

BAB IV HASIL PENELITIAN

A. Gambaran Umum Tempat Penelitian

1. Kantor Akuntan Publik ASRI

Kantor Akuntan Publik Asri adalah kantor akuntan publik yang menawarkan layanan audit laporan keuangan, asuransi, dan konsultasi akuntansi dan perpajakan. Lokasi KAP Asri berada di Jalan Langgau Lorong 8, Kelurahan Karuwisi Utara, Kecamatan Panakkukang, Kota Makassar, Sulawesi Selatan. Klien KAP Asri termasuk perusahaan swasta dan organisasi lainnya yang membutuhkan jasa audit independen. Yang mana aktivitas tersebut menuntut tingkat ketelitian, kejujuran, dan kepatuhan terhadap standar profesional akuntan publik, sehingga deteksi penipuan keuangan menjadi aspek penting dari proses audit (Juni et al., 2023).

Proses kerja audit KAP Asri juga mengalami perubahan seiring dengan kemajuan teknologi informasi. Ini termasuk penggunaan sistem berbasis data dan teknologi digital untuk membantu analisis, pengujian, dan evaluasi laporan keuangan klien. Kondisi ini menjadikan KAP Asri sebagai lokasi yang relevan untuk meneliti penerimaan Artificial Intelligence (AI) sebagai alat bantu dalam mendeteksi penipuan keuangan dari sudut pandang auditor

dan profesional keuangan yang terlibat langsung dalam proses tersebut (Hartono & Laksito, 2022).

Penelitian ini tidak hanya melibatkan KAP Asri, tetapi juga beberapa kantor akuntan publik lain di Kota Makassar sebagai dukungan untuk mendapatkan gambaran yang lebih luas tentang penerimaan AI di lingkungan profesi akuntan publik. Karakteristik KAP-KAP serupa, yaitu menyediakan jasa audit dan asuransi. Mereka juga menghadapi masalah yang sama terkait peningkatan kompleksitas transaksi dan risiko penipuan keuangan.

2. Karakteristik responden

Penelitian ini menjelaskan karakteristik responden meliputi jenis kelamin, usia, kelompok generasi, dan pengalaman menggunakan AI. Secara singkat karakteristik terponden dapat dilihat pada tabel berikut:

Tabel 1
Karakteristik Responden

No	Karakteristik Responden	Frekuensi	Persentase (%)
1	Jenis Kelamin		
	Laki-laki	40	80%
	Perempuan	10	20%
2	U s i a		
	25 tahun	5	10%

No	Karakteristik Responden	Frekuensi	Persentase (%)
	26 – 28 tahun	31	62%
	29 – 35 tahun	14	28%
3	Kelompok Generasi		
	Generasi Y	14	28%
	Generasi Z	36	72%
4	Pengalaman Menggunakan AI		
	< 1 tahun	7	14%
	1 – 2 tahun	22	44%
	3 tahun	21	42%

Sumber: Data Primer diolah (2026)

Penelitian ini melibatkan 50 responden yang berprofesi di bidang keuangan, seperti akuntan, auditor, analis keuangan, maupun praktisi yang memiliki keterlibatan dalam proses pengelolaan dan pengawasan keuangan. Keterlibatan responden dengan latar belakang profesional di bidang keuangan menjadi penting dalam penelitian ini, mengingat fokus penelitian adalah pada pemanfaatan *Artificial Intelligence* (AI) dalam mendeteksi penipuan keuangan. Dengan demikian, responden yang dipilih memiliki relevansi langsung terhadap konteks penelitian serta diharapkan mampu memberikan penilaian yang objektif berdasarkan pengalaman profesional mereka.

Berdasarkan jenis kelamin, mayoritas responden adalah laki-laki, sebanyak 40%, dan perempuan, sebanyak 10%. Komposisi ini menunjukkan bahwa responden

laki-laki lebih banyak terlibat dalam penelitian. Perbedaan proporsi ini dapat menunjukkan bahwa laki-laki masih mendominasi tenaga kerja di beberapa industri keuangan dan teknologi, terutama yang berkaitan dengan sistem informasi dan analisis data.

Dilihat dari usia, mayoritas responden berada pada rentang usia 26–28 tahun, 31 orang (62 persen), 14 orang (28 persen), dan 5 orang (10 %) berusia 25 tahun. Data ini menunjukkan bahwa sebagian besar responden masih muda. Kelompok usia ini biasanya lebih siap untuk inovasi berbasis sistem cerdas dan lebih adaptif terhadap kemajuan teknologi digital.

Dalam penelitian ini, tidak ada responden dari Generasi X; sebagian besar responden termasuk dalam Generasi Z, atau 36 persen dari peserta, dan 14 persen dari peserta, atau 28 persen. Dominasi Generasi Z dalam kelompok responden menunjukkan bahwa sebagian besar responden adalah generasi yang tumbuh dan berkembang di era digital. Mereka yang termasuk dalam generasi ini lebih akrab dengan teknologi dan lebih terbiasa menggunakan aplikasi dan sistem berbasis kecerdasan buatan dalam aktivitas sehari-hari dan dalam pekerjaan profesional mereka.

Selanjutnya, berdasarkan pengalaman menggunakan AI, sebanyak 22 responden (44%) telah menggunakan AI selama 1 hingga 2 tahun, 21 responden (42%) selama 3 tahun, dan 7 responden (14%) kurang dari 1 tahun. Distribusi ini menunjukkan bahwa mayoritas responden tidak hanya mengenal AI secara teoritis

tetapi juga telah menggunakannya dalam praktik, baik untuk tujuan bisnis maupun pribadi.

Secara keseluruhan, karakteristik responden dalam penelitian ini menunjukkan bahwa sebagian besar merupakan profesional muda di bidang keuangan yang termasuk dalam Generasi Z dengan pengalaman penggunaan AI yang cukup memadai. Kondisi ini mendukung tujuan penelitian, karena responden memiliki latar belakang yang relevan dan pemahaman teknologi yang baik dalam mengevaluasi efektivitas penggunaan AI sebagai alat bantu dalam mendeteksi penipuan keuangan. Dengan demikian, data yang diperoleh dari responden diharapkan mampu memberikan gambaran yang representatif mengenai penerapan AI dalam konteks pengawasan dan pengendalian risiko kecurangan di sektor keuangan.

B. Hasil Penelitian

1. Analisis Statistik Deskriptif

Analisis statistik deskriptif dengan menginterpretasikan nilai rata-rata dari masing-masing indikator pada variabel penelitian ini dimaksudkan untuk memberikan gambaran mengenai indikator - indikator yang membangun konsep model penelitian secara keseluruhan.

