

Terbit online pada laman : <http://teknosi.fti.unand.ac.id/>

Jurnal Nasional Teknologi dan Sistem Informasi

| ISSN (Print) 2460-3465 | ISSN (Online) 2476-8812 |



Artikel Penelitian

Design Science Research on Human–AI Collaboration in Decision Support Systems: A Hybrid Decision-Making Framework

Muhammad Damas Fatih ^{a,*}, Aulia Rachmawati ^b, Hamzah Alghifari ^c, Andicho Haryus Wirasapta ^d

^{a,b,c} Program Studi Sistem Informasi, Universitas Jambi, Jl. Jambi - Muara Bulian No.KM. 15, Mendalo Darat, Kec. Jambi Luar Kota, Kabupaten Muaro Jambi, 36361

^d Program Studi Teknik Elektro, Universitas Jambi, Jl. Jambi - Muara Bulian No.KM. 15, Mendalo Darat, Kec. Jambi Luar Kota, Kabupaten Muaro Jambi, 36361

INFORMASI ARTIKEL

Sejarah Artikel:

Diterima Redaksi: 27 Juni 2026

Revisi Akhir: 20 Agustus 2026

Diterbitkan Online: 29 Agustus 2026

KATA KUNCI

Sistem Pendukung Keputusan,
Kolaborasi Manusia dan AI,
Design Science Research,
Pengambilan Keputusan Hibrida,
Explainable AI

KORESPONDENSI

E-mail: muhammaddamasfatih@unja.ac.id*

A B S T R A C T

"Sistem Pendukung Keputusan (*Decision Support Systems/DSS*) mengalami perkembangan yang pesat seiring kemajuan teknologi kecerdasan buatan (*Artificial Intelligence/AI*), yang memungkinkan organisasi mengolah data dalam jumlah besar dan menghasilkan rekomendasi secara lebih efisien. Meskipun demikian, pengambilan keputusan yang sepenuhnya berbasis AI masih menghadapi berbagai tantangan, seperti kurangnya transparansi, keterbatasan pemahaman konteks, bias algoritmik, serta rendahnya kepercayaan pengguna. Di sisi lain, pengambilan keputusan yang sepenuhnya dilakukan manusia juga rentan terhadap keterbatasan kognitif, kelebihan informasi, dan bias penilaian. Penelitian ini mengusulkan suatu *Human–AI Collaboration Framework* untuk DSS menggunakan pendekatan *Design Science Research (DSR)*. Fokus penelitian adalah merancang dan mengembangkan kerangka konseptual yang mengintegrasikan kemampuan analitis AI dengan keahlian manusia, penalaran kontekstual, dan pertimbangan etis. Framework yang diusulkan terdiri atas tiga lapisan utama, yaitu *AI Layer*, *Human Layer*, dan *Interaction Layer* yang didukung mekanisme *Explainable Artificial Intelligence (XAI)* serta proses umpan balik berkelanjutan. Framework didemonstrasikan secara konseptual melalui skenario pengambilan keputusan dan dievaluasi berdasarkan analisis literatur. Model yang diusulkan memberikan kontribusi terhadap pengembangan kajian *Human–AI Collaboration* dengan menyediakan arsitektur terstruktur untuk pengambilan keputusan hibrida sekaligus memberikan panduan praktis bagi pengembangan DSS yang adaptif, transparan, dan berpusat pada manusia."

1. PENDAHULUAN

Transformasi digital pada berbagai sektor organisasi meningkatkan kebutuhan akan sistem yang dapat mendukung pengambilan keputusan secara cepat, dapat ditelusuri, dan berbasis data. Dalam konteks tersebut, *Decision Support Systems (DSS)* berkembang dari sistem berbasis aturan menjadi lingkungan keputusan yang mengintegrasikan data, model analitis, dan antarmuka pengguna untuk membantu manajer membandingkan alternatif tindakan. Sejarah dan evolusi DSS menunjukkan bahwa tujuan utamanya bukan menggantikan pengambil keputusan, melainkan memperkuat kualitas

pertimbangan manusia melalui informasi yang relevan dan tepat waktu [1], [2].

Perkembangan *Artificial Intelligence (AI)*, *machine learning*, dan *data analytics* memperluas kemampuan DSS untuk mengolah data berskala besar, mengenali pola yang kompleks, menghasilkan prediksi, dan menyusun rekomendasi. Kapabilitas tersebut mendukung keputusan strategis, taktis, maupun operasional ketika organisasi harus merespons perubahan pasar, risiko operasional, dan dinamika lingkungan secara cepat [2], [3]. Namun, nilai AI dalam DSS tidak hanya ditentukan oleh akurasi model, tetapi juga oleh bagaimana rekomendasi dapat dipahami, diuji, dan digunakan secara bertanggung jawab oleh manusia.

Kemajuan AI tidak serta-merta menghilangkan risiko pengambilan keputusan. Model dapat beroperasi sebagai *black box*, mereproduksi bias pada data, atau gagal menangkap konteks sosial, budaya, regulasi, dan etika yang menentukan kesesuaian keputusan. Ketika rekomendasi disajikan tanpa dasar penalaran, tingkat ketidakpastian, atau batas penerapan yang jelas, pengguna dapat menolak sistem secara berlebihan atau justru menerima keluaran AI tanpa evaluasi kritis [4]-[7].

