

BAB II

KERANGKA PENELITIAN

2.1 Tinjauan Teori

Penelitian ini berada dalam ranah Manajemen Teknologi dan keamanan Siber Organisasi, yang merupakan interseksi antara manajemen sumber daya manusia, tata kelola teknologi dan perilaku pengguna dalam lingkungan kerja digital. Fenomena *Shadow AI* dimana penggunaan alat kecerdasan buatan generatif (*GenAI*) secara informal oleh karyawan tanpa otorisasi departemen TI dapat menjadi ancaman sistemik terhadap keamanan data, kepatuhan regulasi dan tata kelola teknologi.

2.1.1 *Shadow AI*

Ekspansi yang besar dari teknologi *artificial intelligence* (AI) di berbagai sektor telah menyebabkan pergeseran besar dalam alur kerja, pengambilan keputusan dan inovasi organisasi (McKinsey & Company, 2023; World Economic Forum, 2024). Namun, penetrasi yang luas ini juga mengakibatkan berkembangnya fenomena “*Shadow AI*”, yaitu penerapan dan penggunaan alat, model atau sistem AI oleh karyawan suatu organisasi tanpa persetujuan yang eksplisit, pengawasan atau manajemen dari departemen TI dan keamanan siber (IBM, 2024b; Puthal et al., 2025). Fenomena ini semakin berkembang seiring dengan kemudahan akses terhadap berbagai platform *generative AI* yang tersedia secara luas di internet, sehingga individu dapat memanfaatkan teknologi tersebut secara mandiri dalam aktivitas pekerjaan mereka. Penelitian menunjukkan bahwa karyawan sering menggunakan alat *generative AI* untuk menyusun laporan, menghasilkan konten, atau mendukung analisis pekerjaan tanpa adanya kerangka kebijakan organisasi yang jelas, sebuah praktik yang dikenal sebagai *shadow use of AI* atau penggunaan AI secara informal dalam organisasi (Combetto, 2026). Implementasi sistem AI yang tidak teregulasi ini sering kali terjadi sebagai akibat dari karyawan atau tim yang

memanfaatkan solusi AI yang mudah diperoleh untuk menangani masalah operasional, mengotomasi suatu proses atau meningkatkan produktivitas secara cepat tanpa melalui prosedur organisasi formal. Walaupun upaya ini dilakukan dengan niatan baik namun upaya tersebut berada di luar *governing framework* yang normal, sehingga menciptakan celah dan kompleksitas dalam struktur organisasi.

Fenomena *Shadow AI* menjangkau berbagai domain lintas fungsi organisasi seperti *marketer* yang menggunakan *tool analytics* berbasis AI dan *recruiter* yang menggunakan *recruitment tools* berbasis AI, sejalan dengan temuan McKinsey & Company (2023) bahwa fungsi *marketing*, *product development* dan *service operations* merupakan area dengan adopsi GenAI tertinggi. Namun, penggunaan yang disebut *shadow* atau informal sering kali dilakukan tanpa adanya ukuran keamanan, akuntabilitas dan transparansi yang memadai. Sehingga membuat organisasi terpapar risiko *data breach*, *regulatory non-compliance* serta ancaman serangan siber (The European Union Agency for Cybersecurity (ENISA), 2025) (Puthal et al., 2025). Dengan demikian, *Shadow AI* bukan sekedar isu adopsi teknologi, tetapi juga persoalan tata kelola, manajemen risiko dan keamanan organisasi yang semakin kompleks di era *generative AI*.

2.1.2 Grand Theory – Expectancy Value Theory (EVT)

Jika menilik kebelakang, sebagai *grand theory* atau akar filosofis dari *Protection Motivation Theory* (PMT) terletak pada sebuah teori bernama *Expectancy Value Theory* (EVT) yang dikemukakan oleh Vroom, (1964). *Expectancy Value Theory* (EVT) menyatakan bahwa motivasi adalah hasil dari perkalian antara keyakinan bahwa tindakan akan menghasilkan outcome (*expectancy*) dan nilai yang diberikan pada outcome tersebut (*value*). R. Rogers (1975) secara eksplisit mengadopsi paradigma ini, dimana *coping appraisal* (*self efficacy* dan *response efficacy*) merepresentasikan *expectancy*, sedangkan *threat appraisal* (*perceived severity* dan *perceived vulnerability*) merepresentasikan *value* negatif yang

ingin dihindari. Studi oleh Shang et al (2023) dan Lyu & Salam (2025) memperkuat hubungan ini, menunjukkan bahwa *Expectancy Value Theory* (EVT) memberikan fondasi teoretis yang kuat untuk memahami penggunaan teknologi, dalam hal ini termasuk AI. Dalam konteks *Shadow AI*, karyawan menghadapi dilema antara *value* positif seperti efisiensi dan produktivitas, dan *value* negatif seperti risiko kebocoran data, dengan *expectancy* ditentukan oleh keyakinan mereka terhadap kemampuan diri dan efektivitas solusi resmi dari perusahaan.

2.1.3 Middle Range Theory - *Protection Motivation Theory* (PMT)

Protection Motivation Theory (PMT) yang dikembangkan oleh R. Rogers (1975) menjadi fondasi utama penelitian ini. *Protection Motivation Theory* (PMT) menjelaskan bagaimana individu termotivasi untuk melindungi diri dari ancaman melalui dua proses kognitif yaitu *threat appraisal* dan *coping appraisal*. Dalam konteks keamanan siber, *Protection Motivation Theory* (PMT) telah terbukti menjadi kerangka teoritis paling kuat untuk menjelaskan perilaku protektif maupun *maladaptive*. Berdasarkan meta analisis yang dilakukan oleh Sommestad et al (2015) menunjukkan bahwa model *Protection Motivation Theory* (PMT) mampu menjelaskan 34% hingga 50% varian niat keamanan siber, dengan komponen *coping appraisal* utamanya *self-efficacy* dan *response efficacy* lebih dominan dari pada *threat appraisal*. Sementara pada penelitian-penelitian terbaru seperti yang dilakukan oleh Kiran et al (2025) dan Alrawhani et al (2025) memperkuat validitas dari *Protection Motivation Theory* dalam konteks kepatuhan kebijakan keamanan informasi, meskipun pada penelitian keduanya menemukan bahwa *perceived vulnerability* sering tidak signifikan untuk sebuah temuan yang relevan dengan konteks *Shadow AI* dimana karyawan cenderung merasa “tidak rentan” meskipun mereka menyadari risikonya.

Komponen *Protection Motivation Theory* jika dibedah secara mendetail terdiri dari dua proses appraisal utama yang membentuk motivasi

protektif individu. Pertama, *Threat Appraisal* (penilaian ancaman) yang mencakup evaluasi individu terhadap tingkat keparahan ancaman (*perceived severity*) dan kerentanan terhadap ancaman tersebut (*perceived vulnerability*). Menurut Floyd et al (2000), *perceived severity* mengacu pada sejauh mana individu mempersepsikan konsekuensi negatif dari suatu ancaman sebagai hal serius atau berbahaya, sementara *perceived vulnerability* merepresentasikan keyakinan individu tentang kemungkinan terjadinya ancaman tersebut.

Kedua, *Coping Appraisal* (penilaian koping) yang terdiri dari tiga komponen utama: (1) *Self-Efficacy* yaitu keyakinan individu terhadap kemampuannya untuk melaksanakan tindakan protektif; (2) *Response Efficacy* yaitu keyakinan bahwa tindakan yang direkomendasikan efektif dalam mengurangi ancaman; dan (3) *Response Cost* yaitu persepsi terhadap biaya atau hambatan yang harus dikeluarkan untuk melaksanakan tindakan protektif (Floyd et al., 2000; R. Rogers, 1983).

Penelitian (Li et al., 2022) secara eksplisit menguji hubungan ini dan menemukan bahwa kesadaran keamanan siber memiliki pengaruh positif yang signifikan terhadap *perceived severity* dan *perceived vulnerability*. Temuan ini menunjukkan bahwa karyawan yang lebih sadar akan risiko siber cenderung mempresentasikan ancaman sebagai lebih serius dan lebih mungkin terjadi pada diri mereka. Lebih lanjut, studi yang sama juga mengungkap bahwa *cybersecurity awareness* berdampak positif pada *response efficacy* dan *self-efficacy*, sekaligus berdampak negatif pada *response cost*.

Dalam konteks perilaku sukarela (*voluntary behaviour*), meta-analisis Sommestad et al (2015) menegaskan bahwa variabel coping appraisal (*self-efficacy* dan *response efficacy*) lebih penting untuk perilaku sukarela, sedangkan *threat appraisal* lebih penting untuk perilaku wajib. Temuan ini diperkuat oleh Kiran et al (2025) yang melaporkan bahwa dalam

lingkungan kerja digital, komponen *coping appraisal* secara konsisten lebih prediktif dibandingkan *threat appraisal*.

