

CHAPTER III

RESEARCH METHODOLOGY

3.1 Type of Research

This study employs a verificative research design, which aims to explain the relationships among variables. In this research, the independent variables consist of Return on Assets (ROA), Current Ratio (CR), and Debt to Equity Ratio (DER), while the dependent variable is stock price.

The method applied is an explanatory survey, which is used to test hypotheses by explaining the relationships among the variables under investigation.

The research technique utilized is quantitative statistical analysis, focusing on numerical data related to amounts, levels, comparisons, and volumes. The development of each variable and the influence among variables are analyzed and validated using statistical methods.

3.2 Research Object, Unit of Analysis, and Research Location

The object of this research refers to the variables or aspects that become the central focus of the study. In this research, the object examined is the effect of Return on Assets (ROA), Current Ratio (CR), and Debt to Equity Ratio (DER) on the stock prices of Food and Beverage sub-sector companies listed on the Indonesia Stock Exchange (IDX).

The unit of analysis used in this study is the organization, in which the research data are obtained from the financial statements and annual reports of Food and Beverage sub-sector companies listed on the Indonesia Stock Exchange for the period 2019–2023.

The research location refers to the place where the research variables are analyzed or where the units of analysis are situated. This study focuses on Food and Beverage sub-sector companies listed on the Indonesia Stock Exchange (IDX) during the period 2019–2023. The Indonesia Stock Exchange is located at Jl. Jendral Sudirman Kav. 52–53, Senayan, Kebayoran Baru, RT 05/RW 03, South Jakarta 12190, Indonesia. To collect the required data and information, the researcher utilized sources available on the official website of the Indonesia Stock Exchange.

3.3 Types and Sources of Data

The type of data used in this research is quantitative data in the form of panel data. The quantitative data consist of the financial statements of Food and Beverage sub-sector companies for a five-year period, from 2019 to 2023. Panel data are a combination of time-series data and cross-sectional data.

The data used in this study are secondary data, which are obtained indirectly by the researcher. Generally, secondary data are gathered from various sources such as mass media, data providers, stock exchanges, previous research, and statistical software. In this research, the data were collected from the official website of the Indonesia Stock Exchange (www.idx.co.id) and the official Yahoo Finance website (www.yahoofinance.com).

3.4 Operational Definition of Variables

Operational definitions are required to determine the indicators, measurements, and data scales of the variables used in the study. In this research, the variables are defined as follows:

1. Dependent Variable

The dependent variable represents the outcome or effect that is influenced by the independent variables. It is observed and measured to determine the extent of the influence exerted by the independent variables. In this study, the dependent variable is Stock Price.

2. Independent Variable

Independent variables are variables that influence or cause changes in the dependent variable. They are used to examine the extent of their impact on the dependent variable. The independent variables used in this study include Return on Assets (ROA), Current Ratio (CR), and Debt to Equity Ratio (DER).

The detailed explanation and measurement of these operational variables are presented in the following table:

Table 3. 1 Operational Definition of Variables

Type of Variable	Indicators	Measurement	Measurement Scale
Independent Variable			
Return on Asset (X_1)	1. Net Income 2. Total Assets	$Return\ on\ Assets = \frac{Net\ Income}{Total\ Assets} \times 100\%$	Ratio
Current Ratio (X_2)	1. Current Assets 2. Current Liabilities	$Current\ Ratio = \frac{Current\ Assets}{Current\ Liabilities} \times 100\%$	Ratio
Debt to Equity Ratio (X_3)	1. Total Liabilities 2. Equity	$Debt\ to\ Equity\ Ratio = \frac{Total\ Liabilities}{Equity} \times 100\%$	Ratio
Dependent Variable			
Stock Price (Y)	Stock Price After the Publication of Financial Statements	$\begin{aligned} & \text{Stock Price} \\ & \text{Closing Price at the End of Jan}_{t+1} + \\ & \text{Closing Price at the End of Feb}_{t-1} + \\ & \text{Closing Price at the End of Mar}_{t+1} \\ & = \frac{\quad}{3} \end{aligned}$	Ratio

3.5 Sampling Method

In this research, the sampling method applied is a non-probability sampling technique using purposive sampling. This technique allows the researcher to select samples based on specific criteria or considerations that align with the research objectives.

