

BAB III

METODE PENELITIAN

A. Jenis Penelitian

Dalam penelitian ini, peneliti menggunakan jenis penelitian kuantitatif. Jenis penelitian kuantitatif ini dipilih karena penulis ingin meneliti suatu populasi atau sampel yang bertujuan untuk membuktikan hipotesis yang ditetapkan dalam penelitian.⁵⁸

Penelitian ini bersifat kuantitatif dengan analisis deskriptif, karena dalam penelitian ini akan menjelaskan gambaran tentang Pelayanan Kesehatan dan konsumsi rumah tangga masyarakat terhadap Indeks pembangunan manusia per provinsi Indonesia dalam 2019 - 2023. Deskriptif ini menjelaskan gejala-gejala, fakta-fakta, atau kejadian-kejadian secara sistematis dan akurat mengenai sifat-sifat populasi atau daerah tertentu.

B. Sumber Penelitian

Sumber data yang diperoleh pada penelitian ini berasal dari data sekunder. Data sekunder adalah sumber atau informasi yang tidak langsung memberikan data kepada pengumpul data. Biasanya data ini bersifat dokumentasi, catatan pribadi, buku literature, data statistik atau arsip dinas.⁵⁹ Adapun data sekunder yang digunakan dalam penelitian ini ialah data yang didapatkan dari BPS Indonesia.

⁵⁸ Dameria Sinaga, *Statistik Dasar, Jurnal Penelitian Pendidikan Guru Sekolah Dasar*, vol. 6, 2016.

⁵⁹ Sinaga.

C. Populasi dan Sampel

Populasi adalah keseluruhan objek atau subjek yang menjadi fokus penelitian yang ditetapkan oleh peneliti untuk dipelajari dan kemudian ditarik kesimpulannya. Populasi yang diambil dalam penelitian ini adalah seluruh laporan data IPM per Provinsi Indonesia, Pelayanan Kesehatan Per Provinsi dan tingkat konsumsi rumah tangga per Provinsi di Indonesia. Sampel adalah bagian dari sejumlah karakteristik yang dimiliki oleh populasi yang digunakan untuk penelitian. Dalam hal ini penulis menggunakan sampel lima tahun yaitu dari tahun 2019-2023.⁶⁰

Metode yang digunakan dalam pengambilan sampel penelitian ini adalah Purposive Sampling yaitu teknik penentuan sampel dengan pertimbangan atau kriteria-kriteria tertentu.⁶¹ Adapun alasan pemilihan sampel dalam penelitian ini adalah data yang diterbitkan dari BPS periode 2019-2023.

D. Variabel Penelitian

Variabel penelitian terdiri dari dua macam, yakni variabel dependent dan variabel independen. Variabel independen (X) dalam penelitian ini adalah Pelayanan Kesehatan dan Tingkat Konsumsi Rumah Tangga (Y) disini ialah IPM.

⁶⁰ and others Sihabudin, Danny Wibowo, Sri Mulyono, Jaka Wijaya Kusuma, Irvana Arofah, Besse Arnawisuda Ningsi, *Ekonometrika Dasar Teori Dan Praktik Berbasis SPSS*, 2021.

⁶¹ Sugiyono, *Metode Penelitian Kuantitatif Kualitatif Dan R & B* (Bandung: Alfabeta, 2009).

Tabel 3.1 Variabel Penelitian

Variabel	Pelayanan Kesehatan (X1)
Definisi Operasional	Sistem atau mekanisme yang digunakan untuk mengumpulkan, mengelola, dan mengalokasikan dana untuk menyediakan layanan kesehatan kepada masyarakat.
Indikator	Total dana Kesehatan yang dikeluarkan provinsi dibidang Kesehatan.
Skala Pengukuran	Ribuan rupiah
Variabel	Konsumsi Rumah Tangga (X2)
Definisi Operasional	Kegiatan menghabiskan pendapatan yang dimiliki oleh anggota rumah tangga untuk memenuhi kebutuhan dan keinginan mereka.
Indikator	Total belanja rumah tangga dalam membiaya sandang pangan, papan dan lainnya dalam pemenuhan kebutuhan rumah tangga.
Skala Pengukuran	Ribuan Rupiah
Variabel	Indeks Pembangunan Manusia (Y)
Definisi Operasional	Suatu indikator yang digunakan untuk mengukur dan membandingkan tingkat kualitas hidup dan kesejahteraan manusia di berbagai negara atau wilayah.
Indikator	Indeks

Skala	Persen
Pengukuran	

E. Metode Analisis Data

Metode analisis data dalam penelitian ini menggunakan analisis regresi data panel. Dalam analisis regresi, selain mengukur kekuatan hubungan, juga menunjukkan arah hubungan antara variabel dependen dan variabel independen.⁶² Analisis regresi pada penelitian ini yaitu melibatkan satu variabel independen (X) yaitu Pelayanan Kesehatan dan Konsumsi Rumah tangga serta satu variabel dependen (Y) yaitu IPM.

