

BAB III

METODA PENELITIAN

3.1 Strategi Penelitian

Strategi penelitian ini adalah penelitian kausalitas. Sugiono (2017:21) mengemukakan bahwa penelitian kausalitas adalah hubungan yang bersifat sebab akibat. Jadi, disini ada variabel yaitu variabel independen (mempengaruhi) dan variabel dependen (dipengaruhi).

Penelitian pada dasarnya untuk menunjukkan kebenaran dan pemecahan masalah atas apa yang diteliti. Untuk mencapai tujuan tersebut, dilakukan suatu metode penelitian yang tepat dan relevan.

Sugiyono (2017:22) mengemukakan bahwa cara ilmiah untuk mendapatkan data dengan tujuan dan kegunaan tertentu. Terdapat empat kata kunci yang perlu diperhatikan yaitu, cara ilmiah, data, tujuan, dan kegunaan. Cara ilmiah berarti kegiatan penelitian itu didasarkan pada ciri-ciri keilmuan, yaitu rasional, empiris, dan sistematis. Rasional berarti kegiatan penelitian itu dilakukan dengan cara-cara yang masuk akal, sehingga terjangkau oleh penalaran manusia. Empiris berarti cara-cara yang dilakukan itu dapat diamati oleh indera manusia, sehingga orang lain dapat mengamati dan mengetahui cara-cara yang digunakan. Sistematis artinya proses yang digunakan dalam penelitian itu menggunakan langkah-langkah tertentu yang bersifat logis. Pada penelitian ini metode penelitian yang digunakan penulis adalah metode kuantitatif dengan penelitian survei.

Sugiyono (2017:8) mengemukakan bahwa metode kuantitatif adalah metode penelitian yang berlandaskan pada filsafat positivisme, digunakan untuk meneliti pada populasi atau sampel tertentu, pengumpulan data menggunakan instrumen penelitian, analisis data bersifat kuantitatif/statistik, dengan tujuan untuk menguji hipotesis yang telah ditetapkan. Sedangkan penelitian survei yaitu penelitian yang digunakan untuk menjelaskan hubungan kausal dan pengujian hipotesis.

Sugiyono (2017:12) mengemukakan bahwa penelitian survei adalah penelitian yang dilakukan pada populasi besar maupun kecil, tetapi data yang dipelajari adalah data dari sampel yang diambil dari populasi tersebut, untuk menemukan kejadian-kejadian relatif, distribusi, dan hubungan-hubungan antar variabel sosiologis maupun psikologis.

3.2 Populasi dan Sampel

3.2.1 Populasi Penelitian

Sugiyono (2017:80) mengemukakan bahwa populasi penelitian adalah wilayah generisasi yang terdiri atas objek atau subjek yang mempunyai kualitas dan karakteristik tertentu yang ditetapkan oleh peneliti untuk dipelajari dan kemudian ditarik kesimpulannya.

Dari definisi diatas dapat disimpulkan bahwa populasi bukan hanya sekedar jumlah yang ada pada obyek atau subyek yang dipelajari, melainkan meliputi seluruh karakteristik atau sifat yang dimiliki oleh subyek atau obyek tersebut. Dalam penelitian ini, populasi yang digunakan adalah seluruh perusahaan manufaktur sektor industri barang konsumsi yang terdaftar di Bursa Efek Indonesia tahun 2015-2019.

Tabel 3.1.
Daftar Populasi Penelitian

No	Kode>Nama Perusahaan	Nama
1	ADES	Akasha Wira International Tbk.
2	BUDI	Budi Starch & Sweetener Tbk.
3	CINT	Chitose Internasional Tbk.
4	DLTA	Delta Djakarta Tbk.
5	DVLA	Darya-Varia Laboratoria Tbk.
6	GGRM	Gudang Garam Tbk.
7	HMSP	H.M. Sampoerna Tbk.
8	ICBP	Indofood CBP Sukses Makmur Tbk
9	INDF	Indofood Sukses Makmur Tbk.

No	Kode>Nama Perusahaan	Nama
10	KAEF	Kimia Farma Tbk.
11	KLBF	Kalbe Farma Tbk.
12	MERK	Merck Tbk.
13	MLBI	Multi Bintang Indonesia Tbk.
14	MYOR	Mayora Indah Tbk.
15	PYFA	Pyridam Farma Tbk
16	ROTI	Nippon Indosari Corpindo Tbk.
17	SIDO	Industri Jamu dan Farmasi Sido
18	SKBM	Sekar Bumi Tbk.
19	SKLT	Sekar Laut Tbk.
20	STTP	Siantar Top Tbk.
21	TCID	Mandom Indonesia Tbk.
22	TSPC	Tempo Scan Pacific Tbk.
23	ULTJ	Ultra Jaya Milk Industry & Trading Company Tbk.
24	UNVR	Unilever Indonesia Tbk.
25	WIIM	Wismilak Inti Makmur Tbk.