2.1.4 Applied Range Theory - Cybersecurity Awareness

Cybersecurity awareness jika dikaitkan pada konteks organisasi atau lingkup kerja mengacu pada tingkat pemahaman karyawan mengenai risiko keamanan siber, kebijakan organisasi terkait perlindungan data serta konsekuensi pribadi dan organisasi dari pelanggaran keamanan informasi (Sulaiman et al., 2022; Tran et al., 2025). Dalam konteks penelitian ini, *cybersecurity awareness* bukan sekedar pengetahuan teoretis, melainkan kemampuan kognitif untuk mengenali ancaman, memahami kerentanan sistem dan menyadari dampak dari tindakan berisiko seperti penggunaan alat AI generatif secara informal (Li et al., 2022). Perspektif ini sejalan dengan literatur keamanan siber yang menekankan bahwa kesadaran pengguna terhadap risiko teknologi merupakan faktor penting dalam mencegah pelanggaran keamanan informasi dan penggunaan sistem digital secara tidak aman dalam organisasi (Rjoub et al., 2023). Penelitian mengenai keamanan AI juga menunjukkan bahwa kompleksitas sistem AI modern sering kali menimbulkan tantangan transparansi dan akuntabilitas, sehingga meningkatkan pentingnya pemahaman pengguna terhadap risiko teknologi tersebut dalam lingkungan kerja (Rjoub et al., 2023).

Menurut Tran et al (2025) definisi operasional ini selaras bahwa kesadaran keamanan siber mencakup pemahaman tentang nilai asset informasi, kewajiban individu, dan mekanisme pengendalian yang memadai untuk melindungi data perusahaan. Dalam kerangka *Protection Motivation Theory* (PMT), *cybersecurity awareness* berperan sebagai antesedan kritis yang membentuk dua proses kognitif utama yaitu *threat appraisal* dan *coping appraisal*. Penelitian empiris oleh Li et al (2022) secara eksplisit menguji hubungan ini dan menemukan bahwa kesadaran keamanan siber memiliki pengaruh positif yang signifikan terhadap *perceived severity* dan *perceived vulnerability*. Temuan ini menunjukkan bahwa karyawan yang

lebih sadar akan risiko siber cenderung mempresentasikan ancaman sebagai lebih serius dan lebih mungkin terjadi pada diri mereka. Lebih lanjut, studi yang sama juga mengungkap bahwa *cybersecurity awareness* berdampak positif pada *response efficacy* dan *self efficacy*, sekaligus berdampak negatif pada *response cost*. Hal ini berarti bahwa pengetahuan yang memadai tidak hanya meningkatkan persepsi ancaman, tetapi juga memperkuat keyakinan karyawan bahwa mereka mampu dan memiliki solusi yang efektif untuk menghadapi ancaman tersebut, sekaligus mengurangi persepsi bahwa tindakan protektif itu merepotkan.

Relevansi *cybersecurity awareness* dalam konteks *Shadow AI* menjadi semakin krusial. Laporan dari Cyberhaven labs, (2024) menunjukkan bahwa 77% karyawan global bingung tentang batasan penggunaan AI di tempat kerja yang mengindikasikan rendahnya tingkat kesadaran keamanan siber terkait teknologi baru ini. Ketidaktahuan ini menciptakan celah dimana karyawan, meskipun memiliki niat baik untuk meningkatkan produktivitas, secara tidak sadar membahayakan asset data perusahaan dengan mengunggah informasi sensitif ke *platform* AI publik seperti ChatGPT atau Gemini. Dalam situasi ini, *cybersecurity awareness* berfungsi sebagai “peringatan dini” yang membantu karyawan membedakan antara penggunaan AI yang aman (melalui saluran resmi perusahaan) dan penggunaan yang berisiko tinggi (melalui akun pribadi).

Dengan demikian, *cybersecurity awareness* dalam penelitian ini dipahami sebagai fundamen kognitif yang mendasari seluruh proses pengambilan keputusan terkait keamanan siber. Tanpa tingkat kesadaran yang memadai, karyawan tidak akan memiliki dasar yang kuat untuk menilai ancaman *Shadow AI* atau untuk mempercayai bahwa solusi resmi yang disediakan perusahaan adalah pilihan yang layak dan efektif. Hal ini juga sejalan dengan literatur mengenai adopsi AI di organisasi yang menekankan bahwa pemahaman karyawan terhadap risiko teknologi merupakan faktor penting dalam memastikan penggunaan AI yang aman

dan bertanggung jawab di tempat kerja (Bankins et al., 2023; Budhwar et al., 2023). Oleh karena itu, variabel ini menjadi titik awal yang kritis dalam model konseptual untuk menjelaskan mengapa beberapa karyawan memilih untuk mematuhi kebijakan keamanan, sementara yang lain justru terdorong untuk menggunakan *Shadow AI*.

2.1.5 Applied Range Theory - AI Literacy

AI Literacy mengacu pada kemampuan individu untuk memahami, mengevaluasi dan menggunakan teknologi kecerdasan buatan secara kritis, aman dan etis (Gillespie et al., 2025; Lyu & Salam, 2025). Dalam konteks penelitian ini, *AI Literacy* tidak hanya mencakup pengetahuan teknis tentang cara kerja AI, tetapi juga kemampuan untuk menilai risiko privasi, mengenali bias algoritmik dan memverifikasi keakuratan output, hal ini merupakan ketrampilan yang krusial dalam mencegah praktik *Shadow AI* yang berisiko tinggi. Berdasarkan Gillespie et al (2025) pada Laporan KPMG & University of Melbourne 2025 menunjukkan bahwa 60% responden global percaya diri dapat menggunakan AI secara efektif, namun hanya 39% yang pernah menerima pelatihan AI dan 48% mengaku tidak paham kapan dan bagaimana AI digunakan dalam aplikasi sehari-hari. Ketimpangan antara self-efficacy yang tinggi dan pengetahuan yang rendah inilah yang menciptakan ilusi kompetensi, mendorong karyawan menggunakan alat AI public tanpa menyadari risiko keamanan siber yang terlibat.

Dalam kerangka *Protection Motivation Theory* (PMT), *AI Literacy* berperan sebagai antesedan kunci yang membentuk coping appraisal, khususnya *self efficacy* dan *response self efficacy*. Penelitian oleh Lee et al (2025) membuktikan bahwa literasi AI meningkatkan kemampuan pengguna untuk mengadopsi strategi adaptif seperti verifikasi informasi dan perlindungan privasi, yang pada gilirannya memperkuat keyakinan bahwa mereka mampu menggunakan AI secara bertanggung jawab. Demikian pula, Lyu & Salam, (2025) menemukan bahwa pelatihan berbasis AI yang

dipersonalisasi secara signifikan meningkatkan self-efficacy peserta didik dewasa dalam menggunakan alat AI untuk pembelajaran, yang berimplikasi langsung pada konteks profesional.

Dengan demikian, *AI Literacy* dalam penelitian ini dipahami sebagai kompetensi kognitif dan praktis yang memungkinkan karyawan membuat keputusan rasional terkait penggunaan AI, karyawan akan memilih apakah melalui saluran resmi dari perusahaan atau melalui alat pribadi yang berisiko. Tanpa literasi yang memadai, karyawan akan cenderung mengandalkan asumsi yang bersifat subjektif tentang keamanan AI, sehingga meningkatkan kerentanan terhadap ancaman *Shadow AI*.

2.1.6 Applied Range Theory - Organizational Support

Organizational Support mengacu pada sejauh mana perusahaan menyediakan infrastruktur, kebijakan, pelatihan dan sumber daya yang memadai untuk memungkinkan karyawan menggunakan teknologi khususnya AI secara aman, efektif dan sesuai kebijakan (Li et al., 2022; Tran et al., 2025). Dalam konteks penelitian ini, variable ini mencakup tiga dimensi utama yaitu ketersediaan alat AI resmi yang setara dengan solusi AI publik seperti Chat GPT, Copilot dll, penyelenggaraan pelatihan keamanan siber dan literasi AI secara berkala dan kejelasan kebijakan organisasi mengenai penggunaan teknologi generatif. Tanpa dukungan organisasi yang memadai, karyawan cenderung mengambil inisiatif sendiri dengan menggunakan alat AI pihak ketiga, yang berpotensi menimbulkan risiko kebocoran data, fenomena ini dikenal sebagai *Shadow AI*. Data empiris dari (Microsoft and LinkedIn, 2024) pada Laporan *Work Trend Index 2024* memperkuat urgensi variabel ini, hanya 39% karyawan global yang menerima pelatihan AI dari perusahaan, sementara 68% merasa kewalahan oleh volume dan kecepatan pekerjaan. Kesenjangan ini menciptakan tekanan untuk mencari solusi cepat sehingga sering kali karyawan mencari solusi melalui *Shadow AI*. Dengan demikian, *Organizational Support* bukan sekedar faktor pendukung, melainkan

penentu utama apakah karyawan memilih jalur yang aman atau berisiko dalam mengadopsi Gen AI.