Sebaliknya, pengambilan keputusan yang sepenuhnya bergantung pada manusia juga memiliki keterbatasan. Kapasitas pemrosesan informasi, kelelahan mental, noise, dan bias penilaian dapat mengurangi konsistensi keputusan, khususnya pada situasi dengan banyak alternatif dan ketidakpastian tinggi [8], [9]. Oleh karena itu, persoalan utama bukan memilih manusia atau AI sebagai aktor tunggal, tetapi merancang pembagian peran, batas kewenangan, dan mekanisme interaksi yang membuat keunggulan keduanya saling melengkapi.

Human-AI Collaboration menawarkan pendekatan untuk menjawab ketegangan tersebut. AI berperan menganalisis data, mengenali pola, dan menyajikan alternatif, sedangkan manusia menafsirkan keluaran berdasarkan pengetahuan domain, kondisi kontekstual, nilai organisasi, serta konsekuensi etis. Pendekatan ini selaras dengan *Human-Centered AI* yang menempatkan manusia sebagai pemegang otoritas keputusan, bukan sekadar penerima keluaran sistem [10], [11].

Transparansi merupakan prasyarat penting dalam kolaborasi tersebut. *Explainable Artificial Intelligence* (XAI) dapat membantu pengguna memahami faktor, asumsi, dan bukti yang mendasari rekomendasi AI [6], [7], [17]. Akan tetapi, penjelasan model saja belum cukup. Dalam konteks DSS, penjelasan perlu disajikan sebagai informasi yang relevan bagi keputusan, termasuk tingkat keyakinan atau ketidakpastian, alternatif yang tersedia, serta ruang bagi pengguna untuk menantang atau mengubah rekomendasi.

Walaupun studi tentang *Human-AI Collaboration*, *Human-Centered AI*, dan XAI terus berkembang, pembahasannya masih kerap terfragmentasi pada performa algoritma, perilaku pengguna, atau teknik penjelasan model secara terpisah [10]-[15]. Masih terbatas artefak konseptual yang secara eksplisit memetakan hubungan antara fungsi analitis AI, *affordance* interaksi, kewenangan keputusan manusia, serta mekanisme umpan balik dan akuntabilitas dalam satu arsitektur DSS. Berangkat dari kondisi tersebut, penelitian ini mengajukan tiga pertanyaan desain: (RQ1) kebutuhan desain apa yang diperlukan untuk DSS kolaboratif manusia-AI; (RQ2) bagaimana kebutuhan tersebut dapat dikonfigurasi sebagai *framework*; dan (RQ3) bagaimana artefak dapat didemonstrasikan serta dievaluasi secara analitis sebelum pengujian empiris dilakukan.

1.1. Gap Penelitian dan Posisi Penelitian

Literatur terdahulu telah membahas penerapan AI, *Human-AI Collaboration*, *Human-Centered AI*, dan XAI dalam konteks pengambilan keputusan [10]-[15]. Kesenjangan penelitian tidak terletak pada ketiadaan setiap konsep tersebut, melainkan pada terbatasnya artefak DSS yang menghubungkan konsep-konsep itu menjadi keputusan yang dapat dijelaskan, diperdebatkan, diaudit, dan tetap berada di bawah otoritas manusia. Dengan demikian,

posisi penelitian ini adalah merancang konfigurasi komponen dan mekanisme kolaborasi, bukan mengklaim bahwa tidak ada penelitian lain yang pernah membahas salah satu komponennya.

Tabel 1 memetakan kontribusi penelitian terdahulu, keterbatasan desain yang masih tersisa, dan posisi *framework* yang diusulkan.

Tabel 1. Pemetaan penelitian terdahulu dan implikasi desain *framework*

Referensi	Fokus dan kontribusi	Keterbatasan desain / implikasi
Power [1]	DSS: memperkuat pertimbangan manajerial.	Belum membahas opacity AI, penjelasan, atau feedback.
Dietvorst et al. [13]	Trust: kesalahan AI dapat menurunkan reliance.	Trust perlu dikalibrasi, bukan diasumsikan.
Cai et al. [10]	Human-AI collaboration: koordinasi peran mendukung keputusan bersama.	Belum memetakan arsitektur DSS lintas domain.
Shneiderman [11]	Human-centered AI: kontrol manusia adalah prinsip utama.	Belum menguraikan komponen operasional DSS.
Arrieta et al. [17]	XAI: interpretabilitas mendukung transparansi.	Penjelasan belum selalu dikaitkan dengan keputusan dan override.
Penelitian ini	DSS design: menghubungkan AI/data, interaksi, keputusan manusia, dan feedback.	Artefak konseptual; perlu validasi ahli dan pengguna.