Sumber: Arsip Peneliti

3.2.2 Sampel Penelitian

Sugiyono (2017:12) mengemukakan bahwa sampel penelitian adalah bagian dari jumlah dan karakteristik yang dimiliki oleh populasi tersebut. Bila populasi besar dan peneliti tidak mungkin mempelajari semua yang ada pada populasi, misalnya keterbatasan dana, tenaga, dan waktu, maka peneliti dapat menggunakan sampel yang diambil dari populasi tersebut.

Teknik pengambilan sampel yang digunakan pada penelitian ini adalah metode *non probability* dengan teknik pengambilan sampel penelitian ini dilakukan dengan *purposive sampling*.

Adapun beberapa pertimbangan atau kriteria penentu sampel yang akan digunakan pada penelitian ini disajikan pada tabel 3.2. berikut ini:

Tabel 3.2.
Kriteria Penentuan Sampel

No	Keterangan	Jumlah Perusahaan
1	Jumlah perusahaan manufaktur sektor industri barang konsumsi yang terdaftar di Bursa Efek Indonesia	(60)
2	Perusahaan manufaktur sektor industri barang konsumsi yang baru masuk bursa pada periode pengamatan	(21)
3	Perusahaan dengan laporan keuangan tidak lengkap	(4)
4	Perusahaan yang mengalami rugi pada periode penelitian	(10)
Jumlah Perusahaan yang digunakan sebagai sampel		25
Periode pengamatan (tahun)		5
Jumlah Sampel		125

Sumber: Data yang diolah

3.3 Data dan Metoda Pengumpulan Data

3.3.1 Data Penelitian

Sugiyono (2017:16) mengemukakan bahwa data yang digunakan dalam penelitian ini adalah data sekunder. Data sekunder adalah sumber data yang tidak langsung memberikan data kepada pengumpul data. Data sekunder ini merupakan data yang sifatnya mendukung keperluan data primer seperti buku-buku, literatur dan bacaan yang berkaitan dengan menunjang penelitian ini. Sumber data yang digunakan dalam penelitian ini adalah sumber data sekunder yang berasal dari data eksternal. Sumber data sekunder eksternal adalah sumber data yang diperoleh dari sumber di luar perusahaan dan diperoleh secara tidak langsung, melalui media perantara. Data yang digunakan dalam penelitian ini di peroleh dari Bursa Efek Indonesia pada perusahaan manufaktur sektor industri barang konsumsi tahun 2015-2019 yang didapat dari *Annual Report* yang diperoleh dari situs resmi Bursa Efek Indonesia (BEI) pada www.idx.co.id.

3.3.2 Metoda Pengumpulan Data

Sugiyono (2017:224) teknik pengumpulan data adalah beberapa cara yang dilakukan untuk memperoleh data dan keterangan-keterangan yang diperlukan dalam penelitian. Dalam pengumpulan data pada penelitian ini adalah sebagai berikut:

1. Metoda Dokumentasi

Metoda ini dilakukan dengan cara mengumpulkan data sekunder yang terdapat dalam laporan keuangan tahunan yang sudah dipublikasikan melalui website resmi pada masing-masing perusahaan yang diperoleh dari statistik total asset perusahaan dan publikasi lainnya yang terkait dengan hipotesis penelitian.

3.4 Operasionalisasi Variabel

Sugiyono (2017:37) mengemukakan bahwa operasionalisasi variabel adalah penggambaran definisi yang ada dalam penelitian. Variabel penelitian pada dasarnya adalah segala sesuatu yang berbentuk apa saja yang ditetapkan oleh peneliti untuk dipelajari sehingga diperoleh informasi tentang hal tersebut, kemudian ditarik kesimpulannya.

