

BAB III

METODE PENELITIAN

A. Metode Penelitian

1. Jenis Metode Penelitian

Metode yang diterapkan dalam penelitian ini yaitu metode asosiatif dengan pendekatan klausal. Pendekatan ini digunakan untuk mengidentifikasi hubungan antara dua variabel atau lebih dalam suatu penelitian. Menurut Sugiyono (2019:65), penelitian asosiatif adalah suatu pendekatan penelitian yang menitik beratkan pada penggalian hubungan antara dua variabel atau lebih dalam suatu permasalahan penelitian dan hubungannya yang bersifat sebab akibat. Dalam penelitian ini penulis menganalisis uji pengaruh antara variabel yang diteliti yaitu media sosial dan lokasi terhadap keputusan penelitian melalui *brand awareness* sebagai variabel *intervening*.

2. Data dan Sumber Data

a. Data Primer

Data primer dalam penelitian ini merujuk pada data yang diperoleh secara langsung dari sumbernya dengan menggunakan metode penyebaran kuesioner kepada responden. Kuesioner ini berisi serangkaian pertanyaan yang dirancang untuk mengumpulkan informasi yang relevan dengan tujuan penelitian mengenai hal yang berkaitan dengan media sosial, lokasi, *brand awareness*, dan keputusan pembelian.

b. Data Sekunder

Data sekunder dalam penelitian ini merujuk pada data yang diperoleh secara tidak langsung atau melalui pihak lain, seperti studi kepustakaan, jurnal, literatur terkait dengan permasalahan, majalah-majalah perekonomian, dan informasi dokumentasi lainnya yang dapat diakses melalui sistem online. Data sekunder ini merupakan laporan historis yang telah disusun dalam arsip dan dapat dipublikasikan atau tidak dipublikasikan sebelumnya.

B. Variabel Penelitian dan Pengukuran

Penelitian ini bertujuan untuk mengukur dampak dari kualitas informasi, berbagi informasi, keselarasan insentif, dan pengambilan keputusan bersama terhadap kinerja operasional. Variabel yang terdapat dalam penelitian ini terdiri dari variabel independen dan variabel dependen. Variabel independen yang dimasukkan dalam penelitian ini meliputi kualitas informasi, berbagi informasi, keselarasan insentif, dan pengambilan keputusan bersama, sedangkan variabel dependen dalam penelitian ini adalah kinerja operasional.

Dalam konteks penelitian ini, terdapat beberapa variabel yang relevan, yang terdiri dari variabel eksogen (variabel independen X_1 dan X_2), variabel endogen (variabel dependen Y_1 dan Y_2), serta variabel *intervening*. Berikut adalah penjelasan lebih lanjut mengenai variabel-variabel tersebut :

- 1) Variabel eksogen (*exogenous*) atau variabel independent adalah variabel yang memiliki pengaruh terhadap variabel terikat (endogen) dan menjadi

penyebab perubahan pada variabel terikat, baik secara positif maupun negatif. Dalam bahasa Indonesia seringkali disebut variabel bebas, (Sugiyono, 2017:39). Dalam penelitian ini, terdapat dua variabel eksogen, yaitu variabel media sosial (X_1) dan variabel lokasi (X_2).

- 2) Variabel Endogen (*endogenous*) atau variabel dependen adalah variabel-variabel yang dipengaruhi atau menjadi akibat karena adanya variabel bebas. Dalam variabel endogen penelitian ini adalah yaitu keputusan pembelian.
- 3) Variabel *Intervening* adalah variabel yang secara teoritis mempengaruhi hubungan antara variabel independen dan dependen, namun tidak terlihat secara langsung dan tidak dapat diukur. Variabel ini berfungsi sebagai perantara antara variabel independen dan dependen, sehingga variabel independen tidak langsung mempengaruhi perubahan atau timbulnya variabel dependen. (Sugiyono, 2019:39). Variabel *intervening* dalam penelitian ini yaitu *brand awareness*.

C. Populasi dan Sampel

1. Populasi

Menurut Sugiyono (2016:80), populasi adalah wilayah generalisasi yang terdiri dari objek atau subjek dengan kualitas dan karakteristik tertentu yang ditetapkan oleh peneliti untuk dipelajari dan kemudian ditarik kesimpulannya. Populasi juga dapat diartikan sebagai himpunan semua unsur atau elemen yang ingin diketahui dalam penelitian.

Dengan demikian, populasi adalah keseluruhan objek penelitian. Populasi dalam penelitian ini adalah orang – orang berdomisili Kota Bogor yang pernah mengunjungi The Gade Coffee and Gold.

2. Sampel

Menurut Sugiyono, (2017 : 81) sampel adalah sebagian dari populasi yang digunakan sebagai sumber data dalam penelitian. Populasi, di sisi lain, merujuk pada keseluruhan individu atau elemen yang memiliki karakteristik yang ingin diteliti. Sedangkan menurut Arikunto (2019:09), sampel merupakan sebagian atau wakil dari populasi yang menjadi objek penelitian. Dengan demikian, dapat disimpulkan bahwa sampel merupakan sekelompok individu yang mewakili secara representatif populasi yang akan diteliti.

