

BAB III

METODE PENELITIAN

3.1 Metode Penelitian

3.1.1 Jenis Penelitian

Penelitian ini menggunakan jenis penelitian dengan pendekatan kuantitatif. Yang berarti suatu metode penelitian yang berfokus pada analisis numerik dan pengumpulan data untuk menguji hipotesis dan menjawab rumusan masalah.

3.1.2 Objek Penelitian

Dalam penelitian ini objek yang digunakan adalah *Tax Avoidance* sebagai variabel independen (X1), *Tax Aggressiveness* sebagai variabel independen (X2), *Leverage* sebagai variabel independen (X3) serta nilai perusahaan yang terdaftar di Bursa Efek Indonesia sebagai variabel dependen (Y).

3.1.3 Sumber Data

Sumber data yang diambil untuk mendukung penelitian ini adalah sebagai berikut:

1. Buku maupun ebook yang berhubungan dengan judul penelitian.
2. Jurnal, Skripsi, Tesis dan data dari internet yang berhubungan dengan judul penelitian
3. Data laporan keuangan perusahaan yang telah dipublikasikan di Bursa Efek Indonesia (BEI) pada periode tahun 2020-2023 yang diambil dari situs resmi BEI yaitu www.idx.co.id, dan web resmi perusahaan.

3.1.4 Jenis Data

Jenis data yang digunakan dalam penelitian ini adalah data sekunder. Data sekunder merupakan suatu jenis data yang digunakan dalam penelitian dan diperoleh secara tidak langsung oleh penulis melalui media perantara, seperti catatan, bukti, atau laporan historis yang tersimpan dalam arsip baik yang dipublikasikan maupun tidak dipublikasikan. Dimana dalam penelitian ini data sekunder diambil dari data laporan tahunan seluruh perusahaan yang terdaftar di Bursa Efek Indonesia pada periode tahun 2020-2023.

3.1.5 Teknik Pengumpulan Data

Penelitian ini menggunakan teknik pengumpulan data dengan metode dokumentasi, yakni penggunaan data yang bersumber dari dokumen-dokumen yang sudah ada sebelumnya.

3.1.6 Populasi dan Sampel

a. Populasi

Menurut Sugiyono (2013) Populasi adalah wilayah generalisasi yang terdiri dari atas obyek atau subyek yang mempunyai kualitas dan karakteristik tertentu yang ditetapkan oleh peneliti untuk dipelajari dan kemudian ditarik kesimpulannya. Populasi pada penelitian ini adalah seluruh perusahaan yang terdaftar dalam Bursa Efek Indonesia pada tahun 2020-2023.

b. Sampel

Adapun sampel pada penelitian ini adalah perusahaan non-keuangan yang terdaftar dalam Bursa Efek Indonesia pada tahun 2020-2023. Penelitian ini menggunakan teknik purposive sampling sebagai teknik penarikan sampel.

Menurut (Sugiyono, 2013) *Purposive sampling* yaitu penarikan sampel didasarkan pada kriteria tertentu yang ditentukan dari penulis kepada yang akan diteliti. Yang menjadi kriteria penarikan sampel di sini adalah perusahaan yang terdaftar secara berturut-turut di BEI periode 2020-2023. Dalam *annual report* BEI tahun 2023 terdapat 903 perusahaan yang tercatat sampai dengan akhir tahun 2023.

Table 3.2 Pemilihan Sampel Penelitian

No	Keterangan	Jumlah
1	Perusahaan yang terdaftar di BEI tahun 2020-2023	903
2	Perusahaan yang termasuk dalam sektor keuangan periode BEI tahun 2020-2023	(94)
3	Perusahaan yang baru melakukan IPO setelah tahun 2020	(192)
4	Perusahaan yang terkena suspend/delisting selama periode 2020–2023	(44)
5	Perusahaan dengan laporan keuangan tidak lengkap	(22)
6	Perusahaan yang mengalami kerugian berturut-turut selama 3 tahun	(128)
	Total Sampel 423 Perusahaan x 4 tahun periode penelitian = 1.692	423

Sumber: BEI yang diolah penulis (2025)