Variabel-variabel yang dibutuhkan dalam penelitian ini ada tiga (3) jenis yaitu (2) variabel independen : Umur Perusahaan, Ukuran Perusahaan, dan (1) variabel dependen yaitu Pengungkapan *Corporate Social Responsibility* (CSR). Dari variabel yang sudah tertera maka dapat diuraikan sebagai berikut:

3.4.1 Variabel Independen (X)

Variabel independen sering disebut sebagai variabel bebas yang artinya variabel yang mempengaruhi atau menjadi sebab perubahannya atau timbulnya variabel terikat. Variabel independen dalam penelitian ini adalah Umur Perusahaan (X1) dan Ukuran Perusahaan (X2).

3.4.2 Variabel Dependen (Y)

Variabel dependen atau biasa disebut dengan variabel terikat yang artinya variabel ini muncul karena dipengaruhi atau menjadi akibat dari variabel bebas atau independen. Variabel dependen dalam penelitian ini adalah pengungkapan *Corporate Social Responsibility* (Y). Jika disimpulkan, masing-masing variabel penelitian secara operasional dapat didefinisikan sebagai berikut:

Tabel 3.3.

Ringkasan Operasionalisasi Variabel

Variabel yang diukur	Alat Ukur(Dimensi)	Skala
Variabel Independen (X)		
Umur Perusahaan (X1)	Dihitung mulai tahun pendirian perusahaan hingga tahun penelitian. (Collins dan Poras, 2001)	Nominal
Ukuran Perusahaan (X2)	Size = Ln (Total Aset) (Pantow, 2015)	Nominal
Variabel Dependen (Y)		
Pengungkapan CSR (Y1)	$CSR = \frac{\text{Jumlah Pengungkapan Item}}{\text{Jumlah Item yang diharapkan}}$ (Sudana, 2015)	Nominal

Sugiyono (2016:38) mengemukakan bahwa variabel penelitian adalah segala sesuatu yang berbentuk apa saja yang ditetapkan oleh peneliti untuk dipelajari sehingga diperoleh informasi tentang hal tersebut, kemudian ditarik kesimpulannya.

3.5 Metoda Analisis Data

Sugiyono (2017:147) mengemukakan bahwa metoda analisis data yaitu mengelompokkan data berdasarkan variabel dan jenis responden, metatulasi data

berdasarkan variabel dari seluruh responden, menyajikan data tiap variabel yang diteliti, melakukan perhitungan untuk menjawab rumusan masalah, dan melakukan perhitungan untuk menguji hipotesis yang telah diajukan.

Metoda analisis data yang akan digunakan dalam penelitian ini adalah analisis regresi parsial dan berganda, dimana pengolahan tersebut menggunakan analisis statistik deskriptif. Penelitian ini menggunakan alat bantu program komputer untuk mengelola data berupa *Software Eviews* versi 11. Sebelum melakukan uji regresi parsial dan berganda dilakukan uji data yaitu uji asumsi klasik.

3.5.1 Uji Asumsi Klasik

Uji asumsi klasik harus dilakukan terlebih dahulu untuk mengetahui apakah data layak untuk dianalisis. Tujuannya adalah untuk menghindari terjadinya estimasi yang bias, karena tidak semua data dapat diterapkan regresi.

Uji asumsi klasik yang digunakan adalah uji normalitas, uji reliabilitas, uji multikolinearitas dan uji heteroskedastisitas. Dalam menganalisis regresi linear untuk menghindari penyimpangan asumsi klasik perlu dilakukan beberapa uji antara lain:

3.5.1.1 Uji Normalitas

Ghozali (2016:154) mengemukakan bahwa uji normalitas dilakukan untuk menguji apakah dalam model regresi variabel independen dan variabel dependen atau keduanya mempunyai distribusi normal atau tidak. Apabila variabel tidak berdistribusi secara normal maka hasil uji statistik akan mengalami penurunan. Pengambilan keputusan pada pengujian ini didasarkan pada nilai probabilitas *Jarque-Bera* hasil pengujian. Jika nilai probabilitas *Jarque-Bera* lebih besar daripada taraf signifikansi sebesar 0,05 yang digunakan maka dapat disimpulkan bahwa data telah

terdistribusi normal. Jika nilai probabilitas *Jarque-Bera* lebih kecil daripada taraf signifikansi yang digunakan maka dapat disimpulkan bahwa data tidak terdistribusi normal.