Menentukan besarnya sampel dalam penelitian adalah suatu proses yang penting dan harus dilakukan secara hati-hati. Besarnya sampel dapat dipengaruhi oleh beberapa faktor, termasuk tujuan penelitian, jenis populasi, tingkat ketelitian yang diinginkan, dan metode analisis yang akan digunakan. Untuk menentukan besarnya sampel bisa menggunakan pendekatan statistik dimana penggunaan rumusnya untuk menghitung besar sampel yang diperlukan berdasarkan tingkat ketelitian atau tingkat kepercayaan yang diinginkan, tingkat keberagaman dalam populasi, dan *margin of error* (marginal kesalahan) yang dapat diterima. Selain itu terdapat pendekatan berbasis literatur yang melibatkan studi sebelumnya dengan topik atau populasi yang serupa. Peneliti dapat menggunakan

informasi dari studi sebelumnya untuk menentukan besarnya sampel yang relevan dan representatif.

Selain itu, penting juga untuk memperhatikan bahwa besarnya sampel harus memadai untuk menghasilkan generalisasi yang dapat diandalkan. Semakin besar sampel, semakin besar kemungkinan untuk mendapatkan hasil yang mewakili populasi secara keseluruhan. Namun, terkadang sampel yang lebih kecil dengan representasi yang baik juga dapat memberikan hasil yang cukup valid. Oleh karena itu, pemilihan besarnya sampel harus mempertimbangkan kebutuhan penelitian secara holistik dengan memperhatikan faktor-faktor yang relevan.

Dalam penelitian ini menggunakan *Structural Equation Modelling* dengan metode estimasi yang digunakan adalah *Maximum Likelihood*. Besarnya sampel memiliki peran penting dalam interpretasi SEM. Dengan metode estimasi menggunakan *Maximum Likelihood* (ML), minimum diperlukan sampel 100 sampai 200 responden. Menurut Hair, et.al (2014: 573), ada beberapa panduan yang dapat digunakan untuk menentukan ukuran sampel dalam analisis *Structural Equation Modelling* (SEM). Pertama, ukuran sampel yang disarankan adalah 100-200 untuk teknik estimasi maximal likelihood (ML). Namun, hal ini juga tergantung pada jumlah parameter yang diestimasi. Pedomannya adalah 5-10 kali jumlah parameter yang diestimasi.

Selain itu, ukuran sampel juga bergantung pada jumlah indikator yang digunakan dalam semua variabel yang dibangun. Jumlah sampel

seharusnya sebesar jumlah indikator variabel yang dibangun, dikali dengan 5 hingga 10. Misalnya, jika terdapat 20 indikator, maka ukuran sampel yang diperlukan adalah antara 100-200. Jika sampel yang tersedia sangat besar, peneliti dapat memilih teknik estimasi tersebut. Maka dengan demikian sampel dalam penelitian ini, peneliti memutuskan sampel yang digunakan sebanyak 210 responden.

Pada penelitian ini menggunakan pendekatan *non probability sampling* yaitu teknik yang tidak memberikan peluang atau kesempatan yang sama bagi setiap elemen atau anggota populasi untuk dipilih sebagai sampel. Pemilihan sampel didasarkan pada kriteria subjektif yang telah direncanakan oleh peneliti. Teknik ini digunakan ketika tujuan penelitian hanya untuk mendeskripsikan suatu objek penelitian tanpa melakukan generalisasi terhadap populasi. Dengan salah satu metodenya adalah *accidental sampling*. Menurut Sugiyono (2019), *accidental sampling* merupakan suatu teknik penentuan sampel yang didasarkan pada kebetulan. Dalam teknik ini, siapa pun yang secara kebetulan bertemu dengan peneliti dapat digunakan sebagai sampel, asalkan mereka sesuai dengan kriteria yang telah ditentukan oleh peneliti.

D. Teknik Pengumpulan Data

Untuk memperoleh data dalam penelitian ini, penulis mengadakan wawancara dan dokumentasi untuk mengetahui fakta dan informasi yang dibutuhkan oleh peneliti dan menyebarkan kuesioner terhadap masyarakat terutama kaum millenials berdomisili Kota Bogor.

- 1) Kuesioner digunakan oleh peneliti untuk menyebarkan angket berisi pernyataan kepada masyarakat, terutama kaum milenials yang berdomisili di Kota Bogor. Tujuan kuesioner ini adalah untuk memahami bagaimana media sosial dan lokasi mempengaruhi keputusan pembelian pada The Gade Coffee and Gold Bogor.
- 2) Wawancara yaitu teknik pengumpulan data yang dilakukan secara langsung kepada masyarakat yang berdomisili di Kota Bogor secara random.
- 3) Dokumentasi adalah metode pengumpulan data yang melibatkan penggunaan catatan atau dokumen yang tersedia di lokasi penelitian.

E. Instrumen Penelitian

Metode penelitian yang digunakan dalam studi ini adalah kuesioner. Kuesioner yang digunakan adalah jenis kuesioner langsung dan tertutup, yang berarti angket tersebut langsung diberikan kepada responden dan mereka dapat memilih satu pilihan jawaban dari beberapa alternatif yang disediakan. Dalam penelitian ini, jawaban yang diberikan oleh responden akan diberi skor yang mengikuti skala Likert.