Dasar interpretasi nilai rata-rata yang digunakan dalam penelitian ini, mengacu pada interpretasi skor yang digunakan Steven, Jr (2004). sebagaimana digambarkan pada tabel berikut ini:

Tabel 2
Dasar Interpretasi Skor Item Dalam Variabel Penelitian

No.	Nilai Skor	Interpretasi
1	1 - 1,8	Sangat tidak baik/ Sangat rendah
2	1,8 - 2,6	Tidak Baik/ Rendah
3	2,6 – 3,4	Cukup Baik/ Sedang
4	3,4 – 4,2	Baik/ Tinggi
5	4,2 – 5,0	Sangat Baik/ Sangat tinggi

Sumber: Modifikasi dari Steven, Jr (2004)

Uraian dari analisis statistik deskriptif dari masing-masing variabel diuraikan sebagai berikut:

1. Kepercayaan pada Artificial Intelligence

Tabel 3
Nilai Rata-rata jawaban responden per indikator dari variabel Kepercayaan pada Artificial Intelligence (X1)

Indikator	Skor Jawaban Responden										Mea n	Ktgr
	Sangat tidak setuju		Tidak setuju		Netral		Setuju		Sangat setuju			
	F	%	F	%	F	%	F	%	F	%		
Saya percaya bahwa sistem Artificial Intelligence	0	0%	0	0%	7	14 %	22	44 %	21	42 %	3,82	Baik

Indikator	Skor Jawaban Responden										Mea n	Ktgr
	Sangat tidak setuju		Tidak setuju		Netral		Setuju		Sangat setuju			
	F	%	F	%	F	%	F	%	F	%		
mampu mendeteksi penipuan keuangan secara akurat. (X1.1)												
Saya dapat memahami proses analisis yang dilakukan oleh AI dalam mendeteksi penipuan keuangan. (X1.2)	0	0 %	3	6%	8	16 %	31	62 %	8	16 %	3,88	Baik
Saya menilai AI dapat diandalkan untuk digunakan dalam berbagai kondisi dan situasi operasional. (X1.3)	0	0%	0	0%	5	10 %	17	34 %	28	56 %	3,98	Baik
Saya percaya bahwa penggunaan AI tidak menyebabkan kebocoran data keuangan yang bersifat sensitif. (X1.4)	0	0%	0	0%	18	36 %	18	36 %	14	28 %	3,92	Baik
Mean Variabel Kepercayaan pada AI											3,90	Baik

Sumber: Data Primer diolah (2026)

Berdasarkan hasil analisis deskriptif terhadap variabel *Artificial Intelligence* (AI), diperoleh nilai rata-rata (mean) masing-masing indikator yang menunjukkan tingkat persepsi responden terhadap pemanfaatan AI dalam mendeteksi penipuan keuangan. Indikator pertama (X1.1), yaitu kemampuan AI dalam mendeteksi penipuan keuangan secara akurat, memperoleh nilai mean sebesar 3,82. Indikator kedua (X1.2), terkait pemahaman responden terhadap proses analisis yang dilakukan AI, memiliki nilai mean sebesar 3,88. Selanjutnya, indikator ketiga (X1.3), mengenai keandalan AI dalam berbagai kondisi dan situasi operasional, memperoleh nilai mean tertinggi yaitu sebesar 3,98. Sementara itu, indikator keempat (X1.4), yang berkaitan dengan keamanan penggunaan AI dalam mendeteksi kebocoran data keuangan yang bersifat sensitif, memiliki nilai mean sebesar 3,92.

Secara keseluruhan, rata-rata (grand mean) variabel AI adalah sebesar 3,90. Nilai ini berada dalam kategori baik, yang menunjukkan bahwa responden memiliki persepsi positif terhadap penggunaan AI dalam mendeteksi penipuan keuangan.

Nilai rata-rata tertinggi pada indikator keandalan AI (3,98) menunjukkan bahwa responden merasa sistem AI mampu bekerja secara konsisten dan fleksibel dalam menghadapi berbagai situasi serta tingkat kekompleksan dari pola kecurangan yang terus berubah. Sementara itu, meskipun nilai rata-rata indikator akurasi (3,82) adalah yang terendah dibandingkan indikator lainnya, nilai tersebut masih termasuk dalam kategori yang baik. Ini menunjukkan bahwa responden masih percaya bahwa AI mampu mendeteksi penipuan dengan baik, meskipun sistem masih perlu diperbaiki.

Temuan ini secara konseptual menunjukkan bahwa penggunaan AI dianggap sebagai alat yang berguna untuk membantu proses pengawasan dan pengendalian risiko kecurangan di sektor keuangan. Persepsi positif ini menunjukkan bahwa AI juga dianggap sebagai kemajuan teknologi dan sebagai alat strategis yang dapat meningkatkan kualitas deteksi dini praktik penipuan keuangan.

2. Pengalaman Kerja (X2)

Tabel 4

Nilai Rata-rata jawaban responden per indikator variabel Pengalaman Kerja (M)

Indikator	Skor Jawaban Responden										Mea n	Ktgr
	Sangat tidak setuju		Tidak setuju		Netral		Setuju		Sangat setuju			
	F	%	F	%	F	%	F	%	F	%		
Pengalaman kerja saya di bidang keuangan membantu saya dalam memahami proses pengelolaan dan pengawasan keuangan. (X2.1)	0	0%	0	0%	23	46 %	12	24 %	15	30 %	3,84	Baik
Saya memiliki pengalaman yang memadai dalam melakukan analisis risiko atau deteksi	0	0%	0	0%	11	22 %	29	58 %	10	20 %	3,98	Baik

Indikator	Skor Jawaban Responden										Mea n	Ktgr
	Sangat tidak setuju		Tidak setuju		Netral		Setuju		Sangat setuju			
	F	%	F	%	F	%	F	%	F	%		
penipuan keuangan. (X2.2)												
Saya memahami teknik pendeteksian penipuan keuangan yang dilakukan secara manual atau tradisional. (X2.3)	0	0%	0	0%	9	18 %	31	62 %	10	20 %	4,02	Baik
Tugas dan tanggung jawab pekerjaan saya memberikan pengalaman langsung dalam pengawasan dan pengendalian keuangan. (X2.4)	0	0%	0	0%	8	16 %	37	74 %	5	10 %	3,94	Baik
Mean Variabel Pengalaman Kerja											3,94	Baik

Sumber: Data Primer diolah (2026)

Hasil dari analisis deskriptif variabel pengalaman kerja menunjukkan bahwa responden melihat bagaimana pengalaman profesional mereka membantu pengawasan dan pengendalian keuangan. Empat indikator digunakan untuk mengukur variabel ini:

pengalaman dalam pengelolaan, analisis risiko, pemahaman teknik penipuan, dan keterlibatan langsung dalam proses pengawasan.

Pengalaman kerja di bidang keuangan yang membantu dalam pemahaman proses pengelolaan dan pengawasan keuangan, indikator pertama (M), memperoleh nilai mean 3,84 dan termasuk dalam kategori baik. Ini menunjukkan bahwa penilaian responden terhadap pengalaman profesional berperan dalam meningkatkan pemahaman mereka tentang mekanisme pengelolaan keuangan. Indikator kedua (M), pengalaman yang memadai dalam melakukan analisis risiko atau deteksi penipuan keuangan, menerima nilai mean 3,98. Nilai ini menunjukkan bahwa responden merasa memiliki pengalaman yang cukup untuk menemukan potensi risiko dan indikasi kecurangan dalam sistem keuangan.

Selain itu, indikator ketiga (M), yang berkaitan dengan pengetahuan responden tentang pendeteksian penipuan keuangan secara manual maupun tradisional, mendapatkan nilai tertinggi secara keseluruhan sebesar 4,02. Nilai ini menunjukkan bahwa responden memahami teknik pendeteksian kecurangan konvensional dengan baik. Oleh karena itu, ini dapat menjadi dasar untuk mengevaluasi seberapa efektif teknologi berbasis kecerdasan buatan. Indikator keempat (M), yang berkaitan dengan tanggung jawab pekerjaan yang mencakup pengalaman langsung dalam pengawasan dan pengendalian keuangan, memperoleh nilai mean 3,94 dan berada dalam kategori baik; ini menunjukkan bahwa sebagian besar responden terlibat dalam aktivitas pengawasan secara langsung, yang relevan dengan konteks penelitian.