3.5.1.2 Uji Multikolinearitas

Ghozali (2016:103) mengemukakan bahwa pengujian multikolinearitas bertujuan untuk menguji apakah model regresi ditemukan adanya korelasi antar variabel bebas (independen). Pengujian multikolinearitas adalah pengujian yang mempunyai tujuan untuk menguji apakah dalam model regresi ditemukan adanya korelasi antara variabel independen. Pengambilan keputusan pada pengujian ini didasarkan pada nilai *Centered VIF*. Jika nilai *Centered VIF* lebih kecil dari 10 maka dapat disimpulkan bahwa tidak terjadi masalah multikolinearitas dalam penelitian ini. Jika nilai *Centered VIF* lebih besar dari 10 maka dapat disimpulkan bahwa terjadi masalah multikolinearitas dalam penelitian ini.

3.5.1.3 Uji Heteroskedastisitas

Ghozali (2016:134) mengemukakan bahwa uji ini bertujuan untuk menguji apakah dalam sebuah model regresi terjadi ketidaknyamanan varian dari residual satu pengamatan ke pengamatan lain. Pengambilan keputusan dalam pengujian ini didasarkan pada nilai probabilitas *chi-square* dari *Obs*R-square*. Jika nilai probabilitas *chi-square* dari *Obs*R-square* lebih besar dari taraf signifikansi sebesar 0.05 yang telah ditetapkan maka dapat disimpulkan bahwa tidak terjadi gejala heteroskedastisitas dalam penelitian ini. Jika nilai probabilitas *chi-square* dari *Obs*R-square* lebih kecil dari taraf signifikansi yang ditetapkan maka dapat disimpulkan bahwa terdapat masalah heteroskedastisitas dalam penelitian ini.

3.5.1.4 Uji Autokorelasi

Ghozali (2016;107) mengemukakan bahwa autokorelasi muncul karena observasi yang berurutan sepanjang waktu berkaitan satu sama lainnya. Permasalahan ini muncul karena residual tidak bebas dari satu observasi ke observasi lainnya. Model regresi yang baik adalah model regresi yang bebas dari autokorelasi.

Pengambilan keputusan dalam pengujian ini didasarkan pada nilai probabilitas *chi-square* dari *Obs*R-square*. Jika nilai probabilitas *chi-square* dari *Obs*R-square* lebih besar dari taraf signifikansi sebesar 0.05 yang telah ditetapkan maka dapat disimpulkan bahwa tidak terjadi gejala autokorelasi dalam penelitian ini. Jika nilai probabilitas *chi-square* dari *Obs*R-square* lebih kecil dari taraf signifikansi yang ditetapkan maka dapat disimpulkan bahwa terdapat masalah autokorelasi dalam penelitian ini.

3.5.2 Analisis Regresi Data Panel

Analisis regresi data panel digunakan untuk mengukur pengaruh antara lebih dari satu variabel. Data panel adalah data yang dikumpulkan secara *cross section* dan diikuti pada periode waktu tertentu. Teknik data panel yaitu dengan menggabungkan jenis data *cross section* dan *time series*. Menurut Basuki dan Prawoto (2017:275) data panel merupakan gabungan antara dua kurun waktu (*time series*) dan data silang (*cross section*). Keunggulan penggunaan data panel memberikan keuntungan diantaranya sebagai berikut (Basuki dan Prawoto, 2017:275) :

1. Data panel mampu memperhitungkan heterogenitas individu secara eksplisit dengan mengizinkan variabel spesifik individu.
2. Data panel dapat digunakan untuk menguji, membangun dan mempelajari model- model perilaku yang kompleks.

3. Data panel mendasarkan diri pada observasi yang bersifat *cross section* yang berulang-ulang (*time series*) sehingga cocok digunakan sebagai *study of dynamic adjustment*.
4. Data panel memiliki implikasi pada data yang lebih informatif, lebih bervariasi dan dapat mengurangi kolinieritas antar variabel derajat kebebasan (*degree of freedom*) yang lebih tinggi sehingga dapat diperoleh hasil estimasi yang lebih efisien.
5. Data panel dapat digunakan untuk meminimalkan bias yang mungkin ditimbulkan oleh agregasi data individu.
6. Data panel dapat mendeteksi lebih baik dan mengukur dampak yang secara terpisah di observasi dengan menggunakan data *time series* ataupun *cross section*.

3.5.3 Metoda Estimasi Regresi Data Panel

Teknik model regresi data panel dapat dilakukan dengan tiga pendekatan alternatif metoda pengolahannya yaitu metoda *Common Effect Model* (CEM), *Fixed Effect Model* (FEM), dan *Random Effect Model* (REM) sebagai berikut :

3.5.3.1 *Common Effect Model* (CEM)

Fairuz (2017:187) mengemukakan bahwa *Common Effect Model* dikatakan sebagai model yang paling sederhana, dimana pendekatannya mengabaikan dimensi waktu dan ruang yang dimiliki oleh data panel. *Common Effect Model* dilakukan dengan mengkombinasikan data *time series* dan *cross-section*. Penggabungan kedua jenis data tersebut dapat digunakan metoda OLS biasa sehingga sering disebut dengan *Pooled Least Square* atau *common OLS* model untuk mengestimasi model data panel.