Tabel 1 menunjukkan bahwa penelitian terdahulu menyediakan landasan penting, tetapi umumnya berhenti pada satu tingkat analisis: fondasi DSS, perilaku kepercayaan, prinsip *human-centeredness*, atau *explainability model*. *Framework* ini memosisikan kontribusinya pada tingkat integrasi desain: rekomendasi AI harus disertai penjelasan dan informasi ketidakpastian; manusia memiliki kewenangan untuk menerima, merevisi, atau mengesampingkan rekomendasi; dan keputusan serta umpan balik dicatat untuk evaluasi dan kalibrasi berikutnya.

1.2. Kontribusi Penelitian

Penelitian ini memberikan kontribusi teoretis dan praktis bagi pengembangan DSS berbasis AI. Secara teoretis, penelitian menghasilkan artefak konseptual berupa *Human-AI Collaboration Framework* yang mengintegrasikan DSS, *Human-Centered AI*, XAI, dan *Human-AI Collaboration* ke dalam arsitektur keputusan yang terstruktur. Integrasi tersebut memandang pengambilan keputusan hibrida bukan sebagai sekadar pembagian tugas, melainkan sebagai pengaturan relasi antara kemampuan analitis, pemahaman konteks, dan akuntabilitas keputusan.

Dari perspektif *Design Science Research* (DSR), kontribusi penelitian ini terletak pada penerjemahan pengetahuan dari literatur menjadi kebutuhan desain, komponen artefak, dan kriteria evaluasi analitis. *Framework* tidak hanya menjelaskan peran AI dalam menghasilkan rekomendasi berbasis data, tetapi juga menjelaskan bagaimana pengguna memahami rekomendasi, menilai batasnya, dan memberi umpan balik yang dapat ditelusuri.

Secara praktis, *framework* menyediakan blueprint awal bagi organisasi dan pengembang sistem untuk merancang DSS yang lebih transparan, adaptif, dan berpusat pada manusia. *Blueprint* ini tidak menetapkan algoritma tunggal atau domain tertentu; sebaliknya, *framework* menawarkan struktur yang dapat disesuaikan dengan risiko keputusan, ketersediaan data, dan kebutuhan tata kelola organisasi.

Kontribusi penelitian dapat dirangkum sebagai berikut:

1. Merumuskan kebutuhan desain DSS kolaboratif manusia-AI yang mencakup transparansi, kontrol manusia, kontestabilitas, umpan balik, dan akuntabilitas.
2. Mengembangkan *Human-AI Collaboration Framework* yang menghubungkan *AI and Data Layer*, *Interaction Layer*, dan *Human Layer*.
3. Menempatkan XAI, komunikasi ketidakpastian, *trust calibration*, alasan *override*, dan pencatatan keputusan sebagai mekanisme desain yang saling terkait.
4. Menyediakan artefak konseptual berbasis DSR sebagai dasar pengembangan prototipe serta penelitian empiris pada berbagai konteks organisasi.

Berdasarkan kesenjangan tersebut, penelitian ini bertujuan merancang *Human-AI Collaboration Framework* untuk DSS menggunakan pendekatan DSR. *Framework* dirancang untuk menggabungkan kemampuan analitis AI dengan penalaran kontekstual, pertimbangan etis, dan tanggung jawab keputusan manusia. Artefak kemudian didemonstrasikan melalui skenario keputusan organisasi dan dievaluasi secara analitis terhadap kebutuhan desain yang diturunkan dari sintesis literatur [17], [18].

1.3. Kebaruan Penelitian

Kebaruan penelitian ini perlu dipahami sebagai kebaruan pada tingkat integrasi dan operasionalisasi desain. Studi sebelumnya telah membahas komponen-komponen penting, seperti kolaborasi manusia-AI, kontrol manusia, dan XAI [10]-[12]. Penelitian ini tidak mengklaim bahwa komponen tersebut sepenuhnya baru; kontribusinya adalah menyatukan komponen tersebut menjadi konfigurasi DSS yang secara eksplisit mendefinisikan alur rekomendasi, keputusan, *override*, umpan balik, dan akuntabilitas.

Framework yang diusulkan menghubungkan tiga lapisan: *AI and Data Layer* untuk menghasilkan rekomendasi dan sinyal risiko; *Interaction Layer* untuk menyajikan alternatif, penjelasan, ketidakpastian, dan ruang contestability; serta *Human Layer* untuk penilaian domain, evaluasi etis dan strategis, serta keputusan akhir. Konfigurasi ini menempatkan manusia sebagai pemilik otoritas keputusan sekaligus menyediakan jejak alasan ketika rekomendasi diterima, direvisi, atau diabaikan.

Selain integrasi lapisan, penelitian ini memposisikan *trust calibration* dan *feedback loop* sebagai mekanisme tata kelola lintas lapisan. *Trust calibration* diarahkan untuk membentuk *appropriate reliance*, bukan sekadar meningkatkan penerimaan pengguna terhadap AI. Umpan balik juga tidak diperlakukan sebagai pembaruan otomatis model, melainkan sebagai masukan yang dicatat, ditinjau, dan digunakan secara terkontrol untuk pemantauan serta kalibrasi sistem.