From a total population of 30 companies, this study uses a sample of 20 Food and Beverage sub-sector companies listed on the Indonesia Stock Exchange (IDX) during the 2019–2023 period. The sample selection is based on the following predetermined criteria:

1. Companies classified under the Food and Beverage sub-sector and listed on the Indonesia Stock Exchange throughout the entire research period, 2019–2023.
2. Food and Beverage sub-sector companies that have been publicly listed (Initial Public Offering/IPO) for more than 5 years.
3. Food and Beverage sub-sector companies that did not experience losses during the 2019–2023 period.
4. Food and Beverage sub-sector companies that were not delisted during the 2019–2023 period.
5. Food and Beverage sub-sector companies that present their financial statements in Indonesian Rupiah (Rp) and provide complete and consistent financial reports for all five years from 2019 to 2023.

The number of population and sample that meet the established criteria is presented by the researcher in Table 3.2, as follows:

Table 3. 2 Sample Selection Procedures

No.	Code	Company Name	Criteria					Total
			1	2	3	4	5	
1	ADES	Akasha Wira International Tbk.	√	√	√	√	√	√
2	AISA	FKS Food Sejahtera Tbk.	√	√	-	√	√	-
3	ALTO	Tri Banyan Tirta Tbk.	√	√	-	√	√	-
4	BTEK	Bumi Teknokultural Unggul Tbk.	√	√	-	√	√	-
5	BUDI	Budi Strach & Sweetener Tbk.	√	√	√	√	√	√
6	CAMP	Campina Ice Cream Industry Tbk.	√	√	√	√	√	√
7	CEKA	Wilmar Cahaya Indonesia Tbk.	√	√	√	√	√	√
8	CLEO	Sariguna Primatirta Tbk.	√	√	√	√	√	√
9	COCO	Wahana Interfood Nusantara Tbk.	√	√	√	√	√	√
10	DLTA	Delta Djakarta Tbk.	√	√	√	√	√	√
11	FOOD	Sentra Food Indonesia Tbk.	√	√	-	√	√	-
12	GOOD	Garudafood Putra Putri Jaya Tbk.	√	√	√	√	√	√
13	HOKI	Buyung Poetra Sembada Tbk.	√	√	√	√	√	√
14	ICBP	Indofood CBP Sukses Makmur Tbk.	√	√	√	√	√	√
15	IIKP	Inti Agri Resources Tbk.	-	√	-	√	-	-

No.	Code	Company Name	Criteria					Total
			1	2	3	4	5	
16	INDF	Indofood Sukses Makmur Tbk.	√	√	√	√	√	√
17	KEJU	Mulia Boga Raya Tbk.	√	√	√	√	√	√
18	MGNA	Magna Investama Mandiri Tbk.	-	√	-	√	-	-
19	MLBI	Multi Bintang Indonesia Tbk.	√	√	√	√	√	√
20	MYOR	Mayora Indah Tbk.	√	√	√	√	√	√
21	PANI	Pantai Indah Kapuk Dua Tbk.	√	√	-	√	√	-
22	PCAR	Prima Cakrawala Abadi Tbk.	-	√	-	√	√	-
23	PSDN	Prasidha Aneka Niaga Tbk.	√	√	-	√	√	-
24	PSGO	Palma Serasih Tbk.	√	√	-	√	√	-
25	ROTI	Nippon Indosari Corpindo Tbk.	√	√	√	√	√	√
26	SKBM	Sekar Bumi Tbk.	√	√	√	√	√	√
27	SKLT	Sekar Laut Tbk.	√	√	√	√	√	√
28	STTP	Siantar Top Tbk.	√	√	√	√	√	√
29	TBLA	Tunas Baru Lampung Tbk.	√	√	√	√	√	√
30	ULTJ	Ultra Jaya Milk Industry and Trading Company Tbk.	√	√	√	√	√	√