Common Effect Model adalah model yang paling sederhana, karena metode yang digunakan dalam *Common Effect*

Model hanya dengan mengkombinasikan data *time series* dan *cross section*. Dengan hanya menggabungkan kedua jenis data tersebut, maka dapat digunakan *Ordinal Least Square Model* (OLS) atau teknik kuadrat terkecil untuk mengestimasi model data panel.

Dalam pendekatan ini tidak memperhatikan dimensi individu maupun waktu, dan dapat diasumsikan bahwa perilaku data antar perusahaan sama dalam rentan waktu. Asumsi ini jelas sangat jauh dari realita sebenarnya, karena karakteristik antar perusahaan baik dari segi kewilayahan jelas sangat berbeda.

Persamaan metode ini dapat dirumuskan sebagai berikut:

$$Y_{it} = \alpha + \beta X_{it} + \varepsilon_{it}$$

Dimana :

Y_{it}	: Variabel terikat individu ke-i pada waktu ke-i X_j
it	: Variabel bebas ke-j individu ke-i pada waktu ke-t
I	: Unit cross-section sebanyak N
J	: Unit <i>time series</i> sebanyak T
ε_{it}	: Komponen error individu ke-i pada waktu ke-t
α	: <i>Intercept</i>
β	: Parameter untuk variabel ke-j

3.5.3.2 *Fixed Effect Model (FEM)*

Fairuz (2017:189) mengemukakan bahwa model ini digunakan untuk mengatasi kelemahan dari analisis data panel yang menggunakan *Common Effect Model*, penggunaan data panel *Common Effect Model* tidak realistis karena akan menghasilkan *intercept* ataupun *slope* pada data panel yang tidak berubah baik antar individu (*cross section*) maupun antar waktu (*time series*). Model ini juga untuk mengestimasi data panel dengan menambahkan variabel dummy. Teknik ini dinamakan *Least*

Square Dummy Variabel (LSDV).

Selain diterapkan untuk efek tiap individu, LSDV ini juga dapat mengkombinasikan efek waktu yang bersifat sismatik. Hal ini dapat dilakukan melalui penambahan variabel *dummy* waktu di dalam model. Model ini digunakan untuk mengatasi kelemahan dari analisis data panel yang menggunakan metoda *Common Effect Model*, penggunaan data panel *Common Effect Model* tidak realistis karena akan menghasilkan *intercept* ataupun *slope* pada data panel yang tidak berubah baik antar individu (*cross section*) maupun antar waktu (*time series*).

Model ini juga untuk mengestimasi data panel dengan menambahkan variabel *dummy*. Model ini mengasumsikan bahwa terdapat efek yang berbeda antar individu. Perbedaan ini dapat diakomodasi melalui perbedaan diintersepnya. Oleh karena itu dalam *Fixed Effect Model*, setiap individu merupakan parameter yang tidak diketahui dan akan diestimasi dengan menggunakan teknik variabel *dummy* yang dapat dirumuskan sebagai berikut:

$$Y_{it} = \alpha_i + \beta_j X_{itj} + \sum_{i=2}^n \alpha_i D_i + \varepsilon_{it}$$

Dimana :

Y_{it}	: Variabel terikat individu ke-i pada waktu ke-i
X_{itj}	: Variabel bebas ke-j individu ke-i pada waktu ke-t
D_i	: <i>Dummy variabel</i>
ε_{it}	: Komponen error individu ke-i pada waktu ke-t
α	: <i>Intercept</i>
β_j	: Parameter untuk variabel ke-j

Teknik ini dinamakan *Least Square Dummy Variabel (LSDV)*. Selain diterapkan untuk efek tiap individu, LSDV ini juga dapat mengkombinasikan efek waktu yang bersifat sismatik. Hal ini dapat dilakukan melalui penambahan variabel *dummy* waktu di dalam model.