Secara keseluruhan, variabel Pengalaman Kerja memiliki nilai rata-rata atau grand mean 3,94, yang merupakan nilai yang baik. Hasilnya menunjukkan bahwa responden memiliki pengalaman profesional yang cukup dalam bidang keuangan, terutama dalam hal analisis risiko dan pengawasan keuangan. Oleh karena itu, pengalaman kerja responden dapat dianggap sebagai komponen penting dalam menilai dan mengevaluasi penggunaan AI untuk mendeteksi penipuan keuangan.

3. Penerimaan Artificial Intelligence (Y)

Tabel 5
Nilai Rata-rata jawaban responden per indikator variabel Penerimaan Artificial Intelligence (Y)

Indikator	Skor Jawaban Responden										Mea n	Ktgr
	Sangat tidak setuju		Tidak setuju		Netral		Setuju		Sangat setuju			
	F	%	F	%	F	%	F	%	F	%		
Saya termotivasi untuk menggunakan Artificial Intelligence dalam mengidentifikasi penipuan keuangan. (Y.1)	0	0%	0	0%	16	32 %	28	56 %	6	12 %	3,8	Baik
Saya terbiasa mengandalkan AI dalam proses pengambilan	0	0%	0	0%	14	28 %	24	48 %	12	24 %	3,96	Baik

Indikator	Skor Jawaban Responden										Mea n	Ktgr
	Sangat tidak setuju		Tidak setuju		Netral		Setuju		Sangat setuju			
	F	%	F	%	F	%	F	%	F	%		
keputusan terkait risiko dan penipuan keuangan. (Y.2)												
Saya berniat untuk terus menggunakan Artificial Intelligence dalam mendeteksi penipuan keuangan di masa depan. (Y.3)	0	0%	0	0%	10	20 %	31	62 %	9	18 %	3,98	Baik
Mean Variabel Penerimaan Artificial Intelligence											3,91	Baik

Sumber: Data Primer diolah (2026)

Berdasarkan hasil jawaban responden terhadap variabel Penerimaan Artificial Intelligence menunjukkan bahwa responden menerima penggunaan AI untuk mendeteksi penipuan keuangan dengan baik. Tiga indikator menunjukkan motivasi responden, kebiasaan penggunaan, dan niat keberlanjutan penggunaan AI.

Motivasi responden untuk menggunakan AI untuk mengidentifikasi penipuan keuangan (Y.1) memperoleh nilai mean 3,80 dan termasuk dalam kategori baik. Distribusi jawaban menunjukkan bahwa sebagian besar responden berada dalam kategori setuju (56%) dan sangat setuju (12%). Ini menunjukkan bahwa responden

memiliki dorongan yang cukup kuat di dalamnya untuk memanfaatkan AI sebagai alat bantu dalam proses identifikasi kecurangan.

Indikator kedua (Y.2), yang menunjukkan kebiasaan responden untuk bergantung pada AI untuk proses pengambilan keputusan terkait penipuan keuangan dan risiko, memperoleh nilai rata-rata 3,96 dan berada dalam kategori baik. 24% dari responden menyatakan sangat setuju, dan 48% menyatakan setuju. Hasilnya menunjukkan bahwa AI telah dianggap sebagai teknologi yang membantu dan telah mulai terintegrasi dalam proses pengambilan keputusan profesional.

Selain itu, indikator ketiga (Y.3), yang berkaitan dengan niat untuk terus menggunakan kecerdasan buatan untuk mendeteksi penipuan keuangan di masa depan, menerima nilai tertinggi dari semua indikator, dengan nilai rata-rata 3,98. Sebagian besar orang yang menjawab setuju (62%) dan sebagian besar sangat setuju (18%). Hal ini menunjukkan komitmen keberlanjutan terhadap pemanfaatan AI dan menunjukkan betapa populernya teknologi ini.

Secara keseluruhan, variabel Penerimaan *Artificial Intelligence* memiliki nilai rata-rata (grand mean) 3,91 dan masuk dalam kategori baik. Hasilnya menunjukkan bahwa responden secara umum menerima penggunaan AI untuk membantu deteksi penipuan keuangan. Keinginan untuk terus menerus menggunakan AI menunjukkan bahwa teknologi ini benar-benar bermanfaat dan sesuai dengan kebutuhan para profesional keuangan.

Hasil ini mendukung gagasan bahwa AI memiliki peluang untuk berkembang dalam sistem pengawasan dan pengendalian risiko kecurangan, terutama jika persepsi dan kesiapan pengguna positif terhadap teknologi tersebut.

a. Evaluasi Model Pengukuran (Outer Model)

Evaluasi model pengukuran (outer model) dilakukan untuk menilai sejauh mana indikator-indikator yang digunakan dalam penelitian ini mampu merepresentasikan konstruk laten yang diukur. Pada pendekatan Partial Least Squares (PLS), pengujian outer model bertujuan untuk memastikan bahwa setiap indikator memiliki tingkat validitas dan reliabilitas yang memadai sebelum dilakukan analisis lebih lanjut pada model struktural. Dengan demikian, hasil pengujian ini menjadi dasar dalam menentukan kelayakan instrumen penelitian yang digunakan.

Pengujian model pengukuran dalam penelitian ini dilakukan melalui beberapa tahapan, yaitu pengujian validitas konvergen (*convergent validity*), validitas diskriminan (*discriminant validity*), serta uji reliabilitas konstruk. Validitas konvergen digunakan untuk melihat sejauh mana indikator dalam satu konstruk memiliki korelasi yang tinggi satu sama lain. Sementara itu, validitas diskriminan bertujuan untuk memastikan bahwa suatu konstruk benar-benar berbeda dari konstruk lainnya dalam model penelitian. Adapun uji reliabilitas dilakukan untuk mengukur konsistensi internal indikator dalam merepresentasikan konstruk laten.

Melalui serangkaian pengujian tersebut, diharapkan setiap konstruk dalam model penelitian ini memenuhi kriteria statistik yang telah ditetapkan dalam metode PLS. Konstruk yang memenuhi kriteria validitas dan reliabilitas dapat dinyatakan layak untuk dianalisis lebih lanjut pada tahap evaluasi model struktural (inner model). Oleh karena itu, evaluasi outer model menjadi tahapan penting dalam memastikan bahwa hasil analisis yang diperoleh bersifat akurat, konsisten, dan dapat dipertanggungjawabkan secara ilmiah.

1. Convergent Validity

Pengujian validitas konvergen dalam penelitian ini dilakukan dengan menganalisis nilai outer loading masing-masing indikator terhadap konstruk laten yang diukur. Nilai outer loading menunjukkan tingkat korelasi antara indikator dengan variabel konstruknya. Dalam pendekatan Partial Least Squares (PLS), indikator dinyatakan memenuhi kriteria validitas konvergen apabila memiliki nilai outer loading lebih dari 0,70. Namun demikian, dalam penelitian yang bersifat pengembangan model, nilai di atas 0,60 masih dapat diterima sepanjang konstruk tetap memenuhi kriteria validitas dan reliabilitas secara keseluruhan.