2.1.7 Applied Range Theory - *Perceived Severity*

Perceived Severity dalam kerangka *Protection Motivation Theory* (PMT) mengacu pada sejauh mana individu mempersepsikan konsekuensi negatif dari suatu ancaman sebagai hal serius atau berbahaya bagi dirinya secara pribadi (Floyd et al., 2000; R. Rogers, 1983). Dalam konteks penelitian ini, *perceived severity* dioperasionalkan sebagai keyakinan karyawan bahwa penggunaan alat AI generatif secara informal (*Shadow AI*) dapat mengakibatkan dampak serius seperti kebocoran data sensitif, sanksi disiplin, kehilangan pekerjaan atau kerusakan reputasi profesional. Item pengukuran yang umum digunakan secara konsisten diadopsi dalam studi keamanan (Li et al., 2022; Sulaiman et al., 2022). Dalam konteks *Shadow AI*, *perceived severity* menjadi variable kritis karena risiko kebocoran data tidak hanya bersifat teknis, tetapi juga personal dan eksistensial seperti kehilangan pekerjaan akibat pelanggaran kebijakan perusahaan atau tuntutan hukum akibat kebocoran data pelanggan. Oleh karena itu, variable ini berperan sebagai komponen utama dalam threat appraisal yang membentuk motivasi protektif karyawan untuk menghindari penggunaan AI ilegal.

2.1.8 Applied Range Theory - *Self Efficacy*

Self Efficacy dalam kerangka *Protection Motivation Theory* (PMT) mengacu pada keyakinan individu bahwa ia mampu melaksanakan tindakan protektif untuk menghadapi ancaman siber (Floyd et al., 2000; R. Rogers, 1983). Dalam konteks penelitian ini, *self efficacy* dioperasionalkan sebagai keyakinan karyawan bahwa mereka mampu menyelesaikan tugas kerja tanpa mengandalkan alat AI generatif pihak ketiga seperti ChatGPT, Copilot, Gemini atau lainnya. Meta-analisis oleh Somestad et al (2015) menunjukkan bahwa *self efficacy* adalah prediktor terkuat dari niat keamanan siber, bahkan lebih dominan dari pada persepsi ancaman.

Temuan ini diperkuat oleh Li et al (2022) yang melaporkan bahwa *self efficacy* memiliki pengaruh signifikan terhadap perilaku perlindungan keamanan siber. Demikian pula, Alrawhani et al (2025) menemukan bahwa *Perceived Self Efficacy* berpengaruh positif terhadap kepatuhan kebijakan keamanan informasi. Dalam konteks *Shadow AI*, *self efficacy* menjadi penentu karyawan yang merasa tidak mampu bekerja efisien tanpa AI eksternal cenderung mengabaikan kebijakan keamanan, meski menyadari risikonya. Oleh karena itu, variabel ini merupakan komponen inti dalam *coping appraisal* yang membentuk motivasi protektif atau *maladaptive*.

2.1.9 Applied Range Theory - Response Self Efficacy

Response Self Efficacy dalam kerangka *Protection Motivation Theory (PMT)* mengacu pada keyakinan individu bahwa tindakan protektif yang direkomendasikan dalam hal ini, penggunaan alat AI resmi perusahaan efektif dalam mengurangi ancaman kebocoran data (Floyd et al., 2000; R. Rogers, 1983). Berbeda dengan *self efficacy* yang berfokus pada kemampuan diri, untuk *response self-efficacy* menilai efektivitas solusi itu sendiri. Meta analisis oleh Sommestad et al., (2015) menunjukkan bahwa *response efficacy* memiliki korelasi dengan niat keamanan siber, menjadikannya salah satu predictor kuat dalam PMT. Temuan ini diperkuat oleh Alrawhani et al (2025) yang melaporkan bahwa *Perceived Response Efficacy* berpengaruh signifikan terhadap kepatuhan kebijakan keamanan informasi. Demikian pula Li et al (2022) menemukan bahwa *response efficacy* berdampak positif terhadap perilaku perlindungan keamanan siber. Dalam konteks penelitian khususnya *Shadow AI*, *response self efficacy* menjadi penentu kritis jika karyawan merasa alat AI internal tidak efektif atau tidak mampu dibandingkan ChatGPT atau Copilot. Mereka akan cenderung mengabaikan kebijakan keamanan. Oleh karena itu, variabel ini merupakan komponen inti juga dalam *coping appraisal* yang membentuk motivasi protektif.

2.1.10 Applied Range Theory - Response Cost

Response Cost dalam kerangka *Protection Motivation Theory* (PMT) mengacu pada persepsi individu mengenai biaya pribadi baik berupa usaha, waktu, ketidaknyamanan, maupun stress yang harus dikeluarkan untuk melaksanakan tindakan protektif terhadap ancaman siber (Floyd et al., 2000; Sommestad et al., 2015). Dalam konteks penelitian ini, *Response Cost* dioperasionalkan sebagai keyakinan karyawan bahwa menghindari penggunaan *Shadow AI* (misalnya, hanya menggunakan alat resmi perusahaan) akan memperlambat pekerjaan, meningkatkan beban kognitif atau mengurangi efisiensi. Meta analisis oleh Sommestad et al., (2015) menunjukkan bahwa *Response Cost* memiliki korelasi negatif yang signifikan dengan niat keamanan siber menjadikannya prediktor. Dalam konteks *Shadow AI*, *Response Cost* menjadi variable kritis karena tekanan produktivitas yang tinggi dimana 68% karyawan merasa kewalahan seperti pada Microsoft and LinkedIn, (2024) pada Laporan *Work Trend Index 2024* yang mendorong karyawan memilih solusi cepat meskipun hal itu berisiko. Jika organisasi tidak menyediakan alat AI resmi yang setara dengan ChatGPT, Copilot, Gemini dan sejenisnya, maka *Response Cost* dari kepatuhan akan terasa sangat tinggi sehingga memicu perilaku maladaptif. Oleh karena itu, *Response Cost* berperan sebagai penghambat utama dalam *coping appraisal* yang menentukan apakah karyawan memilih jalur aman atau berisiko.

2.1.11 Applied Theory - Theory of Planned Behavior (TPB)

Sebagai *applied theory*, *Theory of Planned Behavior* (TPB) oleh Ajzen (1991) berperan sebagai jembatan antara proses kognitif pada *Protection Motivation Theory* (PMT) dan perilaku nyata. TPB menyatakan bahwa niat (*intention*) adalah prediktor utama perilaku dan niat ini dibentuk oleh sikap, norma subjektif dan kontrol perilaku yang dirasakan. Dalam penelitian ini, TPB membenarkan penggunaan *Intention to use Shadow AI* sebagai mediator antara persepsi kognitif pada PMT dan perilaku aktual. Integrasi PMT dan TPB telah diuji secara empiris oleh Tran et al (2025)

dalam konteks kepatuhan keamanan siber di Vietnam, yang menemukan bahwa kesadaran membentuk sikap yang kemudian membentuk niat dan perilaku protektif. Pendekatan serupa diadopsi oleh Alshammari & Almamary (2025) yang menggabungkan PMT, TPB dan *Operant Conditioning Theory* untuk menjelaskan kepatuhan kebijakan keamanan informasi. Sommestad et al (2019) juga menemukan bahwa penambahan variabel seperti *anticipated regret* dan *habit* ke dalam TPB meningkatkan daya prediksi model secara signifikan. Dalam konteks *Shadow AI*, *intention* menjadi variabel kritis karena mencerminkan rasionalisasi diam-diam karyawan sebelum mengambil keputusan untuk menggunakan AI informal, sehingga menjembatani kesenjangan antara pemahaman risiko dan tindakan nyata.

2.1.12 Applied Theory - *Intention to Use Shadow AI*

Intention to use Shadow AI mengacu pada niat karyawan untuk menggunakan alat kecerdasan buatan generatif (Gen AI) secara informal melalui akun atau *platform* pribadi tanpa sepengetahuan atau persetujuan dari departemen TI perusahaan. Dalam kerangka teoretis, variabel ini berperan sebagai mediator kritis antara proses kognitif individu (*threat appraisal* dan *coping appraisal*) dalam *Protection Motivation Theory* (PMT) dan perilaku nyata (*Shadow AI Security Behavior*). Konsep ini didasarkan pada prinsip dasar *Theory of Planned Behavior* (TPB) yang menyatakan bahwa *intention* adalah prediktor utama perilaku, terutama dalam konteks tindakan sukarela dan berisiko (Ajzen, 1991; Sommestad et al., 2019). Penelitian oleh (Tran et al., 2025) secara eksplisit menguji peran mediasi *intention* dalam hubungan antara kesadaran keamanan siber dan perilaku protektif dan menemukan bahwa niat memiliki pengaruh signifikan terhadap tindakan nyata. Demikian pula, Kiran et al (2025) membuktikan bahwa *intention* memediasi hubungan antara persepsi PMT dan perilaku keamanan siber, dengan kekuatan prediktif yang diperkuat melalui algoritma *machine learning*. Dalam konteks *Shadow AI*, *intention* mencerminkan rasionalisasi diam-diam karyawan, meskipun mereka

menyadari risiko kebocoran data, mereka tetap berniat untuk menggunakan ChatGPT, Gemini atau Copilot dan lain-lainnya karena tekanan produktivitas, kurangnya alternatif resmi atau keyakinan bahwa risiko itu kecil bagi mereka.

Data dari Cyberhaven labs (2024) yang tertuang pada *AI Adoption and Risk Report* menunjukkan bahwa 57% karyawan sengaja menyembunyikan penggunaan AI dari atasan (Linkedin, 2025), yang mengindikasikan adanya kesenjangan antara kesadaran risiko dan niat bertindak, dan pada titik inilah intention menjadi variable penjelas utama. Oleh karena itu, dalam model penelitian ini, *Intention to Use Shadow AI* bukan sekadar sikap, melainkan komitmen psikologis yang menjembatani pemahaman kognitif dan tindakan nyata, sehingga menjadi kunci untuk merancang intervensi keamanan siber yang efektif.