Source: www.idx.co.id, data processed by the author (2024)

The samples selected by the researcher based on the predetermined criteria are as follows:

Table 3. 3 Research Sample

No.	Code	Company Name
1	ADES	Akasha Wira International Tbk.
2	BUDI	Budi Strach & Sweetener Tbk.
3	CAMP	Campina Ice Cream Industry Tbk.
4	CEKA	Wilmar Cahaya Indonesia Tbk.
5	CLEO	Sariguna Primatirta Tbk.
6	COCO	Wahana Interfood Nusantara Tbk.
7	DLTA	Delta Djakarta Tbk.
8	GOOD	Garudafood Putra Putri Jaya Tbk.
9	HOKI	Buyung Poetra Sembada Tbk.
10	ICBP	Indofood CBP Sukses Makmur Tbk.
11	INDF	Indofood Sukses Makmur Tbk.
12	KEJU	Mulia Boga Raya Tbk.
13	MLBI	Multi Bintang Indonesia Tbk.
14	MYOR	Mayora Indah Tbk.
15	ROTI	Nippon Indosari Corpindo Tbk.
16	SKBM	Sekar Bumi Tbk.
17	SKLT	Sekar Laut Tbk.
18	STTP	Siantar Top Tbk.

No.	Code	Company Name
19	TBLA	Tunas Baru Lampung Tbk.
20	ULTJ	Ultra Jaya Milk Industry and Trading Company Tbk.

Source: Data processed by the author (2024)

3.6 Data Collection Method

The data used in this study consist of secondary data, which are obtained indirectly by the researcher through intermediary sources. In this research, the data were collected from the Indonesia Stock Exchange (IDX), Yahoo Finance, and various scientific journals. The data required in this study include:

1. Accounting data in the form of annual financial statements, obtained from the Indonesia Stock Exchange through its official website (www.idx.co.id) and from the official websites of the respective companies being examined.
2. Stock price data as the dependent variable, specifically the stock prices following the publication of the financial statements. The data used are the average stock prices at the end of January, February, and March, obtained online through the official website www.yahoofinance.com.
3. Supporting references, including scientific journals, books, and other relevant publications used to strengthen the theoretical foundation and support the research analysis.

3.7 Data Processing and Analysis Techniques

This study employs a quantitative analysis method and tests the hypotheses using statistical procedures. The analysis aims to determine the effect of the independent variables Return on Assets (ROA), Current Ratio (CR), and Debt to Equity Ratio (DER) on the dependent variable represented by stock prices. The data processing is conducted using panel data regression analysis with the assistance of EViews 12 software.

3.7.1 Descriptive Statistical Analysis

Descriptive statistical analysis is used to describe or summarize the collected data as they are, without drawing general conclusions or making broader generalizations. Descriptive statistics are appropriate when the researcher intends only to describe the sample data and does not aim to infer conclusions about the entire population from which the sample is drawn. The descriptive statistical measures examined include minimum value, maximum value, mean, standard deviation, and the total number of observations (Sugiyono, 2020).

3.7.2 Panel Data Regression Model

In panel data regression, both the intercept and slope may vary across companies and time periods. According to Basuki & Prawoto (2019), there are three primary approaches used to estimate parameters in panel data models:

1. **Common Effect Model:** This is the simplest method for estimating panel data parameters. Data from the cross-sectional and time-series dimensions are combined into a single dataset without considering differences across individuals or time. The model uses the Ordinary Least Squares (OLS) method as its estimation technique.
2. **Fixed Effect Model:** This model accounts for differences in intercepts across companies by using dummy variables. It assumes that each company has a different intercept, while the intercept remains constant over time. Meanwhile, the slope coefficients are assumed to be constant across both companies and time. Parameter estimation in this model is carried out using the Least Square Dummy Variable (LSDV) technique.
3. **Random Effect Model:** In this approach, parameter estimation allows for correlation across time and individuals within the error components. Differences across companies and time periods are accommodated through the error term. Because the error components are correlated, OLS is not suitable for this estimation. The advantage of the Random Effect Model is its ability to eliminate heteroscedasticity. This model is also known as the Error Component Model (ECM) and is estimated using the Generalized Least Square (GLS) technique.