3.1.7 Operasional Variabel

3.1.7.1 Definisi Operasional Variabel

a. Variabel Dependen

Variabel dependen yang sering disebut juga, sebagai variabel terikat merupakan suatu variabel yang dipengaruhi oleh variabel lain. Penulis menggunakan nilai perusahaan (*firm value*) sebagai variabel dependen yang sering dimaknai sebagai variabel Y ini. Nilai perusahaan merupakan suatu ukuran yang mencerminkan seberapa besar harga atau nilai suatu perusahaan menurut para investor. Untuk meneliti nilai perusahaan pada penelitian ini penulis menggunakan proksi *Price to Book Value* (PBV). Dimana nilai perusahaan diukur menggunakan rasio harga saham terhadap nilai buku per lembar sahamnya. Sehingga rumus yang digunakan sebagai proksi dari *Price to Book Value* adalah sebagai berikut:

$$PBV = \frac{\text{Harga Per Lembar Saham}}{\text{Nilai Buku Per Lembar Saham}}$$

b. Variabel Independen

Variabel independen atau biasa dikenal sebagai variabel bebas merupakan suatu variabel yang dianggap memengaruhi atau menjadi sebab perubahan pada variabel lain atau variabel dependen. Dalam penelitian ini terdapat tiga variabel independen yang terdiri dari :

b.1 Perencanaan pajak (*tax planning*)

Yang mencakup dua bagian di dalamnya, dengan proksi yang terdiri dari ETR (*Effective Tax Rate*) untuk mengukur *tax avoidance* dan BTD (*Book Tax Difference*) untuk mengukur *tax aggressiveness*. Perencanaan pajak (*tax planning*) Merupakan suatu strategi yang biasa digunakan oleh para manajemen perusahaan untuk mengoptimalkan pembayaran pajak mereka secara legal sesuai yang diatur oleh undang-undang perpajakan. Perencanaan pajak memiliki tujuan untuk mengoptimalkan jumlah pembayaran pajak tanpa melanggar hukum, sehingga suatu perusahaan dapat meningkatkan efisiensi keuangannya. Perencanaan pajak sendiri terdiri dari beberapa bagian diantaranya adalah *tax avoidance* dan *tax aggressiveness*. Dimana dalam penelitian ini proksi yang digunakan untuk mengukur *tax planning* terdiri dari *Effective Tax Rate (ETR)* dan *Book Tax Difference (BTD)* dengan penjelasan sebagai berikut:

b.1.1 *Effective Tax Rate (ETR)*

Merupakan suatu tingkat efektif perusahaan dalam pembayaran pajak yang dapat dihitung dengan cara membagi laba sebelum pajak dan beban pajak penghasilan suatu perusahaan. Menurut Dyreng, et al. dalam (Shella, et al., 2021) ETR sering digunakan untuk menggambarkan aktivitas penghindaran pajak suatu perusahaan secara menyeluruh. Dimana ETR adalah suatu proksi negatif, jika nilai ETR tinggi berarti tingkat *tax*

planning yang dilakukan suatu perusahaan rendah dan berlaku sebaliknya.

Rumus untuk menentukan ETR adalah sebagai berikut:

$$ETR = \frac{\text{Laba Sebelum Pajak}}{\text{Beban Pajak Penghasilan}}$$

b.1.2 *Book Tax Difference (BTD)*

Menurut (Hadrian, et al., 2022) *Book tax Difference* merupakan selisih jumlah laba yang dihitung berdasarkan akuntansi dengan laba yang dihitung dalam perpajakan terhadap nilai rata rata total aktiva yang diharapkan mampu menggambarkan aktivitas perencanaan pajak. Dalam penelitian ini menggunakan rumus sebagai berikut:

$$BTD = \frac{(\text{Laba Akuntansi} - \text{Laba Pajak})}{\text{Total Aset}}$$

b.2 *Leverage*

Merupakan kemampuan perusahaan dalam melunasi berbagai utangnya, baik jangka pendek maupun jangka panjang. Keputusan menggunakan utang yang tinggi dapat meningkatkan kualitas nilai perusahaan, karena ada nya pengurangan atas pajak penghasilan (Pratiwi dan Stiawan, 2022). Dalam penelitian ini *leverage* dihitung menggunakan *debt equity ratio* dengan rumus sebagai berikut:

$$DER = \frac{\text{Total Utang}}{\text{Total Ekuitas}}$$

3.1.8 Rencana Analisis Data

3.1.8.1 Analisis Statistik Deskriptif

Menurut (Ghozali, 2018) Analisis statistik deskriptif merupakan suatu cara analisis dengan menggunakan statistik untuk menggambarkan atau mendeskripsikan data sebagaimana adanya. Sehingga analisis statistik deskriptif dapat menjelaskan variabel-variabel yang ada dalam penelitian ini.