2.1.13 Applied Theory - *Shadow AI Security Behaviour*

Shadow AI Security Behavior mengacu pada frekuensi dan intensitas penggunaan nyata alat kecerdasan buatan generatif (GenAI) secara informal oleh karyawan melalui akun atau platform pribadi tanpa sepengetahuan, izin atau pengawasan resmi dari departemen TI perusahaan. Dalam penelitian ini, variabel ini berperan sebagai variabel dependen (Y) yang merepresentasikan *maladaptive behavior* sukarela dalam kerangka *Protection Motivation Theory* (PMT), yaitu tindakan yang secara aktif melanggar kebijakan keamanan siber demi memenuhi tuntutan produktivitas pribadi (Balogun et al., 2025; Puthal et al., 2025). Perilaku ini mencakup aktivitas seperti mengunggah dokumen internal ke ChatGPT, menggunakan Copilot pribadi untuk menulis email kerja atau memproses data pelanggan melalui Gemini tanpa otorisasi. Data dari Cyberhaven labs (2024) menunjukkan bahwa 73,8% penggunaan ChatGPT di tempat kerja dilakukan melalui akun non korporat, dan 27,4% data perusahaan yang dikirim ke AI bersifat sensitive, termasuk *source code*, catatan layanan pelanggan dan dokumen hukum. Temuan ini diperkuat oleh IBM (2025) dan

Kompas (2025a) yang melaporkan bahwa lebih dari sepertiga karyawan secara diam-diam membagikan informasi pekerjaan sensitif ke alat AI publik.

Dalam literatur keamanan siber perilaku nyata (*actual behavior*) sering kali diukur sebagai kelanjutan dari intention, sebagaimana diuji oleh Kiran et al (2025) yang menemukan bahwa *intention to secure device* memiliki hubungan signifikan dengan perilaku keamanan aktual. Demikian pula Tran et al (2025) membuktikan bahwa niat kepatuhan membentuk perilaku protektif. Dalam konteks *Shadow AI*, hubungan ini dibalik, niat untuk menggunakan AI informal secara langsung mendorong perilaku pelanggaran kebijakan, sehingga *Shadow AI Security Behavior* menjadi indikator kritis dari kegagalan tata kelola keamanan siber organisasi.

2.2 Penelitian Terdahulu

Di bawah ini merupakan penelitian terdahulu yang digunakan oleh peneliti sebagai referensi penelitian :

Tabel 2.1 Penelitian Terdahulu

No	Penulis (Peneliti)	Judul Paper	Jurnal dan Edisi	Temuan
1	(Kiran et al., 2025)	Explanatory and predictive modeling of cybersecurity behaviors using protection motivation theory	Journal Computers & Security Vol 149 (2025)	PMT adalah kerangka kuat. Algoritma KNN menunjukkan kekuatan prediktif tertinggi. Komponen <i>coping appraisal</i> (efikasi diri & efikasi respons) lebih kuat dari threat appraisal.

No	Penulis (Peneliti)	Judul Paper	Jurnal dan Edisi	Temuan
2	(Alrawhani et al., 2025)	Evaluating the role of protection motivation theory in information security policy compliance: Insights from the banking sector using PLS-SEM approach	Journal of Open Innovation: Technology, Market, and Complexity Vol 11 (2025)	<i>Perceived Self-Efficacy</i> , <i>Perceived Response Efficacy</i> , dan <i>Perceived Severity</i> berpengaruh signifikan. <i>Perceived Response Cost</i> dan <i>Perceived Vulnerability</i> tidak berpengaruh signifikan.
3.	(Lee et al., 2025)	<i>Generative AI</i> risks and resilience: How users adapt to hallucination and privacy challenges	Telematics and Informatics Reports Vol 19 (2025)	Risiko GenAI tidak selalu menghalangi penggunaan. Risiko memicu perilaku adaptif. Verifikasi informasi secara signifikan memoderasi hubungan antara sikap dan niat penggunaan.
4.	(Sulaiman et al., 2022)	Cybersecurity Behavior among Government Employees: The	Information (Switzerland) Vol 13 (2022)	Faktor <i>threat appraisal</i> dan <i>coping appraisal</i> PMT secara

No	Penulis (Peneliti)	Judul Paper	Jurnal dan Edisi	Temuan
		Role of Protection Motivation Theory & Responsibility in Mitigation Cyberattacks		signifikan memengaruhi perilaku keamanan siber. Kesadaran dan kebiasaan memiliki peran penting.
5.	(Li et al., 2022)	The effects of antecedents and mediating factors on cybersecurity protection behavior	Computers in Human Behavior Reports Vol 5 (2022)	Kesadaran berperan penting dalam membentuk sikap dan niat. Sikap memiliki hubungan positif dengan niat kepatuhan dan perilaku protektif. Niat membentuk perilaku protektif.
6.	(Puthal et al., 2025)	<i>Shadow AI: Cyber Security Implications, Opportunities and Challenges in the Unseen Frontier</i>	SN Computer Science Vol 6 (2025)	<i>Shadow AI</i> secara signifikan meningkatkan ancaman keamanan siber, seperti model <i>poisoning</i> dan <i>data leakage</i> . Mengidentifikasi tiga kategori ancaman utama.
7.	(Balogun et al., 2025)	The Ethical and Legal Implications	Journal of Engineering	Pelanggaran privasi paling

No	Penulis (Peneliti)	Judul Paper	Jurnal dan Edisi	Temuan
		of <i>Shadow AI</i> in Sensitive Industries: A Focus on Healthcare, Finance and Education	Research and Reports Vol 27 (2025)	sering di pendidikan. Isu bias paling dominan di keuangan. Risiko keamanan siber tertinggi di layanan kesehatan.
8.	(Tran et al., 2025)	From awareness to behaviour: understanding cybersecurity compliance in Vietnam	International Journal of Organizational Analysis Vol 33 (2025)	Kesadaran berperan penting dalam membentuk sikap dan niat. Sikap memiliki hubungan positif dengan niat kepatuhan dan perilaku protektif. Niat membentuk perilaku protektif.
9.	(Jan et al., 2023)	A Meta-Analysis of Research on Protection Motivation Theory	Journal of Applied Social Psychology	Efek ukuran rata-rata ($d^+ = 0.52$) menunjukkan pengaruh moderat. Semua konstruk PMT memiliki hubungan signifikan dan sesuai arah teoritis dengan niat/perilaku. <i>Self-efficacy</i> dan

No	Penulis (Peneliti)	Judul Paper	Jurnal dan Edisi	Temuan
				<i>response efficacy</i> memiliki efek terkuat. Efek lebih besar terlihat pada tahap " <i>cessation</i> " dibandingkan " <i>initiation</i> ".
10.	(Sommestad et al., 2015)	A Meta-Analysis of Studies on Protection Motivation Theory and Information Security Behaviour	International Journal of Information Security and Privacy Vol 9 (2015)	PMT menjelaskan sekitar 40% variasi dalam niat. Efek lebih kuat jika perilaku bersifat sukarela, spesifik, dan ancumannya ditujukan langsung kepada individu. Variabel <i>coping appraisal</i> (<i>self-efficacy</i> , <i>response efficacy</i>) lebih penting untuk perilaku sukarela, sedangkan threat appraisal lebih penting untuk perilaku wajib.
11.	(Sommestad et al., 2019)	The Theory of Planned Behavior and Information	Journal of Computer	Model inti TPB menjelaskan 43.3% varians niat

No	Penulis (Peneliti)	Judul Paper	Jurnal dan Edisi	Temuan
		Security Policy Compliance	Information Systems Vol 59 (2019)	kepatuhan. Penambahan variable <i>Anticipated Regret</i> meningkatkan penjelasan varians sebesar 3.4 poin persentase, dan <i>Habit</i> sebesar 2.6 poin persentase. <i>Perceived Behavioural Control</i> tidak signifikan dalam model ini.
12.	(Bankins et al., 2023)	A multilevel review of artificial intelligence in organizations: Implications for organizational behavior research and practice	Journal of Organizational Behavior	Literatur AI organisasi masih didominasi perspektif makro/teknis; secara eksplisit menyerukan pergeseran riset ke level individu untuk memetakan mekanisme psikologis, kepatuhan, dan respons perilaku karyawan terhadap AI.

No	Penulis (Peneliti)	Judul Paper	Jurnal dan Edisi	Temuan
13.	(Budhwar et al.,2023)	Human resource management in the age of generative artificial intelligence: Perspectives and research directions	Human Resource Management Journal	Menyoroti kesenjangan antara kecepatan adopsi GenAI dan kesiapan perilaku/keamanan karyawan; menekankan perlunya integrasi faktor kontekstual (<i>AI literacy, organizational support</i>) dalam tata kelola SDM digital.
14.	(Naqbi et al., 2024)	Enhancing Work Productivity through Generative Artificial Intelligence: A Comprehensive Literature Review	Jurnal/Proceeding Terkait Manajemen Teknologi & Produktivitas	Mengonfirmasi adanya <i>trade-off</i> antara dorongan produktivitas instan dan kepatuhan keamanan, yang menjadi pendorong utama karyawan memilih penggunaan AI informal (<i>Shadow AI</i>) meskipun risiko telah disadari.