Dari perspektif metodologis, DSR digunakan untuk menghasilkan artefak konseptual yang dapat dikembangkan lebih lanjut menjadi prototipe. Dengan demikian, penelitian tidak hanya mengidentifikasi faktor-faktor yang memengaruhi *Human-AI Collaboration*, tetapi juga menawarkan rancangan yang dapat diuji melalui *expert review*, *usability evaluation*, maupun studi keputusan terkontrol pada tahap berikutnya.

Secara ringkas, posisi kebaruan penelitian dijelaskan sebagai berikut:

1. Mengintegrasikan DSS, *Human-Centered AI*, XAI, dan *Human-AI Collaboration* dalam satu konfigurasi artefak keputusan.
2. Menetapkan batas kewenangan manusia melalui mekanisme *accept*, *revise*, dan *override* atas rekomendasi AI.
3. Menghubungkan penjelasan, komunikasi ketidakpastian, *trust calibration*, dan alasan *override* ke dalam *Interaction Layer*.
4. Menggunakan DSR untuk menurunkan kebutuhan desain, membangun artefak, mendemonstrasikan skenario, dan melakukan evaluasi analitis.

Tabel 2. Posisi kebaruan *framework* yang diusulkan

Dimensi desain	Penekanan penelitian terdahulu	Kontribusi <i>framework</i>
Otoritas keputusan	Sering implisit atau bergantung pada domain.	<i>Accept</i> , <i>revise</i> , dan <i>override</i> oleh manusia dinyatakan eksplisit.
Explainability	Banyak berfokus pada interpretasi model.	Penjelasan dibuat <i>decision-facing</i> dan disertai ketidakpastian.
Trust calibration	Sering diperlakukan sebagai sikap pengguna.	<i>Appropriate reliance</i> didukung oleh bukti, batas, dan contestability.
Feedback dan audit	Umpan balik sering bersifat umum.	<i>Decision log</i> , alasan <i>override</i> , <i>review</i> , dan kalibrasi terintegrasi.
Artefak DSR	Konsep atau prototipe tertentu.	Kebutuhan desain, <i>framework</i> , demonstrasi, dan evaluasi analitis dihubungkan.
Pendekatan DSR	Tidak digunakan	Menjadi metode utama pengembangan artefak

2. METODE

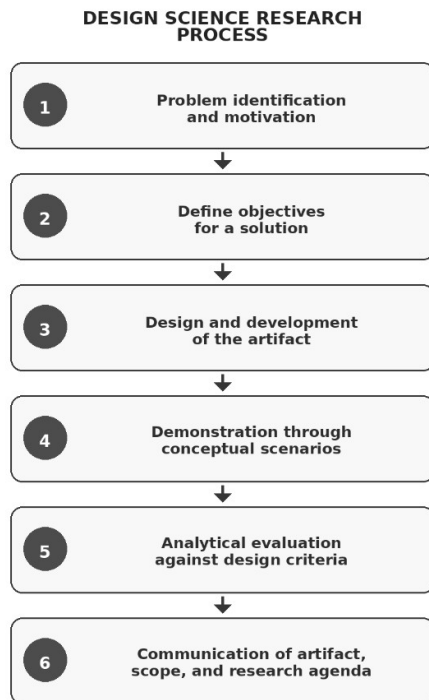
2.1. Pendekatan Penelitian

Penelitian ini menggunakan *Design Science Research* (DSR) untuk menghasilkan artefak berupa *Human-AI Collaboration Framework* bagi *Decision Support Systems* (DSS). DSR dipilih karena fokus penelitian adalah merancang solusi konseptual terhadap masalah desain, bukan menguji hubungan kausal antarvariabel. Dalam DSR, pengetahuan dari literatur dan konteks masalah diterjemahkan menjadi tujuan desain, artefak, dan evaluasi yang dapat dipertanggungjawabkan [17], [18]. Pada penelitian ini, sintesis literatur berfungsi sebagai design knowledge base untuk membangun artefak; sintesis tersebut tidak

diposisikan sebagai meta-analisis atau pengujian efektivitas kuantitatif.

2.2. Tahapan Penelitian

Penelitian mengadopsi enam tahap *Design Science Research Process* (DSRP) dari Peffers et al. [18]: *problem identification and motivation, define objectives for a solution, design and development, demonstration, evaluation, dan communication*. Literatur ditelaah pada tahap identifikasi masalah dan perumusan tujuan, sedangkan kontribusi serta keterbatasan artefak dikomunikasikan pada tahap akhir. Gambar 1 menyajikan alur enam tahap tersebut secara konsisten dengan proses penelitian.



Adapted from Peffers et al. [28].

Gambar 1. Tahapan *Design Science Research Process* (adaptasi dari [18])

2.2.1. Problem Identification and Motivation

Tahap pertama mengidentifikasi masalah desain pada DSS berbasis AI, yaitu *opacity*, bias algoritmik, keterbatasan pemahaman konteks, serta risiko *overreliance* dan *algorithm aversion*. Masalah tersebut ditelaah bersama keterbatasan keputusan manusia, seperti bias kognitif dan keterbatasan kapasitas pemrosesan informasi [4]-[9], [15], [16]. Hasil tahap ini adalah pernyataan masalah: DSS perlu memberi rekomendasi berbasis AI tanpa menghilangkan kemampuan pengguna untuk memahami, menilai, dan mempertanggungjawabkan keputusan.