3.1.8.2 Uji Pemilihan Model

Regresi data panel merupakan uji yang menggabungkan jenis *cross-section* dan *time series* data. Menurut Wibisono (2005), ada beberapa keuntungan menggunakan regresi data panel. Yang pertama data panel dapat secara eksplisit memperhitungkan heterogenitas individu dengan menggunakan variabel spesifik individu. Yang kedua adalah fakta bahwa metode ini bergantung pada observasi *cross-section* berulang-ulang (*time series*), yang berarti bahwa data panel dapat digunakan untuk menguji dan membuat model perilaku lebih kompleks. Keempat, lebih banyak observasi berarti lebih banyak data yang informatif, lebih variatif, dan kolinearitas antara data yang lebih rendah dan derajat kebebasan yang lebih tinggi. Akibatnya, hasil estimasi yang lebih efektif dapat dicapai. Kelima, data panel dapat digunakan untuk mempelajari model perilaku yang kompleks. Terakhir, bias yang dapat ditimbulkan oleh agregasi data individu dapat dikurangi dengan menggunakan data panel (Agus T.B dan Imammudin Y, 2015).

1. Penentuan model estimasi

Penentuan model estimasi data panel dapat dilakukan melalui beberapa pendekatan yaitu:

1.1 *Common Effect atau Pooled Least Square (PLS)*

Merupakan pendekatan model data panel yang paling sederhana karena hanya mengkombinasikan data time series dan cross section. Pada model ini tidak perhatikan 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 kecil untuk mengestimasi model data panel. Untuk model data panel, sering diasumsikan $\beta_{it} = \beta$ yaitu pengaruh dari perubahan dalam X diasumsikan bersifat konstanta dalam waktu kategori *cross section*. Secara umum, bentuk model linear yang dapat digunakan untuk memodelkan data panel adalah :

$$Y_{it} = X_{it}\beta_{it} + e_{it}$$

Catatan:

Y_{it} : observasi dari unit ke-i dan diamati pada periode waktu ke-t
(yakni variabel dependen yang merupakan suatu data panel)

X_{it} : variabel independen dari unit ke-i dan diamati pada periode waktu ke-t disini diasumsikan X_{it} memuat variabel konstanta

e_{it} : komponen error yang diasumsikan memiliki harga mean 0 dan

variansi homogen dalam waktu serta independen dengan X_{it} .

1.2 *Fixed effect Model (FEM)*

Model ini mengasumsikan bahwa perbedaan antar individu dapat diakomodasi dari perbedaan intersepnya. Model *Fixed effect* merupakan teknik mengestimasi data panel dengan menggunakan variabel dummy untuk menangkap adanya perbedaan intercept. Intercept antar perusahaan, perbedaan intercept bisa terjadi karena perbedaan budaya kerja, manajerial, dan insentif. Disamping itu, model ini juga mengasumsikan bahwa koefisien regresi tetap antara perusahaan dan waktu.

Pendekatan dengan variabel dummy ini dikenal dengan sebutan *least square dummy variables (LSDV)*. Persamaan *Fixed effect Model* dapat ditulis sebagai berikut :

$$Y_{it} = X_{it}\beta + C_i + \dots + \epsilon_{it}$$

Catatan:

C_i : variabel dummy

1.3 *Random effect Model (REM)*

Model ini mengestimasi data panel dimana variabel gangguan mungkin saling berhubungan antar waktu dan antar individu. Pada model *Random effect* perbedaan *intercept* diakomodasi oleh error terms masing-masing perusahaan. Keuntungan menggunakan model *Random effect* yakni menghilangkan heteroskedastisitas. Model ini juga disebut

dengan teknik *Generalized Least Square* (GLS). Sebagai estimatornya, berikut bentuk persamaannya adalah:

$$Y_{it} = X_{it}\beta + V_{it}$$

$$\text{Dimana } V_{it} = C_i + D_i + \varepsilon_{it}$$

C_i diasumsikan bersifat *independent and identically distributed* (iid) normal dengan mean 0 dan variansi σ^2_c (komponen cross section), D_i diasumsikan bersifat iid normal dengan mean 0 dan variansi σ^2_d (komponen time series error). ε_{it} diasumsikan bersifat iid dengan mean 0 dan variansi σ^2_e