1. Uji Asumsi Klasik

Alat yang digunakan adalah uji asumsi Klasik yaitu untuk mengetahui apakah terdapat masalah dalam di dalam data regresi. Uji asumsi Klasik yang digunakan untuk mengetahui bagaimana pengaruh variabel bebas (X) terhadap variabel terikat (Y). Ada empat pengujian dalam uji asumsi Klasik yaitu:

a. Uji Normalitas

Uji normalitas dilakukan untuk menguji apakah dalam model regresi, variabel pengganggu atau residual memiliki distribusi normal atau tidak. Regresi yang baik adalah regresi yang memiliki data yang berdistribusi normal. Metode yang baik yang layak digunakan dalam penelitian ini adalah metode kolmogrovsмирnov untuk mengetahui normal atau tidaknya data yang digunakan. Uji kolmogrov smirnov

⁶² Sinaga, *Statistik Dasar*.

adalah uji beda antara data yang di uji normalitasnya dengan data normal baku. Dengan pengambilan keputusan:

- 1) Jika $Sig > 0,05$ maka data berdistribusi normal.
- 2) Jika $Sig < 0,05$ maka data tidak berdistribusi normal.

b. Uji Autokorelasi

Autokorelasi adalah korelasi antara anggota serangkaian observasi yang diurutkan menurut waktu (seperti dalam data deretan waktu) atau ruang (data *cross section*). Uji ini dilakukan untuk menguji apakah dalam model regresi linier ada korelasi antara kesalahan pengganggu pada periode waktu atau ruang dengan kesalahan pengganggu waktu atau ruang sebelumnya. Jika data tidak memiliki masalah autokorelasi maka persamaan tersebut baik atau layak. Jika $Obs*R-squared < 0,05$ maka asumsi ditolak, tapi jika $Obs*R-squared > 0,05$ maka asumsi diterima dan tidak ada autokorelasi.

c. Uji Heterokadeksitas

Uji heteroskedastisitas menguji terjadinya perbedaan variance residual suatu periode pengamatan pada periode pengamatan lainnya. Masalah heteroskedastisitas biasanya terjadi dalam data cross section dibandingkan dengan data time series. Dalam penelitian ini, peneliti menggunakan uji park untuk mendeteksi ada atau tidaknya heteroskedastisitas. Uji park pada prinsipnya meregres residual yang dikuadratkan dengan variable bebas pada model, dengan ketentuan:

- 1) Jika statistik $> t$ -tabel atau nilai probabilitas $< 0,05$ maka ada heteroskedastisitas.
- 2) Jika t -statistik $< t$ -tabel atau nilai probabilitas $> 0,05$ maka tidak ada heteroskedastisitas.⁶³

d. Uji Multikolinearitas

Multikolinearitas dapat diartikan sebagai suatu keadaan dimana satu atau lebih variabel bebas dapat dinyatakan sebagai kombinasi kolinier dari variabel yang lainnya. Uji ini bertujuan untuk mengetahui apakah dalam regresi ini ditemukan adanya korelasi antar variabel independen. Jika terjadi korelasi maka dinamakan terdapat problem multikolinieritas. Cara mendeteksi adanya multikolinieritas dilakukan dengan uji Variance Inflation Factor (VIF) yang dihitung dengan rumus sebagai berikut:

Jika $VIF > 10$, maka antar variabel bebas (independent variabel) terjadi persoalan multikolinearitas.