2.2.2. Define Objectives for a Solution

Tujuan solusi diturunkan menjadi lima kebutuhan desain: (1) AI harus menghasilkan rekomendasi yang dapat ditelusuri ke data dan asumsi relevan; (2) antarmuka harus menyajikan penjelasan, alternatif, dan sinyal ketidakpastian; (3) manusia harus memiliki otoritas untuk menerima, merevisi, atau mengesampingkan rekomendasi; (4) sistem harus mencatat alasan dan umpan balik keputusan; serta (5) mekanisme pemantauan harus mendukung

kalibrasi kepercayaan dan pembaruan yang terkendali. Kebutuhan tersebut selaras dengan prinsip *Human-Centered AI* [11].

2.2.3. Design and Development

Tahap *design and development* menerjemahkan kebutuhan desain menjadi artefak konseptual yang terdiri atas *AI and Data Layer*, *Interaction Layer*, *Human Layer*, serta mekanisme tata kelola dan umpan balik lintas lapisan. Sintesis teori DSS, *Human-AI Collaboration*, *XAI*, *Human-Centered AI*, dan *trust in automation* digunakan untuk menentukan fungsi setiap komponen dan relasi antarkomponen [10]-[12], [15].

2.2.4. Demonstration

Artefak didemonstrasikan melalui *walkthrough* skenario seleksi vendor, prioritas investasi proyek, dan evaluasi risiko bisnis. Demonstrasi tidak dimaksudkan sebagai uji efektivitas empiris; tujuannya adalah memperlihatkan bagaimana rekomendasi AI, penjelasan, penilaian manusia, keputusan, dan umpan balik dapat mengalir secara konsisten dalam *framework*.

2.2.5. Evaluation

Evaluasi dilakukan secara analitis dan berbasis literatur dengan memeriksa keterkaitan antara komponen artefak dan kebutuhan desain. Kriteria evaluasi meliputi relevansi, kegunaan, konsistensi konseptual, transparansi, kontrol manusia, dan adaptabilitas [10]-[12], [14], [17]. Evaluasi ini menilai cakupan dan koherensi desain, bukan membuktikan peningkatan akurasi, kepercayaan, atau kualitas keputusan secara kuantitatif.

2.2.6. Communication

Tahap *communication* mendokumentasikan permasalahan, kebutuhan desain, artefak, demonstrasi, hasil evaluasi analitis, keterbatasan, dan agenda pengujian empiris melalui artikel ilmiah ini.

2.3. Pengumpulan Data

Data penelitian berupa literatur ilmiah yang digunakan sebagai basis pengetahuan desain. Pengumpulan dilakukan melalui sintesis literatur terstruktur dan didokumentasikan dengan alur PRISMA 2020 [19]. Prosedur tersebut digunakan untuk menunjukkan transparansi identifikasi dan seleksi sumber, sedangkan tujuan akhir sintesis adalah membangun *design knowledge base* bagi artefak DSR, bukan menghitung ukuran efek atau menghasilkan generalisasi statistik.

Kriteria inklusi mencakup artikel jurnal atau prosiding ilmiah yang membahas DSS, AI atau *machine learning*, *Human-AI Collaboration*, *Human-Centered AI*, *XAI*, *trust in automation*, atau *hybrid decision-making* serta memiliki relevansi langsung terhadap desain pengambilan keputusan. Dokumen yang tidak relevan, bersifat nonilmiah, duplikat, atau hanya membahas AI tanpa kaitan dengan proses keputusan dan peran manusia dikeluarkan dari analisis.

2.4. Proses Seleksi Literatur

Proses seleksi mengikuti empat fase PRISMA 2020: identifikasi, *screening*, penilaian kelayakan, dan inklusi. Setiap artikel disaring berdasarkan judul dan abstrak, kemudian artikel yang lolos diperiksa melalui *full-text review* untuk memastikan

keterkaitan dengan desain DSS, kolaborasi manusia-AI, atau transparansi dan tata kelola rekomendasi AI.

Pencarian dilakukan pada *Scopus*, *IEEE Xplore*, *ACM Digital Library*, *ScienceDirect*, dan *Web of Science* dengan blok konsep yang menggabungkan istilah DSS, AI, Human-AI, Human-Centered AI, XAI, trust, dan hybrid decision-making. Logika pencarian umum adalah: ("decision support system" OR DSS) AND ("artificial intelligence" OR "machine learning") AND ("human-AI" OR "human in the loop" OR "human-centered AI" OR explainab* OR "trust in automation"). Sintaks spesifik disesuaikan dengan fitur masing-masing basis data. Untuk reproduksibilitas penuh, tanggal pencarian, batas bahasa/periode, dan string final per basis data perlu diarsipkan bersama naskah pendukung.