2. Tahapan analisis data

2.1 Uji Chow

Uji chow adalah pengujian untuk menentukan model apa yang akan dipilih antara *Common Effect Model* atau *Fixed Effect Model*. Apabila nilai F hitung lebih besar dari F kritis maka H_0 ditolak yang artinya model yang tepat untuk regresi data panel adalah model *Fixed Effect*. Hipotesis yang dibentuk dalam Uji Chow adalah sebagai berikut :

H_0 : *Common Effect Model*

H_1 : *Fixed Effect Model*

2.2 Uji Hausman

Uji Hausman adalah pengujian statistik untuk memilih apakah model *Fixed Effect* atau *Random Effect* yang paling tepat digunakan. Apabila nilai statistik Hausman lebih besar dari nilai kritis *Chi-Squares* maka artinya model yang tepat untuk regresi data panel adalah model

Fixed Effect. Hipotesis yang dibentuk dalam Hausman test adalah sebagai berikut:

H_0 : *Random Effect Model*

H_1 : *Fixed Effect Model*

2.3 Uji Lagrange Multiplier

Uji *Lagrange Multiplier* adalah pengujian statistik untuk mengetahui apakah model random effect lebih baik dari pada metode *common effect*. Apabila nilai LM hitung lebih besar dari nilai kritis *Chi-Squares* maka artinya model yang tepat untuk regresi data panel adalah model *Random Effect*. Hipotesis yang dibentuk dalam LM test adalah sebagai berikut :

H_0 : *Common Effect Model*

H_1 : *Random Effect Model*

3.1.8.3 Uji Asumsi Klasik

Uji asumsi klasik merujuk pada serangkaian tes yang dilakukan untuk memverifikasi apakah suatu model regresi yang digunakan sudah memenuhi asumsi-asumsi dasar yang diperlukan supaya hasil dari analisis menjadi valid. Tujuan dari uji asumsi klasik ini adalah untuk memberikan kepastian bahwasanya regresi yang didapatkan sudah tepat, tidak bias dan konsisten. Dalam uji asumsi klasik ini terdapat empat yang harus dipenuhi yaitu uji normalitas, heteroskedastisitas, tidak adanya autokorelasi dan tidak adanya multikolinearitas.

a. Uji Normalitas

Merupakan uji asumsi klasik yang paling fundamental, dimana analisis normalitas menentukan apakah distribusi nilai residual itu normal atau tidak. Nilai residual yang berdistribusi normal merupakan tanda model regresi yang baik. Oleh karena itu, uji normalitas tidak dilakukan pada masing-masing variabel penelitian, tetapi hanya pada nilai residualnya.

Data dikatakan berdistribusi normal apabila data yang berupa titik menyebar disekitar garis diagonal dan mengikuti arah diagonal. Sedangkan dikatakan tidak berdistribusi normal apabila data menyebar jauh dari arah garis atau tidak mengikuti arah diagonal

b. Uji Multikolinearitas

Dalam asumsi multikolinearitas, kita melihat apakah ada korelasi yang tinggi antara variabel bebas dalam model regresi linear berganda. Jika ada korelasi yang tinggi, maka hubungan antara variabel bebas dan variabel terikatnya terganggu

Menurut (Ghozali, 2018) Untuk mengetahui apakah ada atau tidaknya gejala multikolinearitas, kita harus melihat besaran nilai faktor inflasi variasi (VIF) dan nilai toleransi. Nilai VIF dibawah 10,00 dan nilai toleransi di atas 0,10 digunakan untuk menunjukkan adanya gejala multikolinearitas.

c. Uji Autokorelasi

Merupakan uji asumsi yang bertujuan untuk melihat apakah terjadi korelasi antara suatu periode (t) dengan periode sebelumnya ($t-1$). Dalam penelitian ini untuk mengetahui ada tidaknya autokorelasi menggunakan Uji Durbin-Watson dengan kriteria pengambilan keputusan dalam pengujian ini adalah sebagai berikut:

- a. Menghitung nilai dl dan du dari t-tabel berdasarkan jumlah sampel penelitian.
- b. Menghasilkan grafik untuk menentukan apakah data penelitian menunjukkan masalah autokorelasi atau tidak

Kriteria DW tabel dengan tingkat signifikansi 5% digunakan untuk menentukan ada atau tidaknya autokorelasi dengan rincian sebagai berikut:

- a. Nilai D-W dibawah -2 menunjukkan autokorelasi positif
- b. Nilai D-W diantara -2 dan +2 menunjukkan autokorelasi tidak ada
- c. Nilai D-W diatas +2 menunjukkan autokorelasi negatif.