F. Operasional Variabel

Penyusunan operasional variabel dapat mengacu pada satu atau lebih referensi yang didukung dengan alasan penggunaan definisi tersebut. Variabel penelitian harus dapat diukur menggunakan skala ukuran yang umum digunakan. Untuk memberikan gambaran yang lebih terperinci tentang variabel penelitian, tabel operasional variabel berikut ini menjelaskan secara rinci :

1. Definisi Operasional Variabel

a. Keputusan Pembelian

Keputusan pembelian adalah suatu proses dimana konsumen dianggap sebagai individu yang secara rasional memaksimalkan utilitas dan keuntungan dalam memilih produk yang akan dibeli. Yang dimana hal tersebut dapat diukur dengan : Pilihan produk, Pilihan merek, Pilihan saluran pembelian, Waktu pembelian, dan Jumlah saluran pembelian. Dalam hal ini terdiri dari 8 item pertanyaan dengan skala likert 1-5.

b. *Brand Awareness*

Brand awareness adalah kemampuan konsumen untuk mengenali dan mengingat suatu merek dalam konteks kategori produk. Hal ini dapat diukur dengan : Recall, Recognition, Purchase, dan Consumption. Dalam hal ini terdiri dari 6 item pertanyaan dengan skala likert 1-5.

c. Lokasi

Lokasi adalah tempat dimana suatu usaha atau aktifitas usaha dilakukan dan keputusan pemilihan lokasi menjadi menjadi salah satu hal yang penting untuk menjadi faktor utama dalam kelangsungan suatu usaha. Yang dimana hal tersebut dapat diukur dengan : Aksesibilitas, Keberadaan Kompetitor, Demografi, Infrastruktur, dan Daya Tarik dan Citra. Dalam hal ini terdapat 8 item pertanyaan dengan skala likert 1-5.

d. Pemasaran Media Sosial

Pemasaran melalui media sosial adalah proses yang mendorong individu untuk mempromosikan situs web, produk, atau layanan merek melalui

saluran sosial *online*, serta berkomunikasi dengan memanfaatkan komunitas yang lebih luas, yang memiliki potensi pemasaran yang lebih besar daripada melalui saluran periklanan tradisional. Hal tersebut dapat diukur dengan : Entertainment, Interaction, Trendiness, Customization, dan Perceived Risk. Dalam hal terdapat 8 item pertanyaan dengan skala likert 1-5.

Tabel 2
Operasional Konstruk

Konstruk	Indikator Konstruk	Kode Indikator	Skala
Media Sosial Instagram	<i>Entertainment</i>		
	1. Saya sering melihat timeline The Gade Coffee and Gold Bogor dan terlihat menarik perhatian di dunia maya (instagram)	MS1	Likert
	2. Saya merasa konten The Gade Coffee and Gold Bogor memiliki konsep yang menarik	MS2	Likert
	<i>Interaction</i>		
	3. Saya mengetahui menu dan promo yang sedang berlangsung di The Gade Coffee and Gold Bogor melalui media sosial instagramnya	MS3	Likert
	4. Saya pernah bertanya pada akun media sosial instagram The Gade Coffee and Gold Bogor terkait menu ataupun promo yang sedang berlangsung	MS4	Likert
	<i>Trendiness</i>		
	5. Saya merasa The Gade Coffee and Gold Bogor sudah tepat dalam menggunakan media sosial instagram yang paling kekinian	MS5	Likert
	<i>Customization</i>		
	6. Saya pernah <i>complain</i> melalui instagram The Gade Coffee and Gold dan direspon dengan solusi yang baik	MS6	Likert
	Perceived Risk		

Konstruk	Indikator Konstruk	Kode Indikator	Skala
	7. Saya merasa website The Gade Coffee and Gold Bogor memahami cara melayani pelanggan dengan baik dan memperhatikan kepuasan pelanggan	MS7	Likert
	8. Saya percaya bahwa website The Gade Coffee and Gold Bogor memberikan saran dan masukan yang baik dalam proses transaksi jual beli produknya dan merchandise yang tersedia	MS8	Likert
Lokasi	Aksesibilitas		
	9. Saya merasa lokasi Menuju The Gade Coffee and Gold Bogor mudah dijangkau	LK9	Likert
	10. Saya merasa kelancaran akses menuju The Gade Coffee and Gold Bogor sangatlah lancar dan dapat dicapai oleh kendaraan umum	LK10	Likert
	Keberadaan Kompetitor		
	11. Saya merasa harga produk The Gade Coffee and Gold lebih terjangkau dibandingkan kompetitor terdekatnya.	LK11	Likert
	Demografi		
	12. Saya merasa lokasi The Gade Coffee and Gold dekat dengan kediaman saya.	LK12	Likert
	Infrastruktur		
	13. Saya merasa tersedianya colokan listrik yang mudah dijangkau untuk mengisi daya perangkat elektronik dan wifi yang mudah diakses sehingga nyaman untuk belajar ataupun bekerja.	LK13	Likert
	14. Saya merasa tempat parkir pada The Gade Coffee and Gold Bogor sangat aman	LK14	Likert
	Daya Tarik dan Citra		
	15. Saya merasa The Gade Coffee and Gold Bogor berada didekat pusat keramaian	LK15	Likert
	16. Saya merasa The Gade Coffee and Gold Bogor saling berdekatan dengan pesaing yang cukup beragam	LK16	Likert
Brand Awareness	Recall		
	17. Nama The Gade Coffee and Gold Bogor mudah untuk diingat dan produk kopi susu gula aren (<i>Van leening Latte</i>) yang	BA1	Likert