Tabel 4 .1 Convergent Validity (Outer Loadings)

	Outer loadings
AIA1 <- Acceptence AI	0.741
AIA2 <- Acceptence AI	0.782
AIA3 <- Acceptence AI	0.807

TAI1 <- Trust in AI	0.889
TAI2 <- Trust in AI	0.795
TAI3 <- Trust in AI	0.763
TAI4 <- Trust in AI	0.938
WE1 <- Work Experience	0.874
WE2 <- Work Experience	0.831
WE3 <- Work Experience	0.732
WE4 <- Work Experience	0.595

Sumber : Hasil olah data PLS (2026)

Berdasarkan hasil pengolahan data, konstruk *Acceptence AI* menunjukkan nilai outer loading yang baik pada seluruh indikatornya, yaitu AIA1 sebesar 0,741, AIA2 sebesar 0,782, dan AIA3 sebesar 0,807. Seluruh nilai tersebut berada di atas batas minimum yang dipersyaratkan, sehingga dapat disimpulkan bahwa indikator-indikator tersebut mampu merepresentasikan konstruk *Acceptence AI* secara memadai dan memiliki tingkat korelasi yang kuat terhadap variabel laten yang diukur.

Pada konstruk *Trust in AI*, seluruh indikator juga menunjukkan nilai outer loading yang sangat tinggi, yaitu TAI1 sebesar 0,889, TAI2 sebesar 0,795, TAI3 sebesar 0,763, dan TAI4 sebesar 0,938. Nilai tersebut menunjukkan bahwa indikator-indikator pada konstruk *Trust in AI* memiliki hubungan yang sangat kuat dengan variabel laten yang diukur. Secara statistik, konstruk ini dapat dikategorikan memiliki validitas konvergen yang sangat baik karena seluruh indikatornya berada jauh di atas batas minimum 0,70.

Sementara itu, pada konstruk Work Experience, sebagian besar indikator menunjukkan nilai yang memenuhi kriteria, yaitu WE1 sebesar 0,874, WE2 sebesar 0,831, dan WE3 sebesar 0,732. Namun, indikator WE4 memiliki nilai outer loading sebesar 0,595 yang berada di bawah batas ideal 0,70 dan sedikit di bawah batas toleransi 0,60. Meskipun demikian, indikator tersebut masih dapat dipertimbangkan untuk dipertahankan apabila hasil pengujian *Average Variance Extracted* (AVE) dan *Composite Reliability* tetap memenuhi standar yang ditetapkan. Secara keseluruhan, hasil pengujian validitas konvergen menunjukkan bahwa model pengukuran dalam penelitian ini secara umum telah memenuhi kriteria yang dipersyaratkan dan layak untuk dianalisis lebih lanjut pada tahap berikutnya.

2. Discriminant Validity

Pengujian validitas diskriminan dalam penelitian ini bertujuan untuk memastikan bahwa setiap konstruk laten yang digunakan benar-benar berbeda secara empiris dari konstruk lainnya dalam model penelitian. Validitas diskriminan menunjukkan sejauh mana suatu konstruk mampu menjelaskan variabel yang diukurnya sendiri dan tidak memiliki korelasi yang lebih tinggi dengan konstruk lain. Dengan demikian, pengujian ini penting untuk menjamin bahwa tidak terjadi tumpang tindih pengukuran antar variabel dalam model.

Dalam pendekatan Partial Least Squares (PLS), validitas diskriminan dapat diuji melalui beberapa metode, di antaranya Heterotrait-Monotrait Ratio (HTMT),

Fornell-Larcker Criterion, dan Cross Loading. Pada penelitian ini, pengujian dilakukan dengan menggunakan metode HTMT sebagai pendekatan utama karena dinilai lebih sensitif dan mampu mendeteksi permasalahan validitas diskriminan secara lebih akurat. Selain itu, metode Fornell-Larcker juga digunakan sebagai pendukung untuk memperkuat hasil analisis.

Kriteria yang digunakan dalam pengujian HTMT adalah nilai korelasi antar konstruk harus berada di bawah 0,90. Apabila seluruh nilai HTMT berada di bawah batas tersebut, maka dapat disimpulkan bahwa validitas diskriminan telah terpenuhi. Sementara itu, berdasarkan kriteria Fornell-Larcker, nilai akar kuadrat Average Variance Extracted (AVE) pada masing-masing konstruk harus lebih besar dibandingkan dengan korelasi antar konstruk lainnya dalam model. Apabila kedua kriteria tersebut terpenuhi, maka konstruk dalam penelitian ini dinyatakan memiliki validitas diskriminan yang baik.

1. Fornell-Larcker

Pengujian validitas diskriminan dalam penelitian ini dilakukan menggunakan kriteria Fornell-Larcker untuk memastikan bahwa setiap konstruk laten memiliki perbedaan yang jelas dengan konstruk lainnya dalam model. Metode ini dilakukan dengan membandingkan nilai akar kuadrat Average Variance Extracted (AVE) pada masing-masing konstruk dengan nilai korelasi antar konstruk. Suatu konstruk

dinyatakan memenuhi validitas diskriminan apabila nilai akar kuadrat AVE lebih besar dibandingkan dengan korelasi konstruk tersebut terhadap konstruk lainnya.

Tabel 4. 2 1 Fornell-Larcker

	Acceptence AI	Trust in AI	Work Experience
Acceptence AI	0.777		
Trust in AI	0.396	0.849	
Work Experience	0.511	0.602	0.766

Berdasarkan hasil pengolahan data, diperoleh nilai akar kuadrat AVE untuk konstruk *Acceptence AI* sebesar 0,777, *Trust in AI* sebesar 0,849, dan *Work Experience* sebesar 0,766. Nilai-nilai tersebut terletak pada diagonal tabel Fornell-Larcker dan menjadi indikator utama dalam pengujian validitas diskriminan. Secara umum, nilai diagonal ini mencerminkan tingkat validitas konvergen masing-masing konstruk dalam menjelaskan variabelnya sendiri.

Jika dibandingkan dengan nilai korelasi antar konstruk, terlihat bahwa seluruh nilai diagonal lebih besar dibandingkan dengan korelasi antar variabel. Konstruk *Acceptence AI* memiliki korelasi sebesar 0,396 dengan *Trust in AI* dan 0,511 dengan *Work Experience*, yang seluruhnya lebih kecil dari 0,777. Konstruk *Trust in AI* memiliki korelasi sebesar 0,602 dengan *Work Experience*, yang juga lebih kecil dibandingkan nilai diagonalnya sebesar 0,849. Demikian pula, konstruk *Work Experience* menunjukkan nilai diagonal sebesar 0,766 yang lebih tinggi dibandingkan korelasinya dengan konstruk lainnya.

Berdasarkan hasil tersebut, dapat disimpulkan bahwa seluruh konstruk dalam model penelitian ini telah memenuhi kriteria Fornell-Larcker. Hal ini menunjukkan bahwa masing-masing variabel memiliki kemampuan yang baik dalam membedakan dirinya dari konstruk lain dalam model. Dengan demikian, validitas diskriminan dalam penelitian ini dapat dinyatakan terpenuhi, sehingga model pengukuran layak untuk dilanjutkan pada tahap evaluasi model struktural.