2.2.1 Kesenjangan Penelitian (*Research Gap*)

Penelitian ini dirancang untuk mengisi beberapa kesenjangan dari penelitian terdahulu. Berdasarkan kerangka Miles (2017), kesenjangan ini dapat di klasifikasikan ke dalam beberapa jenis yaitu *Evidence Gap*, *Theoretical Gap*, *Empirical Gap* dan *Population Gap*, dimana masing-masing memiliki kontribusi pada keunikan dan signifikansi studi ini.

2.2.1.1 *Evidence Gap*

Salah satu *evidence gap* yang teridentifikasi adalah ketidakkonsistenan temuan empiris mengenai peran *Perceived Vulnerability* dalam kerangka *Protection Motivation Theory* (PMT). Meta-analisis oleh Sommestad et al., (2015) dan studi terkini seperti Alrawhani et al., (2025) serta Kiran et al., (2025) secara konsisten melaporkan bahwa persepsi kerentanan pribadi memiliki korelasi lemah atau tidak signifikan dalam memprediksi niat dan perilaku keamanan siber, terutama pada konteks kepatuhan sukarela. Namun, beberapa penelitian tetap mempertahankan konstruk ini sebagai variabel inti tanpa justifikasi kontekstual yang memadai, sehingga menimbulkan ambiguitas mengenai relevansi empirisnya.

Dalam fenomena *Shadow AI*, ketidakkonsistenan ini semakin mengemuka. Studi oleh Li et al (2022) dan Tran et al (2025) mengungkap adanya paradoks kognitif: karyawan menyadari risiko kebocoran data dan pelanggaran kebijakan, namun tidak merasa rentan secara personal karena dampak tersebut lebih dipersepsikan sebagai liabilitas organisasional. Sementara itu, Puthal et al. (2025) dan Balogun et al (2025) telah memetakan kerentanan teknis dan implikasi makro *Shadow AI*, namun tidak menjelaskan mengapa individu tetap mengabaikan ancaman tersebut. PMT tradisional mengasumsikan bahwa *threat appraisal* harus mencakup dimensi *severity* dan *vulnerability* secara simultan, namun bukti empiris pada konteks teknologi sukarela menunjukkan dominasi *perceived*

severity. Oleh karena itu, penelitian ini mengisi *evidence gap* tersebut dengan mengecualikan *Perceived Vulnerability* dari model, sebuah keputusan yang didukung oleh konsistensi temuan meta-analisis dan relevansi fenomenal *Shadow AI* di lingkungan kerja digital.

2.2.1.2 Theoretical Gap

Kesenjangan teoretis utama yang disebut dalam penelitian ini adalah perluasan *Protection Motivation Theory* (PMT) dari ranah perilaku protektif (*protective behaviour*) menuju perilaku maladaptif sukarela (*voluntary maladaptive behaviour*). Secara historis, PMT telah divalidasi secara kuat untuk menjelaskan kepatuhan terhadap kebijakan keamanan atau tindakan preventif kesehatan (Jan et al., 2023; Sommestad et al., 2019). Namun, dalam konteks *Shadow AI*, karyawan secara sadar memilih untuk melanggar kebijakan keamanan bukan karena ketidaktahuan, melainkan untuk memenuhi tuntutan produktivitas, mengatasi keterbatasan alat resmi, atau mengoptimalkan efisiensi kerja. Sommestad et al (2015) menegaskan bahwa PMT memang lebih prediktif pada perilaku sukarela, namun aplikasi teori ini untuk menjelaskan *voluntary violation* yang dimediasi oleh tekanan kerja dan kesenjangan dukungan organisasional masih belum dieksplorasi secara mendalam.

Selain itu, tinjauan mutakhir oleh Bankins et al (2023) dan Budhwar et al., (2023) menyoroti bahwa literatur adopsi AI organisasi masih didominasi perspektif makro, strategis, dan teknis, sementara mekanisme psikologis dan keamanan siber di level individu tetap menjadi *blind spot*. Penelitian oleh Puthal et al (2025) serta Balogun et al (2025) berhasil mengidentifikasi *Shadow AI* sebagai ancaman sistemik dan regulasi, namun berhenti pada level organisasional tanpa menyentuh proses kognitif yang mendorong

keputusan pelanggaran. Penelitian ini mengisi *theoretical gap* tersebut dengan mengembangkan model PMT yang diperluas, yang tidak hanya memetakan mengapa karyawan mematuhi kebijakan, tetapi juga mengurai mekanisme *threat* dan *coping appraisal* yang mendorong niat dan tindakan menggunakan AI informal meskipun risiko telah disadari.

2.2.1.3 Empirical Gap

Secara empiris, sebagian besar penelitian terdahulu cenderung menguji hubungan parsial atau langsung antar konstruk, misalnya bagaimana *AI Literacy* memengaruhi *Self-Efficacy* atau bagaimana *Organizational Support* memengaruhi kepatuhan kebijakan. Jarang terdapat studi yang menguji model terintegrasi yang secara simultan memetakan jalur dari faktor kontekstual (*contextual antecedents*) → proses kognitif mediasi (*cognitive appraisal*) → niat (*intention*) → perilaku aktual (*behaviour*) dalam satu kerangka statistik yang rigor. Penelitian oleh Li et al (2022) dan Tran et al (2025) telah meletakkan fondasi jalur mediasi *awareness–intention–behaviour*, namun masih terbatas pada konteks keamanan siber umum dan belum mengintegrasikan mekanisme *appraisal* kognitif PMT maupun spesifikasi teknologi *Generative AI*.

Di sisi lain, Kiran et al (2025) telah memvalidasi pendekatan *explanatory-predictive* menggunakan PLS-SEM untuk perilaku keamanan siber, namun fokus penelitiannya masih pada kepatuhan standar dan prediksi algoritmik, bukan pada fenomena *Shadow AI* yang bersifat sukarela dan maladaptif. Dengan demikian, penelitian ini mengisi *empirical gap* dengan menguji model terintegrasi delapan variabel yang menggabungkan *Cybersecurity Awareness*, *AI Literacy*, dan *Organizational Support* sebagai anteseden kognitif, *Perceived Severity*, *Response Cost*, *Self-Efficacy*, dan *Response Self-Efficacy* sebagai mekanisme mediasi PMT, serta *Intention* dan

Shadow AI Security Behaviour sebagai outcome. Pengujian model ini menggunakan PLS-SEM yang divalidasi secara metodologis untuk menangkap kompleksitas hubungan komplementer dan mediasi penuh, sebagaimana direkomendasikan dalam literatur pemodelan perilaku teknologi yang terkini.

2.2.1.4 Population Gap

Kesenjangan populasi yang signifikan terletak pada minimnya penelitian yang berfokus pada pekerja profesional di perusahaan digital negara berkembang, khususnya di Indonesia. Mayoritas studi mengenai *Shadow AI*, tata kelola GenAI, dan perilaku keamanan siber karyawan dilakukan di negara-negara WEIRD (*Western, Educated, Industrialized, Rich, Democratic*) seperti Amerika Serikat, Eropa, atau Australia, yang memiliki ekosistem regulasi siber yang matang, infrastruktur TI terpusat, dan kebijakan penggunaan AI yang sudah terstandarisasi. Kondisi ini berbeda secara fundamental dengan lanskap digital di Indonesia, di mana hanya 31% perusahaan memiliki kebijakan AI yang eksplisit (Amazon, 2025) dan 77% karyawan menyatakan kebingungan mengenai batasan penggunaan alat AI di tempat kerja (Cyberhaven labs, 2024).

Wilayah Jabodetabek sebagai episentrum startup dan perusahaan digital di Asia Tenggara menawarkan konteks unik: tingkat adopsi teknologi yang sangat tinggi, tekanan kompetitif yang mendorong produktivitas instan, namun diimbangi dengan kematangan tata kelola keamanan siber yang masih dalam tahap perkembangan. Google, Temasek (2024) dan Mayer et al (2025) menegaskan bahwa kesenjangan antara kecepatan adopsi dan kesiapan kebijakan ini menciptakan lingkungan yang rentan terhadap praktik *Shadow AI*. Penelitian ini mengisi *population gap* dengan menargetkan sampel karyawan perusahaan digital di

Jabodetabek, sehingga temuan yang dihasilkan tidak hanya memperkaya literatur perilaku siber di ekonomi digital berkembang, tetapi juga memberikan implikasi manajerial yang kontekstual bagi organisasi yang sedang menavigasi transisi AI tanpa kerangka kepatuhan yang rigid.

2.3 Kerangka Berpikir atau Kerangka Konseptual

Kerangka konseptual penelitian ini dibangun melalui sintesis terhadap sejumlah penelitian terdahulu yang relevan, khususnya dalam domain perilaku keamanan siber, adopsi teknologi AI dan aplikasi *Protection Motivation Theory* (PMT). Model konseptual yang diusulkan merupakan pengembangan dan integrasi dari kerangka-kerangka yang telah diuji secara empiris, dengan modifikasi strategis untuk menjawab fenomena *Shadow AI*.