Konstruk	Indikator Konstruk	Kode Indikator	Skala
	menjadi ciri khas dan mudah untuk diingat.		
	Recognition		
	18. Saya mengingat slogan “Nabung Emas Dapat Kopi” yang di gunakan The Gade Coffee and Gold	BA2	Likert
	Purchase		
	19. Saya selalu menghabiskan waktu bersama teman-teman di The Gade Coffee and Gold Bogor karena memiliki suasana yang nyaman	BA3	Likert
	20. Saya merasa The Gade Coffee and Gold Bogor menjadi pilihan utama untuk berkumpul bersama teman.	BA4	Likert
	Consumption		
	21. Saya memikirkan The Gade Coffee and Gold Bogor ketika meminum prodak kopi susu gula aren di tempat lain	BA5	Likert
	22. The Gade Coffee and Gold Bogor memiliki nama yang familiar dari coffeeshop lain	BA6	Likert
Keputusan Pembelian	Pilihan produk		
	23. Saya memutuskan untuk melakukan pembelian minuman dan makanan The Gade Coffee and Gold Bogor karena memiliki berbagai macam pilihan minuman non coffee maupun coffee dan berbagai pilihan makanan yang tersedia.	KP7	Likert
	24. Saya memutuskan untuk melakukan pembelian kopi The Gade Coffee and Gold karena memiliki kualitas terbaik dari jenis biji kopi Arabika.	KP8	Likert
	Pilihan Merek		
	25. Saya memutuskan untuk melakukan pembelian produk The Gade Coffee and Gold Bogor karena merek dari <i>beans coffee</i> dari roaster yang terkenal dan banyak yang menggunakannya	KP9	Likert
	Pilihan Saluran Pembelian		
	26. Saya memutuskan melakukan pembelian produk The Gade Coffee and Gold Bogor karena menyediakan berbagai metode pembayaran yang mudah.	KP10	Likert

Konstruk	Indikator Konstruk	Kode Indikator	Skala
	Waktu Pembelian		
	27. Saya membeli banyak produk The GadeCoffee and Gold di waktu tertentu karena terdapat diskon untuk hari tertentu, seperti peringatan di hari besar atau <i>peak season</i>	KP11	Likert
	28. Saat HUT Pegadaian The Gade Coffee and Gold banyak memberikan diskon dengan produk tertentu	KP12	Likert
	Jumlah Pembelian		
	29. Saya membeli produk The Gade Coffee and Gold Bogor sesuai dengan pesanan	KP13	Likert
	30. The Gade Coffee and Gold Bogor menjual berbagai macam produk <i>non coffee</i> dan <i>coffee</i> dengan jenis <i>beans</i> bervariasi. Produk makanan beragam dan banyak pilihan	KP14	Likert

G. Metode Pengambilan Data

Untuk mengumpulkan data dalam studi ini, penulis menggunakan metode penyebaran kuesioner. Kuesioner ini berupa angket yang berisi serangkaian pertanyaan yang diberikan kepada responden. Variabel yang diukur menggunakan skala Likert dengan pilihan sebagai berikut ini:

Tabel 3
Metode Pengambilan Data

Predikat	Nilai
Sangat Setuju (SS)	5
Setuju (S)	4
Netral (N)	3
Tidak Setuju (TS)	2
Sangat Tidak Setuju (STS)	1

Menurut Edward dan Kenney dalam Imam Ghozali (2017:70) skala Likert dapat dianggap sebagai variabel kontinu atau interval, tanpa melanggar asumsi SEM.

H. Metode Analisis Data

Analisis data adalah suatu proses yang melibatkan pengolahan data guna mengidentifikasi informasi yang relevan dan bermanfaat yang dapat digunakan sebagai dasar untuk pengambilan keputusan dalam memecahkan masalah tertentu. Proses analisis data meliputi kegiatan pengelompokan data berdasarkan karakteristiknya, membersihkan data yang tidak akurat atau tidak relevan, melakukan transformasi data, membangun model data, dan mencari informasi penting yang terkandung dalam data tersebut. Tujuan dari analisis data adalah untuk mendapatkan wawasan yang lebih baik tentang masalah yang sedang diteliti dan mengambil keputusan yang lebih efektif berdasarkan pemahaman yang lebih mendalam terhadap data yang ada.

Dalam penelitian ini, Structural Equation Modelling (SEM) digunakan untuk mengolah data karena penelitian ini menggunakan model reflektif. Model reflektif mengacu pada hubungan antara variabel laten dan indikatornya (Ghozali dan Latan, 2020:7). Menurut Ghazali dan Latan (2020:7) analisis SEM biasanya terdiri dari dua sub bab model yaitu model pengukuran yang disebut outer model dan model struktural yang disebut inner model. Model pengukuran menunjukkan bagaimana variabel manifest atau observed variabel merepresentasikan variabel laten untuk diukur. Sedangkan model struktural menunjukkan kekuatan stimulasi antar variabel laten atau konstruk.

Terdapat 6 langkah tahapan dalam permodelan dan analisis persamaan struktural menurut Imam Ghozali (2017:3) yaitu sebagai berikut :

1. Pengembangan Model Berdasarkan Teori

Model SEM merupakan suatu proses pengembangan model yang didasarkan pada justifikasi yang kuat. Setelah model tersebut dibuat, model tersebut kemudian diuji secara empiris melalui program SEM pada populasi yang relevan (Ghozali, 2017: 59). Model persamaan struktural didasarkan pada hubungan kausalitas, di mana diasumsikan bahwa perubahan satu variabel akan berpengaruh pada perubahan variabel lainnya.