Cara untuk mengetahui multikolinearitas dalam suatu model. Salah satunya adalah dengan melihat koefisien korelasi hasil output komputer. Jika terdapat koefisien korelasi yang lebih besar dari 0,9 maka terdapat gejala multikolinearitas. Untuk mengatasi masalah multikolinearitas, satu variabel independen yang memiliki korelasi

⁶³ Rahmad Solling, Samsul Bachri, Salju, and Muhammad Ikbil Hamid, "PANDUAN PRAKTIS EKONOMETRIKA: Konsep Dasar Dan Penerapan Menggunakan EViews 10," 2020.

dengan variabel independen lain harus dihapus. Dalam hal metode GLS, model ini sudah diantisipasi dari multikolinearitas.⁶⁴

F. Teknik Analisi Data

Penelitian ini menggunakan teknik analisis data panel dengan menggunakan program Stata Versi 17. Analisis data panel merupakan analisis data yang berstruktururut waktu (time series) sekaligus kerat lintang (cross section). Regresi panel merupakan sekumpulan teknik untuk memodelkan pengaruh peubah penjelas terhadap peubah respon pada data panel. Data panel dapat menjelaskan dua macam informasi yaitu: informasi cross section pada perbedaan antar subjek, dan informasi time series yang merefleksikan perubahan pada waktu. Maka jika kedua data tersebut tersedia maka data panel dapat digunakan.

Keuntungan menggunakan analisis data panel antara lain:

- 1) Memberikan jumlah pengamatan yang besar pada penelitian, meningkatkan *degree of freedom* (derajat kebebasan), data memiliki variabelitas yang besar, mengurangi kolineritas antara variable penjelas.
- 2) Dapat memberikan informasi lebih banyak yang tidak dapat diberikan jika hanya menggunakan data *time series* atau *cross section* saja. Data panel dapat memberikan penyelesaian yang lebih baik dalam inferensi perubahan dinamis jika dibandingkan dengan *cross section*.⁶⁵

⁶⁴ Amalia Anisa Fairuz, "Pengaruh Rasio Aktivitas, Rasio Solvabilitas, Rasio Pasar, Inflasi Dan Kurs Terhadap Return Saham Syariah," *UIN Journal*, 2018, 2017.

⁶⁵ Aris Slamet Widodo, *Kewirausahaan Entrepreneur Agribusiness Start Your Own Buisiness*, 1st ed. (Yogyakarta: Jaring Inspiratif, 2012); Anna Paula Soares Cruz, "Processing Data

Dalam model panel data, persamaan model yang menggunakan data *cross section* ditulis sebagai berikut:

$$Y_i = \beta_0 + \beta_1 + \epsilon_i ; i = 1, 2, \dots, N \dots$$

Di mana: N adalah banyaknya data *cross section*

Sedangkan persamaan model dengan *time series* adalah:

$$Y_t = \beta_0 + \beta_1 X_1 + \epsilon_i ; t = 1, 2, \dots, T \dots$$

Di mana: T adalah banyaknya data *time series*

Data panel merupakan gabungan dari *time series* dan *cross section* maka dapat diambil model yaitu:

$$Y_{it} = \beta_0 + \beta_1 X_{1it} + \beta_2 X_{2it} + \epsilon_i ; t \dots$$

$$i = 1, 2, \dots, N; t = 1, 2, \dots, T$$

Di mana:

Y_{it} : variabel dependen untuk individu i pada waktu t

β : koefisien untuk variabel independen $\beta_1 X_{it}$

X_{it} : variabel independen untuk individu i pada waktu t

ϵ_i : error term (gangguan) untuk individu i

N: banyaknya observasi

T: banyaknya waktu

N x T: banyaknya data panel

G. Model Estimasi Data Panel

Dalam metode estimasi model regresi dengan menggunakan data panel dapat dilakukan melalui tiga pendekatan, antara lain

1. *Common Effect Model*

Merupakan pendekatan model data panel yang paling sederhana karena hanya mengkombinasikan data *time series* dan *cross section*. Pada model ini tidak diperhatikan dimensi waktu maupun individu, sehingga diasumsikan bahwa perilaku data perusahaan sama dalam berbagai kurun waktu. Metode ini bisa menggunakan pendekatan *Ordinary Least Square* (OLS) atau teknik kuadrat terkecil untuk mengestimasi model data panel.

2. *Fixed Effect Model*

Model ini mengasumsikan bahwa perbedaan antar individu dapat diakomodasi dari perbedaan intersepnya. Untuk mengestimasi data panel model *Fixed Effects* menggunakan teknik *variable dummy* untuk menangkap perbedaan intersep antar perusahaan, perbedaan intersep bisa terjadi karena perbedaan budaya kerja, manajerial, dan insentif. Namun demikian slopnya sama antar perusahaan. Model estimasi ini sering juga disebut dengan teknik *Least Squares Dummy Variable* (LSDV).