3.7.2.1 Estimation of the Panel Data Regression Model

To determine the most appropriate estimation model, three diagnostic tests are employed to identify the best panel data estimation technique, namely the Chow Test, the Hausman Test, and the Lagrange Multiplier Test.

1. Chow Test

The Chow Test is used to determine whether the Fixed Effect Model or the Common Effect Model is more appropriate for panel data analysis. This test is performed using EViews 12 software. The decision is based on comparing the probability value of the cross-section Chi-square with the significance level (0.05). If the probability value of the cross-section Chi-square < 0.05 , the Fixed Effect Model is selected. Conversely, if the probability value > 0.05 , the Common Effect Model is selected, and the Hausman Test is not required.

The hypotheses for the Chow Test are as follows:

H_0 : The appropriate model is the Common Effect Model.

H_1 : The appropriate model is the Fixed Effect Model.

Decision criteria:

- a. If the probability value of the cross-section Chi-square > 0.05 , then H_0 is accepted, and the chosen model is the Common Effect Model.
- b. If the probability value of the cross-section Chi-square < 0.05 , then H_0 is rejected, and the chosen model is the Fixed Effect Model.

2. Hausman Test

The Hausman Test is conducted to determine whether the Fixed Effect Model or

the Random Effect Model is more suitable. This test is performed using the Correlated Random Effects Hausman Test.

The hypotheses for the Hausman Test are:

H_0 : : The appropriate model is the Random Effect Model.

H_1 : : The appropriate model is the Fixed Effect Model.

Decision criteria:

- a. If the probability value of the cross-section statistic < 0.05 , then H_0 is rejected, and the Fixed Effect Model is selected.
 - b. If the probability value of the cross-section statistic > 0.05 , then H_0 is accepted, and the Random Effect Model is selected.
3. Lagrange Multiplier (LM) Test

The Lagrange Multiplier Test is used to determine whether the Random Effect Model is more appropriate than the Common Effect Model (OLS). This test is carried out using EViews 12 and is based on the Breusch–Pagan method, which is widely used in panel data regression analysis.

The hypotheses for the LM Test are as follows:

H_0 : The appropriate model is the Common Effect Model.

H_1 : The appropriate model is the Random Effect Model.

Decision criteria:

- a. If the probability value of the cross-section Breusch–Pagan statistic < 0.05 , then H_0 is rejected, and the Random Effect Model is selected.
- b. If the probability value of the cross-section Breusch–Pagan statistic > 0.05 , then H_0 is accepted, and the Common Effect Model is selected.

3.7.3 Unit Root Test (Stationarity Test)

The regression analysis in this study employs panel data; therefore, stationarity testing remains necessary because panel data contain a time-series dimension within each cross-sectional unit. According to Das (2019), a time-series dataset is considered stationary when its statistical properties, such as mean and variance, are constant over time. Non-stationary data exhibit time-varying variance and tend to follow certain patterns, which may lead to spurious relationships in econometric model estimation. In the context of panel data, non-stationarity problems may arise from the time dimension of each cross-sectional unit, making stationarity testing an essential preliminary step prior to conducting panel data regression analysis.

Non-stationarity in time-series data is generally caused by the presence of a unit root. Das (2019) explains that data containing a unit root follow a random walk process, in which the current value of a variable is highly dependent on its value in the previous period. Mathematically, the random walk process can be expressed as follows:

$$\mu_t = \mu_{t-1} + u_t$$

If the variance of the error component is equal to zero, the data are stationary around a deterministic trend. However, if the error variance is greater than zero, the data are non-stationary. Ekananda (2014) further states that the presence of a unit root causes the variance of the data to be non-constant, and shocks occurring in a particular period will have permanent effects on the variable's value in subsequent periods. Such conditions may occur in panel data variables observed over time across individual cross-sectional units.