3.5.3.3 *Random Effect Model (REM)*

Fairuz (2017:192) mengemukakan bahwa metoda *Random Effect Model* adalah metoda yang menggunakan residual yang diduga memiliki hubungan antar waktu dan antar individu/perusahaan. Dalam metoda ini mengasumsikan bahwa setiap variabel mempunyai perbedaan intersep tetapi intersep tersebut bersifat *random*/stokastik. Dalam metoda ini perbedaan karakteristik individu dan waktu diakomodasikan dengan *error* dari model.

Mengingat terdapat dua komponen yang mempunyai kontribusi pada pembentukan *error* yaitu individu dan waktu, maka pada metoda ini perlu diuraikan menjadi *error* dari komponen individu, *error* untuk komponen waktu dan *error* gabungan.

Persamaan *Random Effect Model* dapat dirumuskan sebagai berikut:

$$Y_{it} = \alpha + \beta X_{it} + \varepsilon_{it} ; \varepsilon_{it} = u_i + V_t + W_{it}$$

Dimana :

u_i : Komponen error *cross-section*

V_t : Komponen *time series*

W_{it} : Komponen *error* gabungan

3.54 Pemilihan Model Regresi Data Panel

Software Eviews versi 11 memiliki beberapa pengujian yang akan membantu menemukan metoda apa yang paling efisien digunakan dari ketiga model tersebut. Pemilihan model untuk menguji persamaan regresi yang akan di estimasi dapat digunakan tiga penguji yaitu uji *Chow*, uji *Hausmant* dan uji *Langrange Multiplier* yang akan diuraikan sebagai berikut :

3.5.4.1 Uji Chow

Ghozali (2017:184) mengemukakan bahwa uji *Chow* adalah untuk menentukan uji mana di antara kedua model yakni *Common Effect Model* dan *Fixed Effect Model* yang sebaiknya digunakan dalam pemodelan data panel. Hipotesis dalam uji *chow* ini sebagai berikut :

H_0 : *Common Effect Model*

H_a : *Fixed Effect Model*

Apabila hasil uji ini menunjukkan probabilitas *F* lebih dari taraf signifikansi 0,05 maka model yang dipilih adalah *Common Effect Model*. Sebaliknya, apabila probabilitas *F* kurang dari taraf signifikansi 0,05 maka model yang sebaiknya dipakai adalah *Fixed Effect Model*.

3.5.4.2 Uji Hausman

Ghozali (2017: 199) mengemukakan bahwa uji *Hausman* yaitu untuk menentukan uji mana diantara kedua *Random Effect Model* dan *Fixed Effect Model* yang sebaiknya dilakukan dalam pemodelan data panel. Hipotesis dalam uji *Hausman* sebagai berikut :

H_0 : *Random Effect Model*

H_a : *Fixed Effect Model*

Jika probabilitas *Chi-Square* lebih kecil dari taraf signifikansi 0,05 maka H_0 ditolak dan model yang tepat adalah *Fixed Effect Model* dan sebaliknya.

3.5.4.3 Uji Langrange Multiplier

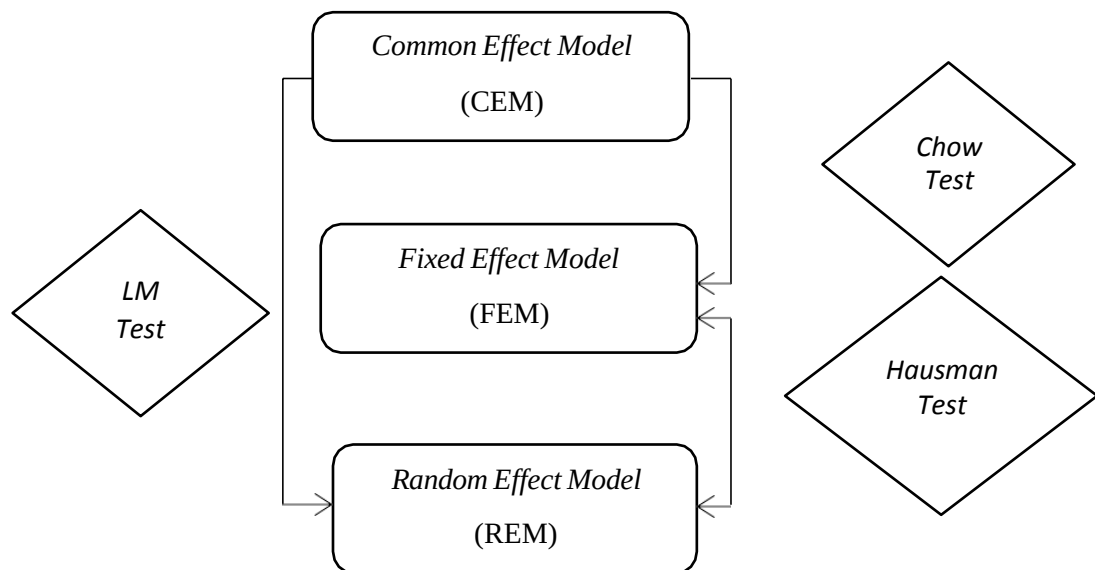
Ghozali (2017:180) mengemukakan bahwa *Lagrange Multiplier* (LM) adalah uji untuk mengetahui apakah *Random Effect Model* atau *Common Effect model* metode *Ordinal Least*