2. Cross Loading

Pengujian validitas diskriminan juga dilakukan melalui analisis cross loading untuk memastikan bahwa setiap indikator memiliki korelasi yang lebih tinggi terhadap konstruk yang diukurnya dibandingkan dengan konstruk lainnya. Metode ini bertujuan untuk melihat apakah suatu indikator benar-benar merepresentasikan variabel laten yang dimaksud dan tidak memiliki hubungan yang lebih kuat dengan variabel lain dalam model. Suatu indikator dinyatakan memenuhi validitas diskriminan apabila nilai loading terhadap konstruk asalnya lebih besar dibandingkan nilai loading terhadap konstruk lainnya.

Tabel 4. 2. 2 analisis cross loading

	Acceptance AI	Trust in AI	Work Experience
AIA1	0.741	0.300	0.179
AIA2	0.782	0.292	0.500
AIA3	0.807	0.337	0.403
TAI1	0.266	0.889	0.453

TAI2	0.220	0.795	0.422
TAI3	0.188	0.763	0.349
TAI4	0.527	0.938	0.696
WE1	0.354	0.621	0.874
WE2	0.411	0.510	0.831
WE3	0.433	0.224	0.732
WE4	0.390	0.403	0.595

Berdasarkan hasil pengolahan data, seluruh indikator pada konstruk *Acceptance AI* menunjukkan nilai loading tertinggi pada konstraknya sendiri dibandingkan dengan konstruk lain. Indikator AIA1 memiliki loading sebesar 0,741 yang lebih tinggi dibandingkan korelasinya dengan *Trust in AI* (0,300) dan *Work Experience* (0,179). Hal yang sama juga terlihat pada AIA2 dengan nilai 0,782 serta AIA3 dengan nilai 0,807, yang keduanya tetap lebih tinggi dibandingkan cross loading pada konstruk lainnya. Hal ini menunjukkan bahwa indikator-indikator tersebut mampu merepresentasikan konstruk *Acceptance AI* secara baik.

Pada konstruk *Trust in AI*, seluruh indikator juga menunjukkan pola yang konsisten. TAI1 memiliki loading sebesar 0,889, TAI2 sebesar 0,795, TAI3 sebesar 0,763, dan TAI4 sebesar 0,938, yang semuanya lebih tinggi dibandingkan nilai loading terhadap konstruk *Acceptance AI* maupun *Work Experience*. Meskipun terdapat beberapa nilai cross loading yang cukup moderat, seperti TAI4 terhadap *Work Experience* sebesar 0,696, nilai tersebut tetap lebih rendah dibandingkan loading

utamanya pada konstruk *Trust in AI*. Dengan demikian, indikator pada konstruk ini dinyatakan memenuhi kriteria validitas diskriminan.

Selanjutnya, pada konstruk *Work Experience*, indikator WE1, WE2, WE3, dan WE4 juga menunjukkan nilai loading tertinggi pada konstruknya sendiri, masing-masing sebesar 0,874, 0,831, 0,732, dan 0,595. Walaupun indikator WE4 memiliki nilai loading yang relatif lebih rendah dibandingkan indikator lainnya, nilai tersebut tetap lebih tinggi dibandingkan cross loading terhadap konstruk lain. Hal ini menunjukkan bahwa indikator-indikator *Work Experience* masih mampu membedakan konstruknya dari variabel lainnya dalam model.

Berdasarkan keseluruhan hasil analisis cross loading, dapat disimpulkan bahwa seluruh indikator dalam penelitian ini memiliki nilai loading tertinggi pada konstruk yang diukur dibandingkan dengan konstruk lainnya. Dengan demikian, model pengukuran telah memenuhi kriteria validitas diskriminan berdasarkan metode cross loading. Hasil ini semakin memperkuat temuan sebelumnya pada uji Fornell-Larcker bahwa masing-masing konstruk dalam penelitian memiliki perbedaan yang jelas dan tidak terjadi permasalahan tumpang tindih antar variabel.

3. Heterotrait-monotrait ratio (HTMT)

Untuk memperkuat pengujian validitas diskriminan, penelitian ini menggunakan metode Heterotrait-Monotrait Ratio (HTMT) sebagai salah satu kriteria evaluasi utama. Metode HTMT merupakan pendekatan yang lebih sensitif dalam

mendeteksi permasalahan validitas diskriminan dibandingkan metode tradisional seperti Fornell-Larcker. Suatu konstruk dinyatakan memenuhi validitas diskriminan apabila nilai HTMT antar konstruk berada di bawah batas yang direkomendasikan, yaitu kurang dari 0,90 atau 0,85 untuk kriteria yang lebih ketat.

Tabel 4. 2 3 Heterotrait-Monotrait Ratio (HTMT)

	Heterotrait-monotrait ratio (HTMT)
Trust in AI <-> Acceptence AI	0.463
Work Experience <-> Acceptence AI	0.653
Work Experience <-> Trust in AI	0.667

Berdasarkan hasil pengolahan data, diperoleh nilai HTMT antara Trust in AI dan Acceptence AI sebesar 0,463, antara Work Experience dan Acceptence AI sebesar 0,653, serta antara Work Experience dan Trust in AI sebesar 0,667. Seluruh nilai tersebut berada di bawah batas maksimum 0,90, sehingga tidak menunjukkan adanya permasalahan diskriminasi antar konstruk dalam model penelitian ini.

Dengan demikian, berdasarkan hasil uji HTMT, seluruh konstruk dalam penelitian ini dapat dinyatakan telah memenuhi kriteria validitas diskriminan. Hasil ini memperkuat temuan sebelumnya melalui uji Fornell-Larcker dan Cross Loading bahwa masing-masing variabel memiliki karakteristik yang berbeda secara empiris. Oleh karena itu, model pengukuran dalam penelitian ini layak untuk dilanjutkan pada tahap evaluasi model struktural.

4. Konstruk Reabilitas dan Validitas

Pengujian reliabilitas konstruk dalam penelitian ini dilakukan untuk memastikan bahwa setiap variabel laten memiliki konsistensi internal yang memadai dalam mengukur konsep yang diteliti. Evaluasi reliabilitas dilakukan dengan melihat nilai Cronbach's Alpha dan Composite Reliability. Suatu konstruk dinyatakan reliabel apabila memiliki nilai Cronbach's Alpha dan Composite Reliability di atas 0,70. Selain itu, nilai Average Variance Extracted (AVE) juga diperhatikan untuk memastikan bahwa konstruk memiliki validitas konvergen yang baik, dengan kriteria nilai AVE di atas 0,50.

Tabel 4. 2. 4 Konstruk Reabilitas dan Validitas

	Cronbach's alpha	Composite reliability (rho_a)	Composite reliability (rho_c)	Average variance extracted (AVE)
Acceptance AI	0.692	0.702	0.820	0.604
Trust in AI	0.872	0.993	0.911	0.721
Work Experience	0.756	0.780	0.847	0.586

Berdasarkan hasil pengolahan data, konstruk Acceptance AI memiliki nilai Cronbach's Alpha sebesar 0,692, Composite Reliability (rho_a) sebesar 0,702, Composite Reliability (rho_c) sebesar 0,820, serta AVE sebesar 0,604. Meskipun nilai Cronbach's Alpha sedikit berada di bawah batas ideal 0,70, nilai tersebut masih dapat diterima dalam penelitian eksploratif. Selain itu, nilai Composite Reliability dan AVE

telah memenuhi kriteria yang dipersyaratkan, sehingga konstruk Acceptance AI tetap dapat dinyatakan reliabel dan valid.