Artikel dikeluarkan apabila tidak berhubungan langsung dengan pengambilan keputusan, tidak menjelaskan keterlibatan manusia atau transparansi AI, bersifat duplikat, atau tidak memenuhi kriteria kelayakan pada full-text review. Seleksi menghasilkan 35 studi yang digunakan dalam sintesis tematik. Alasan eksklusi dicatat pada tingkat screening dan full-text agar alur seleksi dapat ditelusuri.

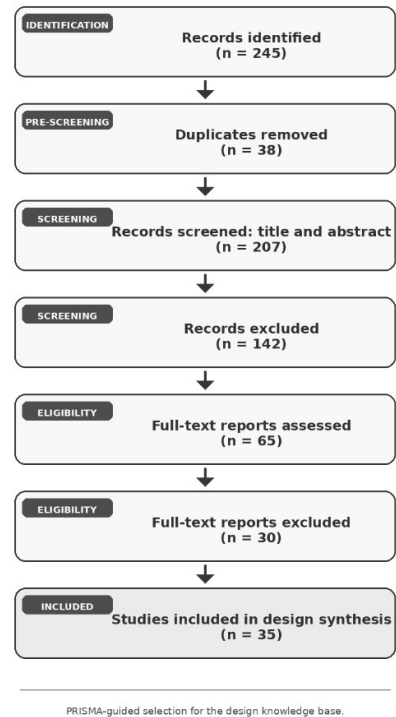
Ringkasan proses seleksi literatur ditunjukkan pada Tabel 3 dan Gambar 2.

Tabel 3. Ringkasan Proses Seleksi Literatur Berdasarkan PRISMA 2020

Tahapan	Jumlah Artikel
Artikel yang teridentifikasi dari basis data	245
Artikel duplikat yang dihapus	38
Artikel setelah penghapusan duplikat	207
Artikel yang disaring berdasarkan judul dan abstrak	207
Artikel yang dikeluarkan pada tahap screening	142
Artikel yang dinilai melalui full-text review	65
Artikel yang tidak memenuhi kriteria inklusi	30
Studi diinklusi dalam sintesis desain	35
Artikel yang digunakan dalam sintesis literatur	35

Sebanyak 35 artikel memenuhi kriteria inklusi dan digunakan sebagai sumber utama sintesis. Studi-studi tersebut berasal dari jurnal internasional, prosiding konferensi, dan publikasi akademik yang relevan dengan DSS, *Human-AI Collaboration*, *Human-Centered AI*, *XAI*, serta *trust in automation*. Penelitian ini tidak melakukan meta-analisis; temuan literatur dikodekan untuk mengidentifikasi kebutuhan desain, komponen, dan mekanisme evaluasi bagi *framework* yang diusulkan.

PRISMA 2020 SELECTION FLOW



Gambar 2. Diagram PRISMA 2020 Proses Seleksi Literatur

2.5. Analisis Data

Analisis dilakukan secara kualitatif-konseptual melalui sintesis tematik. Tahap analisis mencakup: identifikasi konsep dan klaim desain; pengelompokan konsep ke dalam tema AI/data, interaksi, dan peran manusia; pemetaan hubungan antartema; serta penerjemahan tema menjadi kebutuhan desain dan komponen artefak. Proses ini menjaga keterlacakan antara bukti literatur, kebutuhan desain, dan elemen *framework*.

2.6. Desain Artefak

Artefak utama penelitian adalah *Human-AI Collaboration Framework* untuk DSS. Artefak mencakup tiga lapisan utama dan mekanisme lintas lapisan yang mengatur penjelasan, komunikasi ketidakpastian, *trust calibration*, alasan *override*, umpan balik, audit trail, serta pemantauan dan kalibrasi model.

AI and Data Layer bertanggung jawab terhadap kualitas dan provenance data, analisis pola, prediksi, rekomendasi, serta pemantauan model. *Interaction Layer* menjadi ruang bagi pengguna untuk memahami alternatif, penjelasan, tingkat ketidakpastian, dan alasan rekomendasi. Human Layer mengelola evaluasi domain, pertimbangan konteks, kaji etis dan strategis, serta keputusan akhir.

Framework menempatkan AI sebagai *decision support tool*, bukan pemegang kewenangan keputusan. Lima kebutuhan desain yang dirumuskan pada tahap DSR diterapkan melalui: traceability data dan model; penjelasan serta komunikasi ketidakpastian; *accept-revise-override*; pencatatan keputusan dan umpan balik; serta pemantauan dan kalibrasi yang terkendali. Dengan rancangan tersebut, *framework* diarahkan untuk

mendukung keputusan yang lebih transparan, adaptif, dan berpusat pada manusia.

