCA, langkah selanjutnya adalah melakukan pengelompokan atau klasterisasi. Klasterisasi adalah proses pengelompokan objek-objek dari suatu kumpulan data menjadi beberapa kelompok (klaster) yang homogen. Tujuan utama klastering adalah mengelompokkan sejumlah data atau objek ke dalam beberapa klaster sehingga objek-objek dalam satu klaster akan memiliki kemiripan yang tinggi satu sama lain (Bahauddin et al., 2021).

Salah satu cara dalam melakukan klasterisasi adalah dengan menggunakan metode K-Means. K-Means clustering adalah salah satu metode pengelompokan data (*clustering*) yang bersifat non-hierarki (Erda et al., 2023). Metode ini dipilih karena kemampuannya dalam mengelompokkan subjeknya berdasarkan kemiripan yang tinggi dengan perbedaan antara klaster yang jelas (Nurohmah et al., 2023).

Metode ini akan membantu mengidentifikasi pola spasial kemiskinan di Jawa Tengah pada tahun 2023 melalui klasterisasi kabupaten/kota yang memiliki karakteristik serupa. Adapun seluruh proses K-Means clustering ini dilakukan menggunakan *software* R dengan data yang terstandarisasi di mana tahapan dalam melakukan analisis K-Means yaitu:

1. Penentuan Jumlah Klaster Optimal

Salah satu kelemahan dari metode klasterisasi K-Means adalah prosedurnya yang menentukan pusat klaster (*centroid*) secara acak pada awal proses. Jika jumlah klaster awal yang ditetapkan tidak sesuai, maka pusat klaster yang terbentuk dapat

berubah selama iterasi. Perubahan ini dapat mengakibatkan kelompok data menjadi tidak optimal (Basri et al., 2023).

Dalam penelitian ini, penentuan jumlah klaster optimal dilakukan dengan menggunakan *Silhouette Analysis*, yang mengukur seberapa baik kesesuaian suatu objek dengan klasternya sendiri dibandingkan dengan klaster terdekat lainnya. Metode ini menghasilkan nilai yang berkisar antara -1 hingga 1, di mana nilai positif mengindikasikan tingkat kesesuaian objek dengan klasternya, sedangkan nilai negatif menunjukkan bahwa objek tersebut mungkin lebih cocok berada di klaster tetangga. Penentuan jumlah klaster optimal dapat dilihat berdasarkan nilai *silhouette coefficient* rata-rata tertinggi yang ditunjukkan oleh garis tertinggi pada plot yang dihasilkan (Erda et al., 2023; Galela, 2023).

2. Proses Klasterisasi

Proses pengelompokan dalam *K-Means Clustering* didasarkan pada penentuan jumlah kelompok awal (k) dan penentuan nilai *centroid* awal untuk masing-masing kelompok tersebut. Nilai k dipilih sebagai titik pusat awal akan dihitung menggunakan rumus jarak *Euclidean*. Selanjutnya *K-Means* akan mengelompokkan data-data tersebut ke dalam k klaster berdasarkan kedekatan jaraknya dengan *centroid* masing-masing klaster. Analisis *K-Means* menggunakan proses iteratif untuk mendapatkan hasil akhir pengelompokan data. Proses ini akan terus berulang hingga didapatkan hasil pengelompokan yang optimal, yaitu ketika tidak ada lagi perubahan yang signifikan pada anggota klaster (Erda et al., 2023).