Model penelitian ini secara eksplisit mengadopsi dan memperluas kerangka kerja dari Li et al (2022) yang menguji pengaruh *moral intensity*, *organizational commitment* dan status *digital native* terhadap perilaku perlindungan keamanan siber, dengan PMT sebagai mediator. Studi tersebut membuktikan bahwa faktor kontekstual organisasi dan individu secara signifikan membentuk persepsi PMT. Namun, Li et al (2022) tidak menguji secara konstruk seperti *AI Literacy* dan *Organizational Support* dalam konteks GenAI, sehingga penelitian ini mengisinya dengan memposisikan kedua variabel tersebut sebagai antesedan utama.

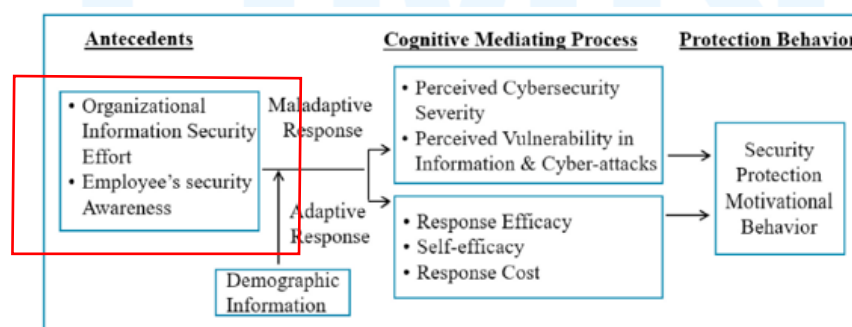
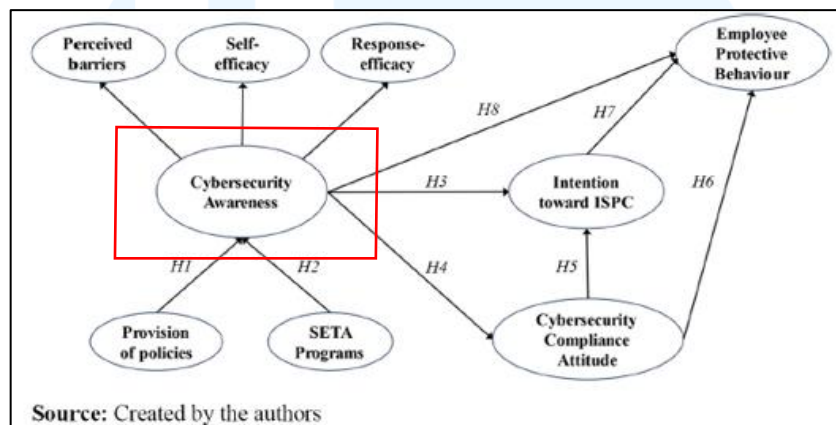


Fig. 1. Extended PMT model for antecedents and mediating factors on cybersecurity actions.

Gambar 2.1 *Framework* Penelitian Terdahulu Pertama

Sumber : Li et al (2022)

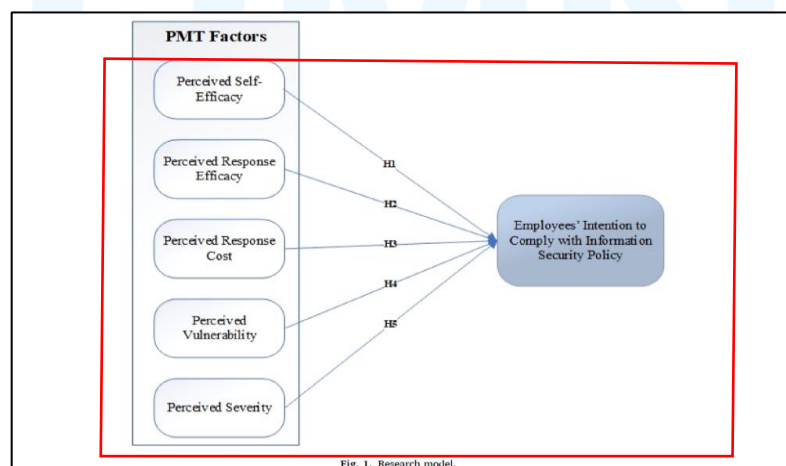
Framework kedua diadopsi dari Tran et al (2025) memberikan fondasi kuat untuk peran *Cybersecurity Awareness* sebagai variabel yang membentuk sikap dan niat kepatuhan. Namun, Tran et al (2025) menggabungkan *self efficacy* dan *response efficacy* ke dalam bentuk konstruk kesadaran, sedangkan penelitian ini memisahnya sebagai bagian dari *coping appraisal* (*Self Efficacy*, *Response Self Efficacy* dan *Response Cost*) untuk menjaga kesesuaian dengan struktur PMT klasik.



Gambar 2.2 Framework Penelitian Terdahulu Kedua

Sumber : Tran et al (2025)

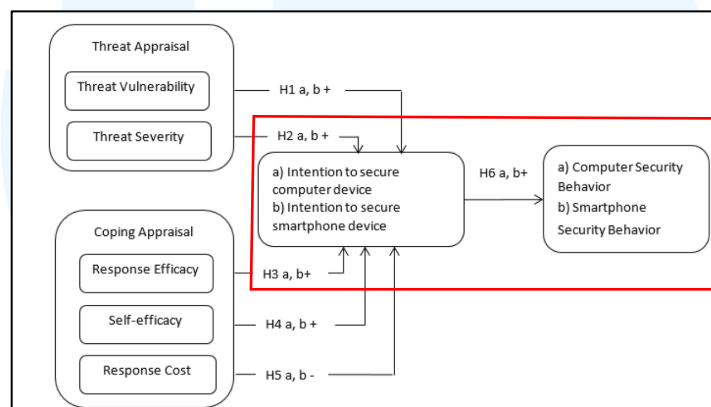
Pada penelitian oleh Alrawhani et al (2025) menemukan bahwa *Perceived Vulnerability* tidak berpengaruh signifikan terhadap niat kepatuhan. Temuan ini menjadi dasar teoretis untuk menghilangkan *Perceived Vulnerability* dari model.



Gambar 2.3 Framework Penelitian Terdahulu Ketiga

Sumber : Alrawhani et al (2025)

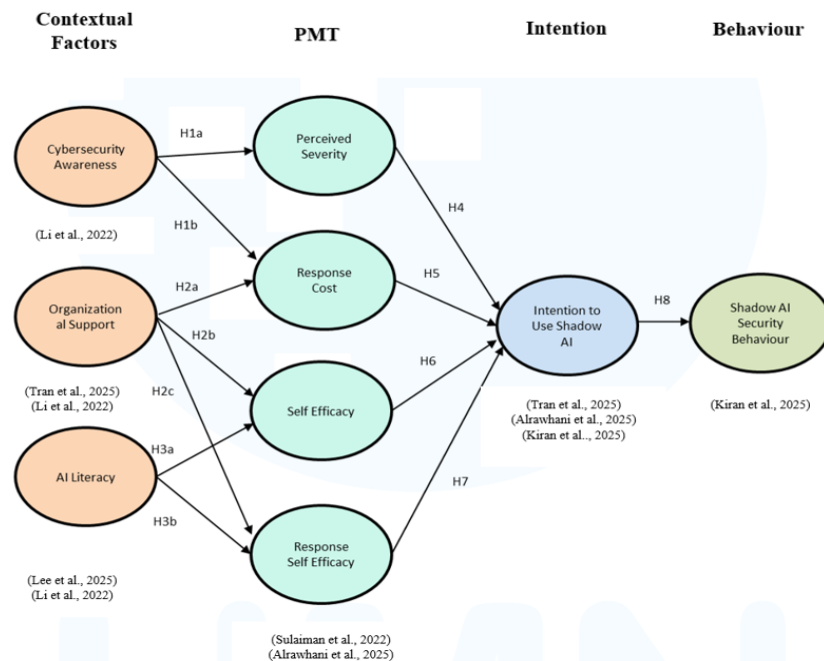
Niat adalah variabel mediasi yang menjembatani proses kognitif dan perilaku nyata. Dalam model ini, *Intention to use Shadow AI* dipengaruhi oleh semua komponen *Threat Appraisal* dan *Coping Appraisal* dalam kerangka PMT. Penelitian sebelumnya oleh Tran et al (2025) menegaskan bahwa niat adalah prediktor perilaku kepatuhan, sementara Kiran et al (2025) juga menemukan bahwa niat juga merupakan mediator kuat dalam model keamanan siber. Dalam konteks *Shadow AI*, niat mencerminkan rasionalisasi karyawan.



Gambar 2.4 *Framework* Penelitian Terdahulu Keempat

Sumber : Kiran et al., (2025)

Dengan menggabungkan beberapa *framework-framework* dari penelitian sebelumnya, maka berikut ini model konseptual yang dikembangkan dalam penelitian ini. Perilaku nyata yaitu frekuensi dan intensitas penggunaan *Shadow AI* adalah *outcome* akhir dari model ini, variabel ini dipengaruhi secara langsung oleh *Intention to use Shadow AI*. Penelitian oleh Kiran et al (2025) dan Balogun et al (2025) mengidentifikasi *Shadow AI* sebagai ancaman sistemik. Model konseptual ini tidak hanya menjelaskan apa yang terjadi, tetapi juga mengapa dan bagaimana penelitian ini dapat memberikan dasar teoretis yang kuat untuk merancang intervensi keamanan siber yang efektif di era *generative AI*, tidak hanya dapat berkontribusi akademis namun juga solusi praktis yang konkret untuk mengelola risiko *Shadow AI* di perusahaan digital di Indonesia.