3. HASIL

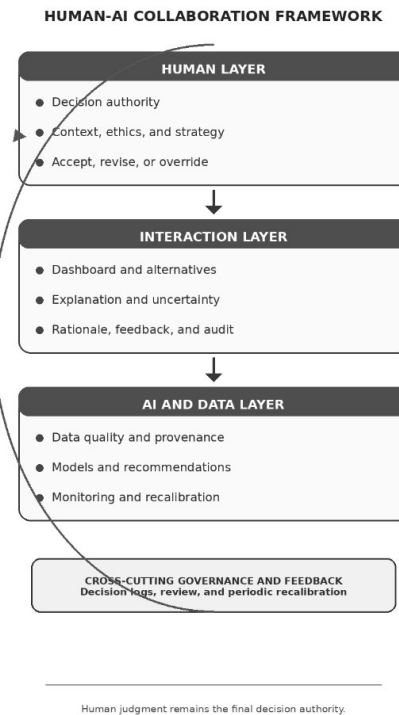
3.1. Framework Human-AI Collaboration

Hasil utama penelitian adalah artefak konseptual berupa *Human-AI Collaboration Framework* untuk DSS. Artefak ini dibangun dari sintesis teori DSS, *Human-Centered AI*, *XAI*, *trust in automation*, dan *Human-AI Collaboration* [10]-[12], [15]. Hasil tersebut bukan prototipe perangkat lunak atau bukti peningkatan performa; framework berfungsi sebagai blueprint yang menjelaskan komponen dan hubungan desain yang perlu diperhatikan ketika AI digunakan untuk mendukung keputusan.

Framework bertujuan menggabungkan kemampuan AI dalam mengolah data dan menghasilkan rekomendasi dengan kemampuan manusia dalam memahami konteks, menguji konsekuensi, dan menimbang aspek etis serta strategis. Karena setiap domain memiliki data, risiko, dan aturan yang berbeda, *framework* tidak mengharuskan penggunaan algoritma tertentu. Sebaliknya, *framework* mengatur peran dan mekanisme yang harus ada agar penggunaan AI tetap dapat dipahami dan dipertanggungjawabkan.

Framework terdiri atas *AI and Data Layer*, *Interaction Layer*, dan *Human Layer*. Ketiga lapisan dihubungkan oleh alur rekomendasi dari data dan model menuju antarmuka dan penilaian manusia, serta alur umpan balik dari keputusan manusia menuju pencatatan, pemantauan, dan kalibrasi. Tata kelola, *auditability*, dan pembelajaran terkendali bersifat lintas lapisan. Gambar 3 memvisualisasikan *Human-AI Collaboration Framework* yang dihasilkan dalam penelitian ini.

Untuk menegaskan human decision authority, Gambar 3 menempatkan Human Layer pada posisi teratas. Secara operasional, aliran analitis bergerak dari AI and Data Layer ke Interaction Layer, lalu ke Human Layer. Sebaliknya, hasil keputusan, alasan override, dan umpan balik bergerak kembali sebagai masukan tata kelola, pemantauan, dan kalibrasi. Dengan demikian, framework tidak memosisikan manusia sebagai tahap verifikasi pasif, melainkan sebagai pemilik keputusan akhir.



Gambar 3. *Human-AI Collaboration Framework for Decision Support Systems*

3.2. Sintesis Literatur

Sintesis literatur dilakukan sebelum perancangan *framework* untuk mengidentifikasi konsep-konsep yang dapat diterjemahkan menjadi kebutuhan desain. Literatur mencakup DSS, *Human-AI Collaboration*, *Human-Centered AI*, *XAI*, *trust in automation*, dan *hybrid decision-making*. Fokus sintesis bukan merangkum seluruh perkembangan bidang, melainkan memilih pengetahuan yang relevan untuk membangun artefak dan mengevaluasi koherensinya.

Hasil sintesis menunjukkan bahwa DSS modern memerlukan lebih dari kemampuan prediksi. Pengguna memerlukan rekomendasi yang dapat dijelaskan, informasi ketidakpastian, pembagian peran yang jelas, serta kemampuan untuk mempertanyakan dan mengubah keluaran AI. Transparansi dan kepercayaan yang proporsional muncul sebagai faktor penting karena keduanya memengaruhi apakah rekomendasi digunakan secara tepat, bukan sekadar seberapa sering rekomendasi diterima. Ringkasan hasil sintesis disajikan pada Tabel 4.

Berdasarkan hasil sintesis pada Tabel 4, terdapat lima tema dasar *framework*: kemampuan analitis dan kualitas data; keterlibatan serta kewenangan manusia; transparansi dan komunikasi ketidakpastian; *trust calibration* serta *contestability*; dan umpan balik yang dapat diaudit. Kelima tema tersebut dipetakan ke dalam *AI and Data Layer*, *Interaction Layer*, *Human Layer*, serta mekanisme tata kelola lintas lapisan..