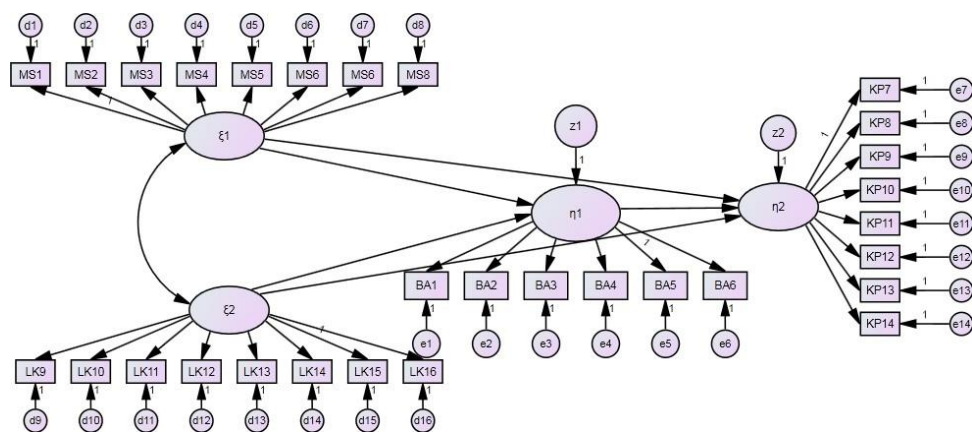
Kevalidan hubungan kausalitas antara dua variabel yang diasumsikan oleh peneliti tidak tergantung pada metode analisis yang dipilih, tetapi terletak pada justifikasi teoritis yang mendukung analisis tersebut. Oleh karena itu, hubungan antar variabel dalam model SEM merupakan deduksi dari teori, dan tanpa dasar teoritis yang kuat, penggunaan SEM tidak dapat dilakukan (Ghozali, 2017: 59).

2. Menyusun Diagram Jalur dan Persamaan Struktural

Dalam penggunaan SEM, langkah-langkah yang perlu dilakukan adalah menggambarkan hubungan kausalitas menggunakan diagram jalur dan menyusun persamaan struktural (Ghozali, 2017: 60). Ada dua hal yang perlu dilakukan dalam proses ini. Pertama, menyusun model struktural dengan menghubungkan antara konstruk laten, baik itu endogen maupun eksogen, untuk membentuk suatu model yang representatif. Kedua, menentukan model dengan menghubungkan konstruk laten atau eksogen

dengan variabel indikator atau manifest yang digunakan (Ghozali, 2017: 60).

Untuk model tersebut dapat dinyatakan dengan bentuk persamaan dengan sebagai berikut :



Gambar 14
Model Persamaan Struktural

Persamaan Struktural

$$\eta_1 = \gamma_{1.1}\xi_1 = \gamma_{1.2}\xi_2 + \zeta_1$$

$$\eta_2 = \gamma_{2.1}\xi_2 + \gamma_{2.2}\xi_2 + \beta_{2.1}\eta_1 + \zeta_2$$

Persamaan Pengukuran Variabel Eksogen

Media Sosial (ξ_1)

$$MS_1 = \lambda_{11}\xi_1 + \delta_1$$

$$MS_2 = \lambda_{21}\xi_1 + \delta_2$$

$$MS_3 = \lambda_{31}\xi_1 + \delta_3$$

$$MS_4 = \lambda_{41}\xi_1 + \delta_4$$

$$MS5 = \lambda_{51} \xi_1 + \delta_5$$

$$MS6 = \lambda_{61} \xi_1 + \delta_6$$

$$MS7 = \lambda_{71} \xi_1 + \delta_7$$

$$MS8 = \lambda_{81} \xi_1 + \delta_8$$

Lokasi (ξ_2)

$$LK9 = \lambda_{92} \xi_2 + \delta_9$$

$$LK10 = \lambda_{102} \xi_2 + \delta_{10}$$

$$LK11 = \lambda_{112} \xi_2 + \delta_{11}$$

$$LK12 = \lambda_{122} \xi_2 + \delta_{12}$$

$$LK13 = \lambda_{132} \xi_2 + \delta_{13}$$

$$LK14 = \lambda_{142} \xi_2 + \delta_{14}$$

$$LK15 = \lambda_{152} \xi_2 + \delta_{15}$$

$$LK16 = \lambda_{162} \xi_2 + \delta_{16}$$

Brand Awareness (η_1)

$$BA1 = \lambda_{11} \eta_1 + \varepsilon_1$$

$$BA2 = \lambda_{21} \eta_1 + \varepsilon_2$$

$$BA3 = \lambda_{31} \eta_1 + \varepsilon_3$$

$$BA4 = \lambda_{41} \eta_1 + \varepsilon_4$$

$$BA5 = \lambda_{51} \eta_1 + \varepsilon_5$$

$$BA6 = \lambda_{61} \eta_1 + \varepsilon_6$$

Kepuasan Pembelian (η_2)

$$KP7 = \lambda_{72} \eta_2 + \varepsilon_7$$

$$KP8 = \lambda_{82} \eta_2 + \varepsilon_8$$

$$KP9 = \lambda_{92} \eta_2 + \varepsilon_9$$

$$KP10 = \lambda_{102} \eta_2 + \varepsilon_{10}$$

$$KP11 = \lambda_{112} \eta_2 + \varepsilon_{11}$$

$$KP12 = \lambda_{122} \eta_2 + \varepsilon_{12}$$

$$KP13 = \lambda_{132} \eta_2 + \varepsilon_{13}$$

$$KP14 = \lambda_{142} \eta_2 + \varepsilon_{14}$$

3. Memilih Jenis Input Matriks dan Estimasi Model yang Diusulkan

Model persamaan struktural memiliki perbedaan dengan teknik analisis multivariat lainnya. Dalam SEM, hanya digunakan data input berupa matriks varian atau kovarian, atau matriks korelasi. Data observasi dapat dimasukkan ke dalam program AMOS, namun program tersebut akan mengubah data mentah menjadi matriks kovarian atau matriks korelasi terlebih dahulu. Analisis terhadap data yang tidak lengkap harus dilakukan sebelum menghitung matriks kovarian atau korelasi (Ghozali, 2017: 61).