Trust in AI menunjukkan hasil yang sangat baik dengan nilai Cronbach's Alpha sebesar 0,872, Composite Reliability (ρ_a) sebesar 0,993, Composite Reliability (ρ_c) sebesar 0,911, serta AVE sebesar 0,721. Seluruh nilai tersebut berada di atas batas minimum yang ditetapkan, sehingga dapat disimpulkan bahwa konstruk Trust in AI memiliki tingkat reliabilitas dan validitas konvergen yang sangat kuat dalam model penelitian ini.

Selanjutnya, konstruk Work Experience memiliki nilai Cronbach's Alpha sebesar 0,756, Composite Reliability (ρ_a) sebesar 0,780, Composite Reliability (ρ_c) sebesar 0,847, serta AVE sebesar 0,586. Seluruh nilai tersebut telah memenuhi kriteria yang dipersyaratkan, sehingga konstruk Work Experience dinyatakan reliabel dan valid. Secara keseluruhan, hasil pengujian reliabilitas dan validitas menunjukkan bahwa seluruh konstruk dalam penelitian ini memiliki konsistensi internal yang baik dan layak untuk dilanjutkan pada tahap evaluasi model struktural.

5. Evaluasi Model Struktural (Inner Model)

Evaluasi model struktural (inner model) dilakukan untuk menganalisis hubungan antar konstruk laten dalam penelitian serta menguji hipotesis yang telah dirumuskan sebelumnya. Berbeda dengan model pengukuran yang berfokus pada validitas dan reliabilitas indikator, model struktural bertujuan untuk menilai kekuatan

hubungan antar variabel serta kemampuan model dalam menjelaskan variabel endogen. Tahapan ini menjadi bagian penting dalam analisis Partial Least Squares (PLS) karena menentukan apakah model penelitian memiliki daya prediksi yang memadai.

Pengujian model struktural dalam penelitian ini dilakukan dengan menganalisis beberapa indikator evaluasi, yaitu nilai koefisien determinasi (R-Square), nilai path coefficient, nilai t-statistic dan p-value hasil bootstrapping, serta effect size (f^2) apabila diperlukan. Melalui pengujian ini, dapat diketahui seberapa besar pengaruh variabel independen terhadap variabel dependen, serta apakah pengaruh tersebut signifikan secara statistik

Nilai R-Square digunakan untuk mengukur seberapa besar variasi variabel endogen mampu dijelaskan oleh variabel eksogen dalam model. Semakin tinggi nilai R-Square, maka semakin besar kemampuan model dalam menjelaskan fenomena yang diteliti. Dalam penelitian berbasis PLS, nilai R-Square sebesar 0,75 dikategorikan kuat, 0,50 dikategorikan moderat, dan 0,25 dikategorikan lemah.

Selanjutnya, pengujian hipotesis dilakukan dengan melihat nilai path coefficient yang menunjukkan arah dan kekuatan hubungan antar konstruk. Signifikansi hubungan tersebut ditentukan berdasarkan nilai t-statistic yang harus lebih besar dari 1,96 pada tingkat signifikansi 5 persen, serta p-value yang harus lebih kecil dari 0,05. Hasil ini akan menentukan apakah hipotesis penelitian diterima atau ditolak.

Melalui evaluasi model struktural ini, penelitian tidak hanya menguji keberlakuan instrumen pengukuran, tetapi juga menguji model konseptual secara keseluruhan. Apabila hasil pengujian menunjukkan hubungan yang signifikan dan model memiliki daya jelaskan yang memadai, maka model penelitian dapat dinyatakan memiliki kekuatan empiris dalam menjelaskan hubungan antar variabel yang diteliti.

6. Variance Inflation Factor

Sebelum melakukan pengujian koefisien determinasi dan uji hipotesis, terlebih dahulu dilakukan pengujian multikolinearitas pada model struktural. Pengujian ini bertujuan untuk memastikan bahwa tidak terdapat korelasi yang tinggi antar variabel prediktor dalam model yang dapat mengganggu estimasi parameter. Dalam pendekatan Partial Least Squares (PLS), multikolinearitas diuji menggunakan nilai Variance Inflation Factor (VIF).

Tabel 4. 2 .5
Variance Inflation Factors (VIF)

	VIF
AIA1	1.750
AIA2	1.146
AIA3	1.722
TAI1	2.905
TAI2	2.131
TAI3	2.040
TAI4	3.039

WE1	2.015
WE2	2.083
WE3	1.885
WE4	1.263

Suatu model dinyatakan bebas dari permasalahan multikolinearitas apabila nilai VIF berada di bawah batas yang direkomendasikan. Secara umum, nilai VIF di bawah 5,00 menunjukkan tidak adanya masalah multikolinearitas yang serius, sedangkan beberapa literatur merekomendasikan batas yang lebih ketat yaitu di bawah 3,00. Oleh karena itu, semakin kecil nilai VIF maka semakin baik kualitas model dalam menjelaskan hubungan antar variabel.

Berdasarkan hasil pengolahan data, seluruh indikator dalam penelitian ini memiliki nilai VIF yang berada di bawah 5,00. Nilai VIF tertinggi terdapat pada indikator TAI4 sebesar 3,039, sementara nilai terendah terdapat pada indikator AIA2 sebesar 1,146. Indikator lainnya juga menunjukkan nilai yang relatif rendah, seperti AIA1 sebesar 1,750, TAI1 sebesar 2,905, serta WE4 sebesar 1,263. Hal ini menunjukkan bahwa tidak terdapat korelasi yang berlebihan antar indikator dalam model.

Dengan demikian, dapat disimpulkan bahwa model penelitian ini bebas dari permasalahan multikolinearitas dan memenuhi kriteria yang dipersyaratkan dalam analisis PLS-SEM. Oleh karena itu, model struktural dapat dilanjutkan ke tahap

evaluasi berikutnya, yaitu pengujian koefisien determinasi (R-Square) dan uji signifikansi hubungan antar konstruk.

7. Nilai R-Square

Pengujian koefisien determinasi (R-Square) dilakukan untuk mengetahui seberapa besar kemampuan variabel eksogen dalam menjelaskan variabel endogen dalam model penelitian. Nilai R-Square menunjukkan proporsi variasi variabel dependen yang dapat dijelaskan oleh variabel independen yang memengaruhinya. Dalam analisis Partial Least Squares (PLS), nilai R-Square sebesar 0,75 dikategorikan kuat, 0,50 dikategorikan moderat, dan 0,25 dikategorikan lemah.

Tabel 4. 2 .6
R-Square

	R-square	R-square adjusted
Acceptance AI	0.273	0.243
Work Experience	0.363	0.350

Berdasarkan hasil pengolahan data, konstruk Acceptance AI memiliki nilai R-Square sebesar 0,273 dan R-Square adjusted sebesar 0,243. Hal ini menunjukkan bahwa sebesar 27,3% variasi pada variabel Acceptance AI dapat dijelaskan oleh variabel-variabel yang memengaruhinya dalam model penelitian, sedangkan sisanya sebesar 72,7% dipengaruhi oleh faktor lain di luar model yang tidak diteliti dalam penelitian ini. Berdasarkan kriteria interpretasi, nilai ini termasuk dalam kategori lemah hingga moderat.