1. Augmented Dickey-Fuller (ADF) Unit Root Test

The unit root test employed in this study is the Augmented Dickey–Fuller (ADF) test. According to Ekananda (2014), the ADF test is an extension of the Dickey–Fuller test that aims to examine the null hypothesis of the presence of a unit root in a variable by incorporating lagged values of the dependent variable in first-difference form to address potential autocorrelation in the error term. The general model of the ADF test can be expressed as follows:

$$\Delta Y_t = \beta_1 + \beta_2 t + \delta Y_{t-1} + \sum_{i=1}^m \alpha_i \Delta Y_{t-i} + \mu_t$$

The hypotheses tested in the ADF procedure are formulated as follows:

H₀: $\delta = 0$, indicating that the data contain a unit root (non-stationary)

H₁: $\delta \neq 0$, indicating that the data are stationary

According to Das (2019), the decision-making process in the ADF test is conducted by comparing the test statistic with the Dickey–Fuller critical values or by examining the probability value. If the probability value is lower than the significance level ($\alpha = 0.05$), the null hypothesis is rejected, and the data are considered stationary. If the data are not stationary at the level form, a differencing process (first difference) is applied. Ekananda (2014) states that data which become stationary after first differencing are classified as integrated of order one, or I(1), whereas data that are stationary without differencing are classified as I(0).

3.7.4 Classical Assumption Tests

Classical assumption tests are statistical requirements that must be fulfilled in panel data regression analysis. Common classical assumption tests include normality test, multicollinearity test, heteroskedasticity test, and autocorrelation test.

1. Normality Test

According to Machali (2021), the normality test aims to determine whether the residuals in the regression model are normally distributed. A good regression model is characterized by residuals that follow a normal distribution. In EViews software, the Jarque-Bera test is a practical method used to detect the normality of residuals. This test is based on large sample asymptotic assumptions and calculates skewness and kurtosis values. The decision criteria for the Jarque-Bera test are as follows:

a. If the Jarque-Bera probability value $> \alpha$ (0.05), H₀ is accepted, indicating that the

residuals are normally distributed.

- b. If the Jarque-Bera probability value $< \alpha$ (0.05), H_0 is rejected, indicating that the residuals are not normally distributed.

2. Multicollinearity Test

Multicollinearity occurs when a regression model includes more than one independent variable and there is a linear relationship among those variables. The presence of multicollinearity may cause most independent variables to become statistically insignificant in explaining the dependent variable, even though the coefficient of determination remains high. Several methods are commonly used to detect multicollinearity, including the Variance Inflation Factor (VIF) and pairwise correlation analysis (Sakti, 2018). In this study, the decision-making criterion is based on the VIF value. Multicollinearity is indicated when the VIF value > 10 . Conversely, if the VIF value < 10 , the regression model is considered free from multicollinearity.

3. Heteroskedasticity Test

According to Sakti (2018), the heteroskedasticity test is used to determine whether the residuals in the regression model have constant variance. The presence of heteroskedasticity can lead to inaccurate t-tests and F-tests. Common detection methods include graphical analysis, the Park test, the Glejser test, Spearman correlation, the Goldfeld–Quandt test, the Breusch–Pagan test, and the White test. The White test is a widely used method to detect heteroskedasticity, either with or without cross terms. The decision criteria are as follows:

- a. If the chi-square probability value > 0.1 , H_0 is accepted, indicating no heteroskedasticity.
- b. If the chi-square probability value < 0.1 , H_0 is rejected, indicating the presence of heteroskedasticity.

4. Autocorrelation Test

According to Machali (2021), the autocorrelation test is conducted to detect correlation among residuals in a regression model, either across cross-sectional units or over time periods. Autocorrelation occurs when sequential observations are correlated. This issue arises because the residual (error) of one observation is influenced by residuals from other observations. Autocorrelation is commonly found in time-series data, where disturbances affecting an individual or group tend to persist and influence subsequent periods. In this study, autocorrelation is tested using the Breusch–Godfrey method, based on the Obs*R-Square value.