Teknik estimasi dalam SEM dilakukan dalam dua tahap. Tahap pertama adalah Estimasi Model Pengukuran, yang digunakan untuk menguji dimensi tunggal dari konstruk eksogen dan endogen dengan menggunakan teknik *Confirmatory Factor Analysis*. Tahap kedua adalah Estimasi Model Persamaan Struktural, yang dilakukan melalui model

lengkap untuk mengevaluasi kesesuaian model dan hubungan kausalitas yang dibangun dalam model tersebut (Ghozali, 2017: 61).

4. Menilai Identifikasi Model Struktural

Selama proses estimasi menggunakan program komputer, seringkali ditemui hasil estimasi yang tidak masuk akal atau tidak memiliki makna, dan hal ini terkait dengan masalah identifikasi model struktural. Problem identifikasi mengacu pada ketidakmampuan model yang diusulkan untuk menghasilkan estimasi unik. Untuk melihat apakah ada masalah identifikasi, kita perlu melihat hasil estimasi yang mencakup :

- a. Adanya masalah identifikasi yang terlihat seperti adanya nilai standar error yang besar untuk 1 atau lebih koefisien.
- b. Ketidakmampuan program untuk membalikkan matriks informasi.
- c. Nilai estimasi yang tidak mungkin memiliki varian kesalahan negatif
- d. Adanya nilai korelasi yang tinggi ($>0,90$) antar estimasi koefisien.

5. Menilai Kriteria *Goodness-Of-Fit*

Pada tahap ini, dilakukan evaluasi terhadap kesesuaian model melalui penilaian terhadap berbagai kriteria *Goodness-of-Fit*. Pengujian prasyarat digunakan untuk memastikan bahwa model struktural telah memenuhi asumsi yang dibutuhkan oleh SEM, menggunakan aplikasi AMOS 24. Selanjutnya, kesesuaian model ditentukan berdasarkan kriteria *Goodness-of-Fit* yang spesifik.

a. Uji Persyaratan Asumsi SEM

1) Uji Normalitas

Pengujian normalitas data dilakukan dengan menghitung distribusi data secara keseluruhan (multivariat). Pengujian tersebut menggunakan perhitungan *critical ratio* (c.r) multivariat. Data tersebut dapat disimpulkan mempunyai distribusi normal *critical ratio skewness value* berkisar antara 1.0 hingga 1.5 atau nilai *critical ratio* (c.r) harus memenuhi syarat $-2,58 < c.r < 2,58$.

2) Uji *Outlier*

Outlier adalah data yang memiliki karakteristik yang sangat berbeda dari observasi lainnya dan muncul dalam bentuk nilai ekstrem, baik itu untuk variabel tunggal maupun kombinasi variabel. Proses deteksi outlier pada data dapat dilakukan dengan menentukan batas nilai yang akan dikategorikan sebagai outlier. Salah satu cara untuk melakukan hal ini adalah dengan mengkonversi nilai data ke dalam skor standar (standardized score) yang sering disebut sebagai *Z-score*.

Outlier dapat diidentifikasi melalui perhitungan jarak Mahalanobis distance, yang kemudian dibandingkan dengan nilai Chi-Square. Selain itu, penilaian juga dapat dilakukan melalui angka p_1 dan p_2 . Jika angka p_1 dan p_2 kurang dari 0,05, maka data dianggap sebagai outlier. Jika nilai jarak Mahalanobis lebih rendah

dari nilai Chi-Square dan semua nilai p^2 lebih besar dari 0,05, maka dapat disimpulkan bahwa tidak terdapat outlier pada data tersebut.

b. Uji Kelayakan Model

Uji kelayakan model dilakukan evaluasi dengan memeriksa berbagai kriteria *Goodness-of-Fit*. Pengujian awal dilakukan untuk memastikan bahwa model struktural telah memenuhi asumsi yang diperlukan oleh SEM dan menentukan apakah model sesuai dengan *criteria goodness of fit* yang spesifik, yang dimana hal tersebut meliputi:

1) *Absolute Fit Measure*

Absolute Fit Measure adalah kriteria yang digunakan untuk mengukur kecocokan keseluruhan model, baik model struktural maupun model pengukuran secara bersama. Kriteria ini mencakup :

a. *Likelihood Ratio Chi Square Statistic*

Uji *Likelihood chi-square* digunakan untuk mengembangkan dan menguji apakah model yang sedang diuji sesuai dengan model yang diestimasi. Jika nilai *Chi-square* yang tinggi relatif terhadap derajat kebebasan, maka hal itu menunjukkan adanya perbedaan yang signifikan antara matriks kovarian atau korelasi yang diamati dengan yang diprediksi. Dalam hal ini, tingkat probabilitas akan lebih rendah daripada tingkat signifikansi yang ditentukan (Ghozali, 2017: 64). Sebaliknya, jika nilai *Chi-square* rendah, maka tingkat probabilitas akan lebih tinggi

daripada tingkat signifikansi yang ditentukan, menunjukkan bahwa tidak ada perbedaan yang signifikan antara matriks kovarian atau korelasi yang diamati dengan yang diprediksi.