Selanjutnya, konstruk Work Experience memiliki nilai R-Square sebesar 0,363 dan R-Square adjusted sebesar 0,350. Artinya, sebesar 36,3% variasi pada variabel Work Experience dapat dijelaskan oleh variabel independen dalam model, sedangkan 63,7% lainnya dipengaruhi oleh faktor di luar penelitian. Nilai tersebut menunjukkan kemampuan penjelasan model yang berada pada kategori lemah menuju moderat.

Secara keseluruhan, hasil pengujian R-Square menunjukkan bahwa model penelitian memiliki kemampuan penjelasan yang cukup dalam menerangkan variabel endogen, meskipun belum tergolong kuat. Hal ini mengindikasikan bahwa masih terdapat faktor-faktor lain di luar model yang berpotensi memengaruhi variabel yang diteliti. Meskipun demikian, nilai tersebut masih dapat diterima dalam penelitian sosial dan perilaku, sehingga analisis dapat dilanjutkan pada tahap pengujian hipotesis melalui analisis koefisien jalur (path coefficient).

8. Uji Hipotesis

Pengujian hipotesis dalam penelitian ini dilakukan dengan menggunakan metode bootstrapping pada analisis Partial Least Squares (PLS). Pengujian ini bertujuan untuk mengetahui signifikansi hubungan antar konstruk yang telah dirumuskan dalam hipotesis penelitian.

Keputusan penerimaan atau penolakan hipotesis didasarkan pada nilai t-statistic yang harus lebih besar dari 1,96 pada tingkat signifikansi 5% serta nilai p-value yang

harus lebih kecil dari 0,05. Selain itu, nilai original sample (O) digunakan untuk melihat arah dan kekuatan hubungan antar variabel.

Tabel 4. 2 .7 Pengujian Hipotesis

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
Trust in AI -> Acceptance AI	0.396	0.415	0.123	3.228	0.001
Trust in AI -> Work Experience	0.602	0.620	0.080	7.506	0.000
Work Experience -> Acceptance AI	0.428	0.488	0.165	2.587	0.010

Hasil pengujian menunjukkan bahwa hubungan antara Trust in AI terhadap Acceptance AI memiliki nilai koefisien jalur sebesar 0,396 dengan nilai t-statistic sebesar 3,228 dan p-value sebesar 0,001. Karena nilai t-statistic lebih besar dari 1,96 dan p-value lebih kecil dari 0,05, maka hipotesis yang menyatakan bahwa Trust in AI berpengaruh terhadap Acceptance AI dinyatakan diterima. Koefisien positif menunjukkan bahwa semakin tinggi tingkat kepercayaan terhadap AI, maka semakin tinggi pula tingkat penerimaan terhadap AI.

Selanjutnya, hubungan antara Trust in AI terhadap Work Experience menunjukkan nilai koefisien sebesar 0,602 dengan nilai t-statistic sebesar 7,506 dan p-value sebesar 0,000. Hasil ini menunjukkan pengaruh yang positif dan sangat signifikan. Dengan demikian, hipotesis yang menyatakan bahwa Trust in AI

berpengaruh terhadap Work Experience dinyatakan diterima. Nilai koefisien yang relatif besar mengindikasikan bahwa Trust in AI memiliki pengaruh yang kuat terhadap Work Experience dalam model penelitian ini.

Terakhir, hubungan antara Work Experience terhadap Acceptance AI memiliki nilai koefisien sebesar 0,428 dengan nilai t-statistic sebesar 2,587 dan p-value sebesar 0,010. Karena nilai t-statistic lebih besar dari 1,96 dan p-value lebih kecil dari 0,05, maka hipotesis ini juga dinyatakan diterima. Hal ini menunjukkan bahwa pengalaman kerja memiliki pengaruh positif dan signifikan terhadap penerimaan AI. Secara keseluruhan, hasil uji hipotesis menunjukkan bahwa seluruh hubungan antar variabel dalam model penelitian ini signifikan dan sesuai dengan hipotesis yang telah dirumuskan.

Tabel 4.3 Keterangan Hipotesis

Hipotesis	Hubungan	Path Coeficients	Keterangan
H1	Trust in AI → AI Acceptance	0.396	Berpengaruh positif dan signifikan
H2	Trust in AI → Work Experience	0,602	Berpengaruh positif dan signifikan
H3	Work Experience → AI Acceptance	0,428	Berpengaruh positif dan signifikan

H1: Trust in AI $\rightarrow$ AI Acceptance (Koefisien jalur = 0,396)

Berdasarkan hasil pengujian hipotesis, Trust in AI berpengaruh positif dan signifikan terhadap AI Acceptance dengan nilai koefisien jalur sebesar 0,396, nilai t-statistics sebesar 3,228, dan p-value sebesar 0,001. Karena nilai t-statistics lebih besar dari 1,96 dan p-value lebih kecil dari 0,05, maka hipotesis pertama diterima. Nilai koefisien tersebut menunjukkan bahwa peningkatan kepercayaan terhadap Artificial Intelligence akan diikuti oleh peningkatan penerimaan terhadap penggunaan AI dalam deteksi penipuan keuangan.

Secara substantif, hasil ini menunjukkan bahwa kepercayaan terhadap kemampuan, keandalan, dan keamanan sistem AI menjadi faktor penting dalam membentuk sikap positif profesional keuangan terhadap adopsi teknologi. Semakin tinggi tingkat keyakinan individu terhadap akurasi dan transparansi sistem AI, maka semakin besar kecenderungan mereka untuk mengintegrasikan teknologi tersebut dalam aktivitas audit dan pengendalian risiko. Dengan demikian, trust berperan sebagai determinan utama dalam proses penerimaan teknologi berbasis kecerdasan buatan.

H2: Trust in AI $\rightarrow$ Work Experience (Koefisien jalur = 0,602)

Berdasarkan hasil pengujian hipotesis, Trust in AI berpengaruh positif dan signifikan terhadap Work Experience dengan nilai koefisien jalur sebesar 0,602, nilai t-statistics sebesar 7,506, dan p-value sebesar 0,000. Karena nilai t-statistics lebih besar dari 1,96 dan p-value lebih kecil dari 0,05, maka hipotesis kedua diterima. Nilai

koefisien ini menunjukkan pengaruh yang kuat, yang berarti peningkatan kepercayaan terhadap Artificial Intelligence secara signifikan mendorong peningkatan pengalaman kerja dalam konteks pemanfaatan teknologi.

Secara substantif, hasil ini mengindikasikan bahwa profesional yang memiliki tingkat kepercayaan tinggi terhadap AI cenderung lebih aktif dan terbuka dalam mengintegrasikan teknologi tersebut ke dalam praktik kerja mereka. Kepercayaan terhadap sistem memungkinkan individu untuk lebih percaya diri dalam menggunakan AI sebagai alat bantu analisis, sehingga pengalaman kerja yang relevan dengan teknologi digital semakin berkembang dan memperkuat kesiapan dalam menghadapi transformasi digital di sektor audit.