Gambar 2.5 Kerangka Penelitian

Sumber : Olahan peneliti (2025)

Secara filosofis, penelitian ini berpijak pada *Expectancy Value Theory* (EVT) yang menyatakan bahwa motivasi terbentuk dari interaksi antara ekspektansi terhadap outcome dan nilai yang melekat pada outcome tersebut (Vroom, 1964). Rogers (1975) mengoperasionalkan paradigma ini ke dalam *Protection Motivation Theory* (PMT), di mana dimensi *expectancy* diterjemahkan menjadi *coping appraisal* (keyakinan individu terhadap kemampuan diri dan efektivitas solusi protektif), sedangkan dimensi *value* direpresentasikan sebagai *threat appraisal* (penilaian terhadap keseriusan dan kemungkinan ancaman yang ingin dihindari). Dalam konteks *Shadow AI*, karyawan perusahaan digital secara simultan menimbang *positive value* (efisiensi, kecepatan, eksplorasi fitur AI) dan *negative value* (risiko kebocoran data, sanksi organisasi). Ketika *coping appraisal* lebih dominan daripada *threat appraisal*, motivasi protektif melemah dan muncul intention untuk menggunakan AI publik secara informal. EVT dengan demikian memberikan landasan makro mengapa *appraisal* kognitif terbentuk, sedangkan PMT menyediakan mekanisme mikro yang memetakan bagaimana *appraisal* tersebut terkristalisasi menjadi niat dan perilaku aktual.

Sehingga dapat dinyatakan model konseptual penelitian ini disusun melalui dekonstruksi sistematis terhadap empat kerangka terdahulu. Pertama, struktur *threat* dan *coping appraisal* dipertahankan sebagai mediator inti sesuai rekomendasi meta-analisis Sommestad et al (2015) yang menunjukkan PMT paling prediktif pada *voluntary behaviour*. Kedua, *Perceived Vulnerability* dieliminasi berdasarkan konsistensi temuan empiris Alrawhani et al (2025) dan Kiran et al (2025) yang melaporkan ketidaksignifikannya pada konteks adopsi teknologi sukarela. Ketiga, faktor kontekstual (*Cybersecurity Awareness, AI Literacy, Organizational Support*) diintegrasikan sebagai anteseden pembentuk appraisal, mengadopsi logika mediasi Li et al., (2022) dan Tran et al (2025) namun menutup celah teoretis yang disoroti Bankins et al (2023) mengenai minimnya penghubungan dinamika organisasi dengan mekanisme psikologis individu. Keempat, orientasi *output* digeser dari *compliance behaviour* menuju *Shadow AI Security Behaviour*, dengan Intention berfungsi sebagai jembatan konseptual ke *Theory of Planned Behavior* (Ajzen, 1991), sejalan dengan validasi integrasi PMT-TPB oleh Kiran et al (2025). Kelima, konteks *policy vacuum* dan tekanan produktivitas di perusahaan digital Jabodetabek dimasukkan sebagai *boundary condition* yang menjelaskan pelemahan *threat appraisal*, mengonfirmasi temuan Lee et al (2025) bahwa pengguna GenAI mengabaikan ancaman abstrak ketika strategi koping mandiri dan urgensi operasional dianggap lebih relevan. Dekonstruksi ini menghasilkan model terintegrasi yang diuji secara rigor melalui *explanatory-predictive PLS-SEM*, sehingga tidak hanya mengonfirmasi teori yang ada, tetapi juga memperluas aplikabilitasnya ke *voluntary maladaptive behaviour* di ekonomi digital berkembang.

2.4 Hipotesis

2.4.1 *Cybersecurity Awareness* memiliki pengaruh terhadap *Perceived Severity*

Dalam konteks keamanan siber, *Cybersecurity Awareness* mengacu pada pemahaman individu terhadap risiko dan konsekuensi pelanggaran data serta tindakan protektif yang diperlukan (Li et al., 2022; Sulaiman et

al., 2022). Menurut *Protection Motivation Theory* (PMT), kesadaran risiko adalah fondasi terbentuknya *threat appraisal*, yaitu penilaian kognitif individu terhadap tingkat keseriusan ancaman R. Rogers (1975). Semakin tinggi tingkat kesadaran seseorang terhadap keamanan siber, semakin besar pula kemampuannya untuk memahami konsekuensi negative dari pelanggaran, seperti kebocoran data sensitif, pelanggaran privasi atau sanksi disiplin (Tran et al., 2025). Li et al (2022) menunjukkan bahwa kesadaran keamanan memiliki pengaruh positif terhadap *perceived severity*, individu yang paham bahaya siber cenderung menilai ancaman tersebut lebih serius. Demikian pula, penelitian Sulaiman et al (2022) dan Alrawhani et al (2025) menegaskan bahwa pemahaman karyawan tentang risiko keamanan meningkatkan persepsi keseriusan ancaman. Dalam konteks *Shadow AI*, kesadaran keamanan siber mendorong karyawan untuk memandang risiko kebocoran data melalui alat AI publik sebagai sesuatu yang serius, bukan sekedar kesalahan teknis. Dengan kata lain semakin tinggi kesadaran keamanan, semakin tinggi pula persepsi keseriusan ancaman terhadap pelanggaran tersebut.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H1a : *Cybersecurity Awareness* berpengaruh terhadap *Perceived Severity*

2.4.2 *Cybersecurity Awareness* memiliki pengaruh terhadap *Response Cost*

Dalam kerangka *Protection Motivation Theory* (PMT), *response cost* menggambarkan persepsi individu terhadap besarnya usaha, waktu atau hambatan yang harus dihadapi ketika melakukan tindakan protektif terhadap ancaman siber (Floyd et al., 2000). Tingkat *Cybersecurity Awareness* yang tinggi membuat karyawan lebih memahami nilai dan pentingnya keamanan data, sehingga tindakan pencegahan tidak lagi dianggap sebagai beban, tetapi sebagai bagian dari tanggung jawab profesional. Penelitian oleh Li et al (2022) menunjukkan bahwa kesadaran

keamanan siber memiliki pengaruh negative terhadap *response cost*, artinya semakin sadar seseorang terhadap risiko siber, semakin rendah persepsinya terhadap hambatan dalam menerapkan perilaku aman. Temuan serupa juga disampaikan oleh Tran et al (2025) yang menegaskan bahwa individu dengan tingkat kesadaran tinggi lebih siap beradaptasi dengan kebijakan keamanan tanpa merasa terbebani. Dalam konteks *Shadow AI*, karyawan yang sadar risiko tidak akan menganggap larangan penggunaan AI publik sebagai hambatan, melainkan sebagai langkah protektif yang perlu dilakukan.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H1b : *Cybersecurity Awareness* berpengaruh terhadap *Response Cost*

2.4.3 *Organizational Support* memiliki pengaruh terhadap *Response Cost*

Organizational Support mencerminkan sejauh mana perusahaan menyediakan kebijakan, pelatihan dan infrastruktur pendukung untuk penggunaan teknologi yang aman (Li et al., 2022; Tran et al., 2025). Dalam kerangka *Protection Motivation Theory* (PMT). Dukungan organisasi dapat menurunkan *response cost* karena karyawan memperoleh sumber daya dan panduan yang jelas untuk bertindak aman tanpa merasa terbebani. Penelitian Tran et al (2025) menunjukkan bahwa dukungan organisasi yang baik mengurangi persepsi hambatan dalam penerapan perilaku keamanan siber. Begitu pula Li et al (2022) menemukan bahwa fasilitas dan pelatihan meningkatkan efisiensi karyawan dalam mematuhi kebijakan keamanan. Dalam konteks *Shadow AI*, ketika perusahaan menyediakan alat AI resmi dan pelatihan keamanan yang memadai, karyawan tidak akan merasa bahwa kepatuhan adalah beban.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H2a : *Organizational Support* berpengaruh terhadap *Response Cost*

2.4.4 *Organizational Support* memiliki pengaruh terhadap *Self Efficacy*

Dalam kerangka *Protection Motivation Theory* (PMT), *self efficacy* mengacu pada keyakinan individu bahwa ia mampu melakukan tindakan protektif terhadap ancaman R. Rogers (1975). *Organizational Support* yang memadai melalui kebijakan yang jelas, pelatihan keamanan siber dan penyediaan alat AI resmi yang dapat memperkuat keyakinan karyawan bahwa mereka mampu bekerja secara aman tanpa melanggar kebijakan perusahaan. Penelitian Li et al. (2022) menunjukkan bahwa dukungan organisasi meningkatkan *self-efficacy* karyawan dalam menerapkan perilaku keamanan. Tran et al. (2025) juga menegaskan bahwa karyawan dengan dukungan manajerial yang kuat lebih percaya diri dalam mengelola risiko teknologi. Dalam konteks *Shadow AI*, semakin besar dukungan organisasi, semakin tinggi keyakinan karyawan untuk bekerja efisien tanpa harus bergantung pada alat AI pribadi yang berisiko.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H2b : *Organizational Support* berpengaruh terhadap *Self-Efficacy*