Data empiris dikatakan identik dengan teori/model atau tingkat signifikan tinggi apabila tingkat *probability chi-square* >0.05 tingkat signifikansi (0.10). Program AMOS akan memberikan nilai chi-square dengan perintah \cminn dan nilai probabilitas dengan perintah \p, serta besarnya *degree of freedom* dengan perintah \df.

b. *CMIN*

Menggambarkan perbedaan antara *unrectristed sample covarian matrix* S dan *restricted covarian matrix* $\Sigma(\theta)$ atau secara esensi menggambarkan *likelihood ratio test statistic*. Nilai statistik ini dapat dihitung sebagai $(N-1) F_{min}$, di mana N adalah ukuran sampel dikurangi 1, dan F_{min} adalah fungsi kesesuaian minimum. Dengan demikian, nilai *chi-square* sangat sensitif terhadap ukuran sampel yang digunakan.

c. *CMIN/DF*

Nilai χ^2 dapat dibandingkan dengan *degrees of freedom* (df) untuk mendapatkan nilai χ^2 -relatif. ilai χ^2 -relatif yang tinggi menunjukkan adanya perbedaan yang signifikan antara matriks kovarians yang diobservasi dan yang diestimasi. Cara untuk mendapatkan nilai ini adalah dengan membagi CMIN (*The*

minimum sample discrepancy function) dengan *degree of freedom* (df). Dalam hal ini, CMIN/DF sebenarnya merupakan statistik *chi-square* (χ^2) yang dinormalisasi dengan df, sehingga disebut χ^2 -relatif. Jika nilai χ^2 -relatif kurang dari 2.0, itu menunjukkan bahwa ada kesesuaian yang dapat diterima antara model dan data. Program AMOS akan memberikan nilai CMIN/DF dengan perintah `\cmindf`.

d. *GFI*

Goodness of Fit Index (GFI) adalah sebuah ukuran *non-statistic* yang memiliki rentang nilai antara 0 (fit yang buruk) hingga 1.0 (fit yang sempurna). Semakin tinggi nilai *GFI*, semakin baik kesesuaian modelnya. Namun, belum ada standar yang pasti untuk menentukan nilai *GFI* yang dapat diterima sebagai fit yang layak dan belum ada standarnya. Banyak peneliti mengusulkan bahwa nilai *GFI* di atas 90% merupakan indikator kesesuaian model yang baik. Program AMOS akan memberikan nilai *GFI* dengan perintah `\gfi`.

e. *RMSEA*

Root Mean Square Error of Approximation (RMSEA) adalah sebuah ukuran yang bertujuan untuk mengatasi kecenderungan statistik *chi-square* untuk menolak model dengan sampel yang besar. Sebuah model dapat dianggap cocok (*fit*) dan dapat diterima jika nilai *RMSEA* (*Root Mean Square Error of*

Approximation) berada dalam rentang $0,05 < \text{nilai } RMSEA < 0,08$. Hasil uji empiris *RMSEA* cocok untuk melakukan pengujian pada model konfirmatori atau strategi model bersaing dengan menggunakan sampel yang besar. Program AMOS akan memberikan nilai *RMSEA* dengan perintah `\rmsea`.

Adapun pengujian merujuk pada kriteria *model fit* yang terdapat pada tabel *Goodness of Fit* dibawah:

Tabel 4
Goodness of fit

No	Goodness of Fit Indeks	Cut-off Value	Kriteria
1	DF	>0	Over Identified
2	Chi-Square	<a.df	<i>Fit</i>
	Probability	>0,05	<i>Fit</i>
3	CMIN/DF	<2	<i>Fit</i>
4	AGFI	$\geq 0,90$	<i>Fit</i>
5	AGFI	$\geq 0,90$	<i>Fit</i>
6	CFI	$\geq 0,90$	<i>Fit</i>
7	TLI atau NNFI	$\geq 0,90$	<i>Fit</i>
8	IFI	$\geq 0,90$	<i>Fit</i>
9	RMSEA	$\geq 0,08$	<i>Fit</i>

Sumber: Ghozali dan Wijayanto dalam Siswoyo (2017:215)

2) Incremental Fit Measure

Incremental Fit Measure merupakan ukuran perbandingan antara *proposed model* dan *null model*. *Null model* adalah model yang realistis di mana model-model lain harus dibandingkan.

a) TLI

Tucker Lewis Index, juga dikenal sebagai *Nonnormed Fit Index (NNFI)*, adalah sebuah ukuran yang menggabungkan ukuran *parsimony* ke dalam *indeks* perbandingan antara model

yang diajukan dan model nol. Nilai *TLI* berkisar antara 0 hingga 1.0. *TLI* merupakan *indeks kesesuaian inkremental* yang membandingkan model yang sedang diuji dengan model dasar (*baseline*). Nilai *TLI* yang direkomendasikan adalah sama atau ≥ 0.90 . Program *AMOS* akan memberikan nilai *TLI* dengan perintah **\tli**.

b) *CFI*

Comparative Fit Index (CFI) adalah sebuah nilai yang tidak dipengaruhi oleh ukuran sampel sehingga dapat digunakan sebagai acuan untuk mengukur tingkat kesesuaian suatu model penelitian. Nilai *CFI* memiliki rentang antara 0 hingga 1. Semakin mendekati 1, nilai tersebut dikategorikan sebagai tingkat kesesuaian yang sangat baik. Dalam penelitian ini, hasil analisis data menunjukkan nilai *CFI* sebesar 0,989, sehingga dapat disimpulkan bahwa model penelitian ini masuk dalam kategori kesesuaian yang baik.

c) *NFI*

Normed Fit Index atau *NFI* adalah sebuah metrik yang digunakan untuk membandingkan model yang diajukan dengan *null model*. Rentang nilai *NFI* biasanya adalah antara 0 (*no fit at all*) hingga 1,0 (*perfect fit*). Jika nilai *NFI* $\geq 0,90$, maka menunjukkan bahwa model memiliki kesesuaian yang baik

(*good fit*), sedangkan jika nilai *NFI* berada di antara $0,80 \leq 0,90$, sering kali disebut sebagai kesesuaian yang cukup (*marginal fit*).