Square (OLS) yang paling tepat digunakan. Uji signifikansi *Random Effect model* ini dikembangkan oleh Breusch Pagan. Metode *Breusch Pagan* untuk uji signifikansi *Random Effect model* didasarkan pada nilai residual dari metode OLS.

Jika nilai LM statistik lebih besar dari nilai kritis statistik *chi-squares* maka kita menolak hipotesis nol, yang artinya estimasi yang tepat untuk model regresi data panel adalah *Random Effect Model* dari pada *Common Effect Model*. Sebaliknya jika nilai LM statistik lebih kecil dari nilai statistik *chi-squares* sebagai nilai kritis, maka kita menerima hipotesis nol, yang artinya estimasi yang digunakan dalam regresi data panel adalah *Common Effect Model* bukan *Random Effect Model*.

Uji LM tidak digunakan apabila pada uji *Chow* dan uji *Hausman* menunjukan model yang paling tepat adalah *Fixed Effect Model*. Uji LM dipakai manakala pada uji *Chow* menunjukan model yang dipakai adalah *Common Effect Model*, sedangkan pada uji *Hausman* menunjukan model yang paling tepat adalah *Random Effect Model*. Maka diperlukan uji LM sebagai tahap akhir untuk menentukan *Common Effect Model* atau *Random Effect Model* yang paling tepat.

Untuk menentukan pendekatan mana yang lebih baik digunakan pengujian baik dari model dan pengujian maka digambarkan sebagai berikut:



Gambar 3.1

Pengujian Kesesuaian Model

3.5.5 Analisis Statistik Deskriptif

Sugiyono (2017:45) mengemukakan bahwa statistik deskriptif berfungsi untuk memberikan gambaran tentang variabel-variabel penelitian yang dilihat dari rata-rata (*mean*), nilai maksimum, nilai minimum, dan standar deviasi. statistik deskriptif adalah statistik yang digunakan untuk menganalisis data dengan cara mendeskripsikan atau menggambarkan data yang berlaku untuk umum atau generalisasi.

Sujarweni (2015:37) mengemukakan bahwa statistik deskriptif berusaha untuk menggambarkan berbagai karakteristik data yang berasal dari suatu sampel. Statistik deskriptif seperti *mean*, *median*, *modus*, *presentile*, *desile*, *quartile*, dalam bentuk analisis angka maupun gambar/diagram.

Statistik deskriptif digunakan untuk menjelaskan dan menggambarkan variabel-variabel berdasarkan data yang dikumpulkan pada periode tertentu. Karakteristik data yang digambarkan dapat dilihat

sebagai berikut:

- a. Nilai maksimum dari sejumlah populasi yang dikumpulkan
- b. Nilai minimum dari sejumlah populasi yang dikumpulkan
- c. Nilai rata – rata (*mean*) dari sejumlah populasi yang dapat mewakili nilai-nilai yang terkumpul, maka akan digunakan rumus sebagai berikut:

1. *Mean* (rata-rata hitung)

Mean (rata-rata hitung) adalah suatu nilai yang diperoleh dengan cara membagi seluruh nilai pengamatan dengan banyaknya pengamatan. Mean dapat dirumuskan sebagai berikut

$$Me = \frac{\sum Xi}{n}$$

Keterangan :

Me = *Mean*

$\sum Xi$ = Jumlah Nilai X ke I sampai ke n

n = Jumlah Sampel atau data

2. Standar Deviasi

Standar deviasi digunakan untuk menilai *disperse* rata-rata atau sampel. Setelah rata-rata diketahui maka perlu ditentukan sebaran datanya. Semakin kecil sebarannya berarti nilai data semakin sama, jika sebarannya bernilai nol, maka nilai semua datanya adalah sama. Semakin besar nilai sebarannya maka nilai yang ada semakin bervariasi. Standar deviasi dapat dirumuskan sebagai berikut:

$$S = \sqrt{\frac{\sum_{i=1}^n (xi - x)^2}{n - 1}}$$

Keterangan :