d. Uji Heteroskedastisitas

Uji heteroskedastisitas digunakan untuk melihat apakah residual dari model yang terbentuk memiliki varians yang konstan atau tidak. Suatu model yang baik adalah model yang memiliki varians dari setiap gangguan atau residualnya konstan. Dampak heteroskedastisitas yaitu tidak efisiennya proses estimasi, sementara hasil estimasinya tetap konsisten dan tidak bias. Uji heteroskedastisitas dalam penelitian dilakukan dengan menggunakan uji Glejser. Uji Glejser adalah uji hipotesis untuk

mengetahui apakah sebuah model regresi memiliki indikasi heteroskedastisitas dengan cara meregresikan nilai absolut residual terhadap variabel independen. Menurut Napitupulu (2021: 191) dasar pengambilan keputusan sebagai berikut:

1. Jika nilai probabilitas $> 0,05$ maka tidak ada masalah heteroskedastisitas.
2. Jika nilai probabilitas $< 0,05$ maka terdapat masalah heteroskedastisitas.

3.1.8.4 Analisis Regresi Data Panel

Merupakan suatu model regresi linear dengan melibatkan lebih dari satu variabel. Biasanya analisis ini digunakan untuk mengetahui pengaruh dari variabel bebas terhadap variabel terikat. Dalam penelitian ini regresi linear berganda dapat dirumuskan sebagai berikut:

$$Y_{it} = \alpha_0 + \beta_1 \text{ETR}_{it} + \beta_2 \text{BTD}_{it} + \beta_3 \text{DER}_{it} + \varepsilon_{it}$$

Keterangan:

Y_{it} : Nilai Perusahaan

α_0 : Konstanta

$\beta_1 \text{ETR}_{it}$: *Tax Avoidance (Effective Tax Rate)*

$\beta_2 \text{BTD}_{it}$: *Tax Aggressiveness (Book Tax Difference)*

$\beta_3 \text{DER}_{it}$: *Leverage (Debt Equity Ratio)*

ε_{it} : Variabel Pengganggu (error) untuk individu ke-i periode ke-t

i : Perusahaan

t : Waktu

3.1.8.5 Uji Hipotesis

1. Koefisien Determinasi (R^2)

Koefisien determinasi biasanya digunakan untuk mengukur seberapa besar variabel independen dapat menjelaskan variabel dependennya. Nilai koefisien determinasi terletak antara 0 dan 1 ($0 < R^2 < 1$), dimana semakin besar nilai R^2 suatu regresi atau nilainya mendekati 1, maka hasil regresi tersebut semakin baik. Tidak ada hubungan antar variabel independen dengan variabel dependen jika $R^2=0$, tetapi terdapat hubungan sempurna jika koefisien determinasi $R^2=1$.

2. Uji Signifikansi Parsial (Uji t)

Menurut (Ghozali, 2018), uji statistik t digunakan untuk menunjukkan seberapa jauh satu variabel independen secara individual dalam menerangkan variabel dependen. Uji statistik t digunakan untuk menguji signifikansi koefisien variabel independen dalam memprediksi variabel dependen bebas. Penelitian ini menggunakan tingkat signifikansi sebesar 0,05 atau ($\alpha=5\%$). Kriteria penerimaan dan penolakan hipotesis yaitu sebagai berikut:

- a. Apabila nilai signifikansi $t < 0.05$, berarti variabel independen secara parsial berpengaruh terhadap variabel dependen.
- b. Apabila nilai signifikansi $t > 0.05$, berarti variabel independen secara parsial tidak berpengaruh terhadap variabel dependen.

3. Uji Signifikansi Simultan (Uji F)

Uji simultan atau F dilakukan dengan tujuan untuk mengukur jika semua variabel independen yaitu perencanaan pajak, penghindaran pajak yang diteliti secara bersama-sama memberikan pengaruh yang signifikan atau tidak pada variabel dependen yaitu nilai perusahaan. Menurut (Ghozali, 2018) uji statistik F pada dasarnya menunjukkan bahwa semua variabel independennya mempunyai pengaruh bersama-sama terhadap variabel dependen. Kriteria untuk pengambilan keputusan dalam uji ini menggunakan quick look yang berarti H_0 dapat ditolak pada tingkat kepercayaan 5%, apabila nilai F lebih besar daripada 4 dan membandingkan nilai F hitung dengan F tabel yang berarti apabila nilai F hitung $> F$ tabel, maka H_0 ditolak dan menerima H_A .