H3: Work Experience $\rightarrow$ AI Acceptance (Koefisien jalur = 0,428)

Hasil pengujian menunjukkan bahwa Work Experience berpengaruh positif dan signifikan terhadap AI Acceptance dengan nilai koefisien jalur sebesar 0,428, nilai t-statistics sebesar 2,587, dan p-value sebesar 0,010. Karena nilai t-statistics lebih besar dari 1,96 dan p-value lebih kecil dari 0,05, maka hipotesis ketiga diterima. Koefisien ini menunjukkan bahwa pengalaman kerja memiliki pengaruh dengan kekuatan sedang terhadap penerimaan AI dalam deteksi penipuan keuangan.

Secara konseptual, semakin tinggi pengalaman profesional individu di bidang keuangan, khususnya dalam proses audit dan pengawasan risiko, maka semakin rasional dan terbuka sikap mereka terhadap penggunaan AI. Pengalaman kerja

memberikan pemahaman yang lebih mendalam terhadap kompleksitas risiko kecurangan, sehingga individu dapat menilai AI sebagai solusi yang mampu meningkatkan efektivitas, efisiensi, dan kualitas pengambilan keputusan dalam praktik profesional.

9. Peran Mediasi Pengalaman kerja (*Specific Indirect Effect*)

Selain pengujian pengaruh langsung antar variabel, penelitian ini juga melakukan pengujian efek tidak langsung untuk mengetahui apakah terdapat mekanisme mediasi dalam model struktural yang dibangun. Pengujian ini bertujuan untuk menganalisis apakah *Trust in AI* memengaruhi *Acceptance AI* secara tidak langsung melalui *Work Experience* sebagai variabel perantara. Analisis dilakukan menggunakan metode bootstrapping pada SEM-PLS untuk menguji signifikansi pengaruh tidak langsung tersebut. Hasil pengujian specific indirect effects disajikan pada tabel berikut.

Tabel 4. 2 8 Specific Indirect Effects

	Specific indirect effects
Trust in AI -> Work Experience -> Acceptance AI	0.258

Berdasarkan hasil pengujian specific indirect effects, diperoleh nilai koefisien pengaruh tidak langsung Trust in AI terhadap Acceptance AI melalui Work Experience sebesar 0,258. Nilai koefisien ini menunjukkan bahwa Trust in AI memiliki pengaruh positif terhadap Acceptance AI melalui perantara Work Experience. Artinya, semakin tinggi tingkat kepercayaan terhadap Artificial Intelligence, maka semakin besar

kecenderungan individu untuk mengembangkan pengalaman kerja yang relevan dalam pemanfaatan teknologi, yang pada akhirnya mendorong peningkatan penerimaan terhadap AI dalam deteksi penipuan keuangan.

Hasil bootstrapping menunjukkan bahwa pengaruh tidak langsung tersebut signifikan secara statistik, dengan nilai t-statistics sebesar 2,253 dan p-value sebesar 0,024. Karena nilai t-statistics lebih besar dari 1,96 dan p-value lebih kecil dari 0,05, maka dapat disimpulkan bahwa Work Experience terbukti memediasi hubungan antara Trust in AI dan Acceptance AI. Dengan demikian, pengalaman kerja berperan sebagai variabel perantara yang memperkuat hubungan antara kepercayaan terhadap AI dan penerimaan teknologi tersebut dalam praktik profesional di Kantor Akuntan Publik.

C. Pembahasan

Berdasarkan hasil analisis yang telah dilakukan, penelitian ini menunjukkan bahwa seluruh hipotesis yang diajukan diterima secara statistik. Hal ini mengindikasikan bahwa kepercayaan terhadap Artificial Intelligence (AI) dan pengalaman kerja memiliki peran penting dalam membentuk penerimaan AI sebagai alat bantu dalam mendeteksi penipuan keuangan. Temuan ini memperlihatkan bahwa faktor psikologis dan profesional secara simultan memengaruhi kesiapan individu dalam mengadopsi teknologi berbasis kecerdasan buatan di lingkungan kerja.

Pengaruh kepercayaan terhadap AI terhadap penerimaan AI terbukti positif dan signifikan. Artinya, semakin tinggi tingkat kepercayaan auditor atau profesional

keuangan terhadap kemampuan, keandalan, dan keamanan sistem AI, maka semakin tinggi pula tingkat penerimaan mereka terhadap penggunaan teknologi tersebut. Temuan ini sejalan dengan teori penerimaan teknologi yang menekankan bahwa trust merupakan faktor krusial dalam proses adopsi sistem berbasis teknologi. Dalam konteks deteksi penipuan keuangan, kepercayaan menjadi fondasi penting karena AI digunakan untuk menganalisis data sensitif dan mendukung pengambilan keputusan profesional.

Selain itu, pengalaman kerja juga berpengaruh positif dan signifikan terhadap penerimaan AI. Hal ini menunjukkan bahwa semakin tinggi pengalaman profesional seseorang di bidang keuangan, terutama dalam pengawasan dan deteksi kecurangan, maka semakin besar pula kecenderungan untuk menerima penggunaan AI. Pengalaman kerja memberikan dasar pemahaman mengenai kompleksitas risiko penipuan, sehingga profesional yang berpengalaman dapat melihat AI sebagai solusi yang membantu meningkatkan efektivitas dan efisiensi proses deteksi.

Meskipun nilai R-Square menunjukkan bahwa kemampuan model dalam menjelaskan variabel endogen berada pada kategori lemah hingga moderat, hasil ini masih dapat diterima dalam penelitian sosial dan perilaku. Hal tersebut mengindikasikan bahwa selain kepercayaan dan pengalaman kerja, masih terdapat faktor lain yang berpotensi memengaruhi penerimaan AI, seperti budaya organisasi, dukungan manajemen, kemudahan penggunaan sistem, atau faktor regulasi. Dengan

demikian, penelitian ini membuka peluang bagi studi lanjutan untuk mengembangkan model yang lebih komprehensif.

Secara keseluruhan, penelitian ini memberikan implikasi teoritis dan praktis yang relevan dalam konteks penerapan *Artificial Intelligence* (AI) pada sektor audit dan pengawasan keuangan. Secara teoritis, hasil penelitian ini memperkuat konsep bahwa kepercayaan (trust) dan pengalaman profesional merupakan determinan penting dalam membentuk penerimaan teknologi berbasis AI, khususnya dalam lingkungan kerja yang menuntut akurasi dan integritas tinggi seperti deteksi penipuan keuangan. Secara praktis, temuan ini menunjukkan bahwa peningkatan penerimaan AI tidak semata-mata bergantung pada kecanggihan sistem yang digunakan, tetapi juga pada upaya membangun kepercayaan pengguna serta meningkatkan kompetensi dan kesiapan sumber daya manusia. Bagi Kantor Akuntan Publik (KAP), hal ini berarti perlunya membangun ekosistem kerja yang mendukung integrasi AI secara menyeluruh, melalui penyediaan pelatihan penggunaan AI, peningkatan literasi digital auditor, transparansi sistem dalam menjelaskan proses analisis, serta jaminan keamanan data. Dengan memperkuat aspek kepercayaan dan kompetensi profesional, implementasi AI dapat berjalan lebih optimal dan berkontribusi pada peningkatan efektivitas deteksi dini kecurangan serta kualitas pengambilan keputusan dalam praktik audit.