S = Standar deviasi (simpangan baku)

Xi = Nilai X ke I sampai ke n

$\bar{x}$ = Rata – rata nilai

n = Jumlah sample

3. Median

Median adalah salah satu teknik penjelasan kelompok yang didasarkan atas nilai tengah dari kelompok data yang telah disusun urutannya dari yang terkecil sampai dengan yang terbesar. Rumus median adalah sebagai berikut :

$$\text{Med} = \frac{x_1 + x_2}{n}$$

Keterangan:

Med = Median

n = Banyaknya nilai tengah

X_1 = Nilai tengah pertama dimana median terletak

X_2 = Nilai tengah kedua dimana median terletak

3.5.6 Uji Regresi Linear Berganda

Analisis regresi linear berganda digunakan untuk mengetahui sejauh mana besarnya pengaruh variabel bebas (X) terhadap variabel terikat (Y). Metode ini menghubungkan satu variabel dependen dengan banyak variabel independen. Dalam penelitian ini yang menjadi variabel terikat adalah pengungkapan CSR, sedangkan yang menjadi variabel bebas adalah umur perusahaan dan ukuran perusahaan. Model hubungan pengungkapan CSR dengan variabel-variabel bebasnya tersebut disusun dalam fungsi atau persamaan sebagai berikut:

$$Y = a + b_1 X_1 + b_2 X_2 + e$$

Keterangan:

Y : Penafsiran Variabel Dependen (Pengungkapan CSR)

b_1 : Koefisien regresi Variabel Independen 1

b_2	: Koefisien regresi Variabel Independen 2
a	: Nilai Konstanta
X_1	: Umur Perusahaan
X_2	: Ukuran Perusahaan
e	: <i>Error</i>

3.5.7 Uji Hipotesis

Uji hipotesis dalam penelitian ini ada tiga pengujian yaitu analisis koefisien determinasi (*adjusted R²*), uji parsial (uji t), uji simultan (uji f), parsial sebagai berikut:

3.5.7.1 Uji Koefisien Determinasi (R^2)

Koefisien determinasi (R^2) digunakan untuk mengetahui besarnya sumbangan efektif yang diberikan variabel independen (umur perusahaan, dan ukuran perusahaan) terhadap variabel dependen (pengungkapan CSR). Semakin besar nilai determinasi maka semakin besar varian sumbangan terhadap variabel terikatnya. Koefisien determinasi yang digunakan adalah nilai (R^2). Rumusnya adalah :

$$KD = r^2 \times 100\%$$

Keterangan :

KD : Koefisien Determinasi

R : Koefisien Korelasi

3.5.7.2 Uji Parsial (t)

Uji parsial (uji t) pada dasarnya digunakan untuk mengetahui pengaruh variabel independen (bebas) terhadap variabel dependen (terikat) secara individual (Ghozali, 2016:99).

Untuk mengetahui nilai apakah nilai t statistik tabel, tingkat signifikan yang digunakan sebesar 5% dengan kriteria pengambilan

keputusan dalam pengujian ini adalah sebagai berikut:

1. Jika nilai probabilitas $< 0,05$ maka H_0 ditolak, artinya variabel independen secara parsial mempengaruhi variabel dependen.
2. Jika nilai probabilitas $> 0,05$ maka H_0 diterima, artinya variabel independen secara parsial tidak mempengaruhi variabel dependen.

3.5.7.3 Uji Simultan (f)

Uji simultan (uji f) digunakan untuk menguji kemampuan seluruh variabel independen secara bersama-sama dalam menjelaskan variabel dependen. Pengujian dapat dilakukan dengan membandingkan nilai f hitung dengan f tabel (Ghozali, 2016:95). Pada tingkat signifikan sebesar $\leq 0,05$ dengan kriteria pengujian sebagai berikut :

1. Apabila nilai *P-Value* f statistik $\leq 0,05$ maka H_0 ditolak dan H_1 diterima yang artinya variabel independen secara bersama-sama mempengaruhi variabel-variabel dependen.
2. Apabila nilai *P-Value* f statistik $\geq 0,05$ maka H_0 diterima dan H_1 ditolak yang artinya variabel independen secara bersama-sama mempengaruhi variabel-variabel dependen.