3. *Random Effect Model*

Model ini akan mengestimasi data panel dimana variabel gangguan mungkin saling berhubungan antar waktu dan antar individu. Pada model *Random Effect* perbedaan intersep diakomodasi oleh *error terms* masing-masing perusahaan. Keuntungan menggunakan model *Random Effect* yakni

menghilangkan heteroskedastisitas. Model ini juga disebut dengan *Error Component Model* (ECM) atau teknik *Generalized Least Square* (GLS).⁶⁶

H. Penentuan Estimasi Data Panel

Untuk memilih model yang paling tepat terdapat beberapa pengujian yang dapat dilakukan, antara lain:

1. Uji Chow

Chow tes adalah pengujian untuk menentukan model apakah Common Effect (CE) ataukah Fixed Effect (FE) yang paling tepat digunakan dalam mengestimasi data panel. Apabila Hasil:

Nilai prob. $F < 0.05$, maka memilih FE

Nilai prob. $F > 0.05$, maka memilih CE

2. Uji Hausman

Hausman tes adalah pengujian statistik untuk memilih apakah model Fixed Effect atau Random Effect yang paling tepat digunakan. Apabila Hasil:

Nilai probabilitas chi squares < 0.05 , maka memilih FE

Nilai probabilitas chi squares > 0.05 , maka memilih RE

3. Uji Lagrange Multiplier

Uji Lagrange Multiplier dilakukan apabila hasil Uji Hausman memilih RE sebagai model terbaik, sehingga dilakukan pengujian selanjutnya mengenai apakah model RE atau CE yang paling tepat digunakan dalam mengestimasi data panel. Apabila Hasil:

Nilai p value < 0.05 , maka memilih RE

⁶⁶ Agus Tri Basuki Nano Prawoto, "Analisis Regresi Dalam Penelitian Ekonomi Dan Bisnis (Dilengkapi Aplikasi SPSS Dan Eviews)," 18 PT Rajagrafindo Persada, Depok § (2019).

Nilai p value > 0.05 , maka memilih CE

I. Uji Hipotesis

1. Uji t (Parsial)

Uji t digunakan untuk menguji pengaruh variabel independen secara parsial terhadap variabel dependen, yaitu pengaruh dari variabel independen yang terdiri atas pelayanan kesehatan dan konsumsi rumah tangga terhadap indeks pembangunan manusia per-Provinsi yang merupakan variabel dependennya. Uji t yang dilakukan untuk mengetahui pengaruh secara signifikan antara variabel bebas dan variabel terikat. Pengambilan keputusan uji hipotesis secara parsial juga didasarkan pada nilai probabilitas yang didapatkan dari hasil dengan kriteria uji:

- a. Jika $\text{sig} > 0,05$ maka H_0 diterima
- b. Jika $\text{sig} < 0,05$ maka H_a ditolak
- c. Jika $\text{sig} < 0,05$ maka H_0 ditolak
- d. Jika $\text{sig} > 0,05$ maka H_a diterima ⁶⁷

2. Uji F (Simultan)

Nilai statistik F adalah untuk menunjukkan apakah semua variabel independen yang dimaksudkan dalam persamaan regresi secara bersamaan berpengaruh terhadap variabel dependen. Apabila F_{hitung} lebih besar dari F_{tabel} maka H_0 ditolak dan H_a diterima yang berarti secara bersamaan variabel independen berpengaruh terhadap variabel dependen. Dengan kriteria:

⁶⁷ Wing Wahyu Winarno, "Analisis Ekonometrika Dan Statistika Dengan EViews (Edisi 5)," *Analisis Ekonometrika Dan Statistika Dengan EViews (Edisi 5)* 102, no. 1 (2017): 53–71.

- a. Jika $F_{hitung} > F_{tabel}$, maka H_0 ditolak dan H_a diterima
 - b. Jika $F_{hitung} < F_{tabel}$, maka H_0 diterima dan H_a ditolak
3. Koefisien Determinasi (R^2)

Koefisien Determinasi (R^2) bertujuan untuk mengukur seberapa jauh kemampuan model dalam menerangkan variasi variabel dependen. Nilai koefisien determinasi adalah antara nol dan satu. Nilai R^2 yang kecil berarti kemampuan variabel independen dalam menjelaskan variasi variabel dependen amat terbatas. Nilai yang mendekati satu berarti variabel independen memberikan hampir semua informasi yang dibutuhkan untuk memprediksi variasi variabel independen.