3) *Measurement Model Fit*

Model pengukuran menggambarkan cara di mana variabel yang dapat diamati (indikator) merefleksikan variabel yang tidak dapat diamati (laten) yang ingin diukur, melalui pengujian validitas dan reliabilitas variabel laten menggunakan analisis faktor konfirmatori. Penelitian ini bertujuan untuk menguji validitas konstruk dengan memperhatikan validitas konvergen. Validitas konvergen dievaluasi melalui analisis SEM (*Structural Equation Modeling*) menggunakan perangkat lunak AMOS, dengan mempertimbangkan nilai faktor loading atau parameter lambda (λ).

Nilai faktor loading yang tinggi menunjukkan bahwa indikator secara konvergen mengarah ke satu titik. Selanjutnya, dalam SEM, terdapat nilai *squared multiple correlations*, yang merupakan kuadrat dari korelasi antara variabel dan indikator yang berkaitan indikatornyanya. Pendekatan untuk menilai measurement model adalah mengukur *composite reliability* dan *variance extracted* untuk setiap konstruk.

Uji reliabilitas dilakukan untuk mengukur konsistensi dan keandalan model. Semakin tinggi reliabilitas pengukuran, semakin besar keyakinan bagi peneliti bahwa semua indikator konsisten dalam pengukuran tersebut. Reliabilitas yang diukur menggunakan

koefisien construct reliability memiliki nilai batas kritis $\geq 0,70$ (menunjukkan reliabilitas yang baik). Namun, jika nilai berada dalam rentang 0,60 - 0,70, masih dapat diterima dengan syarat bahwa validitas indikator dalam model tersebut baik. Sementara itu, koefisien *variance extracted* memiliki batas nilai kritis $\geq 0,50$ (menunjukkan reliabilitas yang baik). Koefisien ini bersifat opsional dalam penelitian.

Uji reliabilitas model dilakukan dengan menghitung *Construct Reliability* (CR) dan *Variance Extracted* (VE) dari nilai *Standardized Factor Loading*, menggunakan rumus berikut :

$$\text{Construct Reliability} = \frac{(\sum \text{Standar loading})^2}{(\sum \text{Standar loading})^2 + \sum \epsilon_j}$$

$$\text{Variance Extracted} = \frac{\sum \text{Standardized loading}^2}{\sum \text{Standardized loading}^2 + \sum \epsilon_j}$$

4) *Multicollinearity and Singularity*

Uji ini berguna untuk mengevaluasi keberadaan multikolinearitas dan singularitas dalam sebuah kombinasi variabel. Tanda-tanda adanya multikolinearitas dan singularitas dapat diidentifikasi melalui nilai determinan matriks kovarian sampel yang sangat kecil atau mendekati nol.

5) *Disriminant Validity*

Discriminant validity digunakan untuk mengevaluasi sejauh mana suatu konstruk berbeda dari konstruk lainnya. Nilai discriminant validity yang tinggi menunjukkan bahwa suatu

konstruk memiliki keunikan dan kemampuan untuk mengukur fenomena yang berbeda. Untuk menguji *discriminant validity*, perbandingan dapat dilakukan antara akar kuadrat dari *Average Variance Extracted (AVE)* dan nilai korelasi antar konstruk.

6. Interpretasi Model dan Modifikasi Model

Setelah model diterima, maka peneliti dapat mempertimbangkan untuk melakukan modifikasi pada model guna meningkatkan penjelasan teoritis atau *goodness-of-fit*. Modifikasi model awal dilakukan setelah mempertimbangkan banyak faktor. Jika terjadi modifikasi pada model, model tersebut harus diuji ulang secara terpisah (*cross-validated*) sebelum model modifikasi dapat diterima. Pengukuran model dapat dilakukan menggunakan *modification indices*. Nilai *modification indices* mengindikasikan penurunan *Chi-squares* yang terjadi jika koefisien diestimasi. Nilai *modification indices* yang sama dengan atau lebih besar dari 3.84 menunjukkan bahwa terjadi penurunan *Chi-squares* secara signifikan.

Dalam Structural Equation Modeling (SEM), terdapat beberapa simbol yang digunakan untuk mewakili pengaruh antar variabel dalam analisis jalur, yaitu:

1. ξ (ksi) : Mewakili variabel laten eksogen.
2. η (eta) : Mewakili variabel laten endogen.
3. λ (lambda) : Hubungna antara variabel laten eksogen ataupun endogen terhadap indikator-indikator.

4. β (beta) : Koefisien pengaruh variabel laten endogen terhadap variabel endogen.
5. γ (gama) : Koefisien pengaruh variabel eksogen terhadap variabel endogen.
6. ϕ (phi) : Koefisien pengaruh variabel eksogen terhadap variabel eksogen.
7. ζ (zeta) : Kesalahan dalam persamaan yaitu antara variabel eksogen dan/atau endogen terhadap variabel endogen.
8. ε (epsilon) : Kesalahan pengukuran variabel manifest untuk variabel laten endogen.
9. δ (delta) : Kesalahan pengukuran variabel manifest untuk variabel laten eksogen.

Simbol- simbol tersebut digunakan untuk menunjukkan pola hubungan antara variabel eksogen dengan variabel endogen maupun dengan indikatornya.