Decision criteria:

- a. If the Breusch-Godfrey probability value > 0.05 , H_0 is accepted, indicating no autocorrelation.
- b. If the Breusch-Godfrey probability value < 0.05 , H_0 is rejected, indicating the presence of autocorrelation.

3.7.5 Panel Data Regression Analysis

According to Das (2019), panel data are a combination of cross-sectional data and time-series data. Cross-sectional data are obtained through surveys or complete enumeration of various units, such as households, companies, or countries, at a specific point in time. This type of data is typically collected over a short period, for example, one year. Time-series data, on the other hand, refer to observations of one or more variables collected periodically over a certain time span.

The choice of panel data in this study is based on the involvement of multiple companies observed over several years. Time-series data are used because this study covers a five-year period from 2019 to 2023. Cross-sectional data are used because the data are collected from multiple companies (pooled). The panel data regression model used in this study, consisting of one dependent variable and three independent variables, is expressed as follows:

$$Y_{it} = a + \beta_1 X_{1it} + \beta_2 X_{2it} + \beta_3 X_{3it} + e_{it}$$

Description of variables:

Y_{it} = Dependent variable (Stock Price)

a = Constant

$\beta_{(1-3)}$ = Regression coefficients, representing the increase or decrease in the dependent variable (Y) resulting from changes in the independent variables (X)

X_{1it} = Return on Assets (ROA) of entity i at time t

X_{2it} = Current Ratio (CR) of entity i at time t

X_{3it} = Debt to Equity Ratio (DER) of entity i at time t

e = Error term

t = Year (time period)

i = Entity (company)

3.7.6 Hypothesis Testing

1. Partial Test (t-Test)

According to Rahmaniar (2019), the partial test (t-test) is used to determine whether each independent variable individually has a significant effect on the dependent variable. The t-test statistic can be calculated using the formula:

$$t = \frac{r_{xy}\sqrt{(n-2)}}{\sqrt{(n-r_x^2)}}$$

The hypotheses for the t-test are as follows:

- If the calculated t-value $>$ t_{table} value at a significance level (α) of 0.05, then H_0 is rejected and H_1 is accepted.
- If the calculated t-value $<$ t_{table} value at significance level (α) of 0.05, then H_0 is accepted and H_1 is rejected.

2. Simultaneous Regression Coefficient Test (F-Test)

The F-test is used to evaluate whether all independent variables collectively have a significant effect on the dependent variable. This test helps to determine whether the

set of independent variables as a whole contributes to explaining variations in the dependent variable within the regression model.

According to Sakti (2018), the F-test serves to test hypotheses regarding regression coefficients simultaneously and to evaluate the overall suitability of the model in explaining the relationship between independent and dependent variables. The F-test is crucial because if the result is insignificant, then the results of individual t-tests may become invalid. The decision criteria for the F-test are as follows:

- a. If the calculated F-value $> F_{\text{table}}$ value or the prob. $F_{\text{-statistik}} < (\alpha) 0,05$, then H_0 is rejected, indicating that the independent variables simultaneously have a significant effect on the dependent variable.
 - b. If the calculated F-value $< F_{\text{table}}$ value or the prob. $F_{\text{-statistik}} > (\alpha) 0,05$, then H_0 is Accepted, indicating that the independent variables collectively do not have a significant effect on the dependent variable.
3. Coefficient of Determination (R^2)

According to Indartini and Mutmainah (2024), the coefficient of determination (R^2) is used to measure the extent to which the independent variables explain variation in the dependent variable. The value of R^2 ranges from 0 to 1 and can also be expressed as a percentage. The closer the R^2 value is to 1, the better the regression model explains the data, as the independent variables provide most of the information needed to predict changes in the dependent variable. Conversely, an R^2 value close to 0 indicates that the model explain only a small portion of the variation, meaning the independent variables contribute minimally in predicting the dependent variable. Thus, R^2 can be interpreted as follows:

- a. R^2 values approaching 1 indicate that the independent variables play a major role in explaining variations in the dependent variable.
- b. R^2 values approaching 0 suggest that the independent variables have limited explanatory power over the variations in the dependent variable.