2.4.5 *Organizational Support* memiliki pengaruh terhadap *Response Self-Efficacy*

Response Self-Efficacy menggambarkan keyakinan individu terhadap efektivitas tindakan protektif atau solusi yang disediakan untuk menghadapi ancaman siber R. Rogers (1975). Dukungan organisasi yang kuat, seperti penyediaan alat AI resmi, pelatihan keamanan berkala, dan kebijakan penggunaan AI yang jelas, meningkatkan kepercayaan karyawan bahwa solusi internal perusahaan efektif melindungi data mereka. Penelitian Li et al. (2022) menemukan bahwa dukungan organisasi memperkuat persepsi efektivitas tindakan keamanan, sedangkan Tran et al. (2025) menunjukkan bahwa kejelasan kebijakan dan fasilitas yang memadai mendorong keyakinan karyawan terhadap solusi resmi. Dalam konteks

Shadow AI, dukungan organisasi akan menumbuhkan keyakinan bahwa penggunaan alat AI perusahaan lebih aman dibandingkan *platform* publik.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H2c : *Organizational Support* berpengaruh terhadap *Response Self Efficacy*

2.4.6 AI Literacy memiliki pengaruh terhadap Self-Efficacy

AI Literacy mencerminkan kemampuan individu untuk memahami, mengevaluasi dan menggunakan teknologi kecerdasan buatan secara aman dan bertanggung jawab (Gillespie et al., 2025). Dalam kerangka *Protection Motivation Theory* (PMT), literasi yang baik meningkatkan *self efficacy* karena individu merasa lebih mampu mengendalikan risiko dan memanfaatkan teknologi secara efektif. Penelitian Lee et al (2025) menunjukkan bahwa literasi AI yang tinggi memperkuat keyakinan pengguna dalam mengelola risiko privasi dan akurasi output AI. Lyu & Salam (2025) juga menemukan bahwa pelatihan AI secara signifikan meningkatkan *self efficacy* peserta dalam konteks profesional. Dalam konteks *Shadow AI*, karyawan dengan literasi AI tinggi lebih percaya diri menggunakan teknologi sesuai kebijakan, bukan melalui alat public berisiko.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H3a : *AI Literacy* berpengaruh terhadap *Self Efficacy*

2.4.7 AI Literacy memiliki pengaruh terhadap Response Self-Efficacy

Dalam *Protection Motivation Theory* (PMT), *response self efficacy* menggambarkan sejauh mana individu percaya bahwa tindakan atau solusi yang direkomendasikan efektif dalam mengurangi ancaman R. Rogers, (1975). *AI Literacy* berperan penting karena pemahaman yang baik tentang cara kerja, batasan dan risiko AI membuat individu yakin bahwa solusi resmi yang disediakan organisasi dapat diandalkan. Penelitian Lee et al

(2025) menunjukkan bahwa pengguna dengan literasi AI tinggi memiliki keyakinan lebih besar terhadap efektivitas strategi mitigasi risiko. Lyu & Salam (2025) juga menemukan bahwa pelatihan literasi AI secara langsung meningkatkan *response self efficacy* melalui peningkatan kemampuan evaluasi kritis. Dalam konteks *Shadow AI*, semakin tinggi literasi AI karyawan, semakin besar keyakinan mereka bahwa alat AI perusahaan cukup efektif sehingga tidak perlu menggunakan platform eksternal.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H3b : *AI Literacy* berpengaruh terhadap *Response Self Efficacy*

2.4.8 *Perceived Severity* memiliki pengaruh terhadap *Intention to Use Shadow AI*

Dalam *Protection Motivation Theory* (PMT), *perceived severity* menggambarkan sejauh mana individu menilai ancaman sebagai sesuatu yang serius dan berpotensi merugikan diri mereka (Roger, 1983). Semakin tinggi persepsi terhadap keseriusan ancaman, semakin kecil kemungkinan individu berniat melakukan perilaku berisiko. Penelitian Li et al (2022) dan Alrawhani et al (2025) menunjukkan bahwa persepsi ancaman yang tinggi menurunkan niat untuk melanggar kebijakan keamanan. Sulaiman et al (2022) juga menegaskan bahwa karyawan yang menilai ancaman siber lebih serius cenderung menghindari tindakan berisiko. Dalam konteks *Shadow AI*, jika karyawan menilai konsekuensi pelanggaran data sebagai ancaman serius terhadap pekerjaan atau reputasinya, maka niat untuk menggunakan alat AI publik akan menurun.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H4 : *Perceived Severity* berpengaruh terhadap *Intention to Use Shadow AI*

2.4.9 *Response Cost* memiliki pengaruh terhadap *Intention to Use Shadow AI*

Response Cost dalam kerangka *Protection Motivation Theory* (PMT) mencerminkan persepsi individu terhadap beban atau pengorbanan yang dilakukan untuk bertindak sesuai kebijakan keamanan (Floyd et al., 2000). Semakin tinggi persepsi terhadap biaya waktu, tenaga atau ketidaknyamanan dalam mengikuti kebijakan, semakin besar kecenderungan individu untuk mencari jalan pintas atau melakukan perilaku maladaptif. Penelitian Sommestad et al (2015) menunjukkan bahwa *response cost* memiliki pengaruh negatif terhadap niat keamanan, artinya semakin tinggi hambatan yang dirasakan, semakin rendah niat untuk patuh. Dalam konteks *Shadow AI*, jika karyawan merasa penggunaan alat resmi perusahaan memperlambat pekerjaan, mereka akan cenderung berniat menggunakan AI publik yang lebih cepat meski berisiko.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H5 : *Response Cost* berpengaruh terhadap *Intention to Use Shadow AI*

2.4.10 *Self Efficacy* memiliki pengaruh terhadap *Intention to Use Shadow AI*

Self efficacy dalam kerangka *Protection Motivation Theory* (PMT) mengacu pada keyakinan individu terhadap kemampuannya untuk menghadapi ancaman tanpa melanggar kebijakan keamanan R. Rogers (1975). Tingkat *Self efficacy* yang tinggi akan menumbuhkan kepercayaan diri karyawan untuk menyelesaikan tugas tanpa perlu menggunakan alat AI publik yang tidak resmi. Penelitian Li et al (2022) dan Kiran et al (2025) menemukan bahwa *self efficacy* berpengaruh positif terhadap niat berperilaku aman dan kepatuhan terhadap kebijakan keamanan informasi. Dan sebaliknya, individu dengan *self efficacy* rendah cenderung mencari jalan pintas berisiko. Dalam konteks *Shadow AI*, karyawan yang yakin akan

kemampuannya bekerja secara aman akan memiliki niat lebih rendah untuk menggunakan AI publik tanpa izin.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H6 : *Self Efficacy* berpengaruh terhadap *Intention to Use Shadow AI*

2.4.11 *Response Self Efficacy* memiliki pengaruh terhadap *Intention to Use Shadow AI*

Dalam *Protection Motivation Theory* (PMT), *response self efficacy* menggambarkan keyakinan individu bahwa tindakan protektif atau solusi yang direkomendasikan mampu secara efektif mengurangi ancaman (Roger, 1975). Jika karyawan percaya bahwa alat atau kebijakan AI resmi dari perusahaan efektif dan dapat diandalkan, maka niat mereka untuk menggunakan AI publik yang tidak diawasi akan menurun. Penelitian Li et al (2022) dan Alrawhani et al (2025) menunjukkan bahwa *response efficacy* memiliki pengaruh positif terhadap niat berperilaku aman dan kepatuhan terhadap kebijakan keamanan informasi. Sebaliknya, keyakinan rendah terhadap efektivitas solusi organisasi dapat meningkatkan kecenderungan perilaku *Shadow AI*.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H7 : *Response Self Efficacy* berpengaruh terhadap *Intention to Use Shadow AI*

2.4.12 *Intention to Use Shadow AI* memiliki pengaruh terhadap *Shadow AI Security Behavior*

Menurut *Theory of Planned Behavior* (TPB) (Ajzen, 1991) dan diperkuat oleh *Protection Motivation Theory* (PMT), *intention* merupakan prediktor utama perilaku nyata. Semakin kuat niat seseorang untuk melakukan suatu tindakan, semakin besar kemungkinan tindakan tersebut diwujudkan dalam perilaku actual. Penelitian Kiran et al (2025) dan Kiran

et al (2025) membuktikan bahwa niat kepatuhan atau pelanggaran berperan signifikan dalam membentuk perilaku keamanan siber aktual. Dalam konteks *Shadow AI*, ketika karyawan memiliki niat kuat untuk menggunakan alat AI publik karena alasan efisiensi atau kenyamanan, maka perilaku nyata berupa penggunaan AI tanpa izin perusahaan juga akan meningkat.

Berdasarkan temuan penelitian sebelumnya di atas, sehingga hipotesis dapat dibentuk sebagai berikut :

H8 : *Intention to Use Shadow AI* berpengaruh terhadap *Shadow AI Security Behaviour*