Tabel 4. Sintesis Literatur sebagai Dasar Pengembangan Framework

Referensi	Fokus	Temuan relevan	Implikasi desain
Power [1]	<i>DSS</i>	DSS mendukung pertimbangan manajerial.	Konteks dasar keputusan.
Dietvorst et al. [13]	<i>Trust</i>	Kesalahan AI dapat mengurangi reliance. Kolaborasi membutuhkan koordinasi peran.	Trust calibration dan evaluasi diperlukan.
Cai et al. [10]	<i>Collaboration</i>	Manusia perlu mempertahankan kontrol.	Dasar integrasi manusia dan AI.
Shneiderman [11]	<i>Human-centered</i>	Interaksi manusia-AI perlu dirancang.	Dasar Human Layer dan decision authority.
Yang et al. [12]	<i>Interaction</i>	Penjelasan meningkatkan interpretabilitas.	Dasar Interaction Layer.
Arrieta et al. [17]	<i>XAI</i>	Trust yang tepat menentukan reliance.	Dasar explanation and uncertainty module.
Lee and See [21]	<i>Trust</i>	Kinerja dipengaruhi tugas, pengguna, dan antarmuka.	Dasar trust calibration.
Lai et al. [14]	<i>Decision-making</i>	Cognitive forcing dapat mengurangi overreliance.	Dasar demonstrasi dan agenda evaluasi.
Buçinca et al. [16]	<i>Overreliance</i>		Dasar contestability dan override.

3.3. Arsitektur Framework

Arsitektur *framework* terdiri atas tiga lapisan yang saling melengkapi dan mekanisme tata kelola lintas lapisan. Pemisahan lapisan digunakan untuk memperjelas akuntabilitas: siapa yang menghasilkan rekomendasi, bagaimana rekomendasi dijelaskan dan ditantang, serta siapa yang menetapkan keputusan.

3.3.1. AI Layer

AI and Data Layer bertanggung jawab terhadap pengelolaan data, analisis, dan penyediaan rekomendasi. Lapisan ini tidak diposisikan sebagai satu model tunggal, melainkan sebagai kumpulan fungsi yang dapat diimplementasikan dengan teknik analitis yang sesuai dengan konteks dan tingkat risiko keputusan.

Komponen utama *AI and Data Layer* meliputi:

1. *Data Provenance and Quality Checks*
2. *Analytical and Modeling Engine*
3. *Pattern and Uncertainty Analysis*
4. *Recommendation and Alternatives Engine*
5. *Performance Monitoring and Recalibration*

Data *provenance and quality checks* memastikan sumber, kelengkapan, dan kualitas data dapat ditelusuri sebelum analisis dilakukan. *Analytical and modeling engine* menjalankan model prediktif atau analitis sesuai kebutuhan domain. *Pattern and uncertainty analysis* mengidentifikasi pola, batas keyakinan, serta kondisi ketika rekomendasi perlu ditinjau secara lebih hati-hati. *Recommendation and alternatives engine* menyajikan rekomendasi berikut alternatif yang relevan, sedangkan

performance monitoring and recalibration memantau perubahan data, kinerja model, dan kebutuhan pembaruan secara terkendali.

Output *AI and Data Layer* bukan hanya peringkat atau rekomendasi tunggal, tetapi paket informasi keputusan yang mencakup rekomendasi, alternatif, faktor pendukung, provenance data, dan sinyal ketidakpastian. Paket ini menjadi masukan bagi *Interaction Layer* dan bukan pengganti keputusan manusia.

3.3.2. Interaction Layer

Interaction Layer berfungsi sebagai ruang komunikasi dan kontestasi antara AI dan pengguna. Lapisan ini bertugas mengubah keluaran teknis menjadi informasi yang dapat dipahami dan digunakan dalam pengambilan keputusan, sekaligus menyediakan sarana bagi pengguna untuk meminta klarifikasi, membandingkan alternatif, dan merekam alasan keputusan.

Komponen utama *Interaction Layer* terdiri atas:

1. *Decision Dashboard and Alternatives*
2. *Explanation and Uncertainty Module*
3. *Trust Calibration and Contestability Mechanism*
4. *Feedback and Override-Rationale Mechanism*
5. *Decision Log and Governance Component*

Decision dashboard menampilkan hasil analisis, alternatif, dan konsekuensi utama yang relevan dengan tugas pengguna. *Explanation and uncertainty module* menyajikan faktor yang memengaruhi rekomendasi, asumsi penting, serta tingkat keyakinan atau batas penerapan model. Dengan demikian, XAI tidak hanya berfungsi sebagai visualisasi teknis, tetapi sebagai dukungan bagi pengguna untuk memahami dasar keputusan AI [6], [7], [12].

Trust calibration diarahkan untuk membangun appropriate reliance: pengguna tidak diminta selalu mempercayai atau selalu menolak AI, tetapi menyesuaikan tingkat penggunaan rekomendasi dengan bukti, ketidakpastian, dan pengalaman kinerja sistem [15], [16]. Ketika pengguna menerima, merevisi, atau mengesampingkan rekomendasi, mekanisme umpan balik mencatat alasan dan konteks keputusan. Catatan tersebut mendukung audit, pembelajaran organisasi, serta pembaruan model yang ditinjau secara terkontrol.

3.3.3. Human Layer

Human Layer berperan dalam evaluasi, interpretasi, dan keputusan akhir. Pada lapisan ini, manusia tidak sekadar memeriksa keluaran AI, melainkan mengintegrasikan rekomendasi dengan pengetahuan domain, konteks organisasi, risiko, nilai, dan konsekuensi keputusan.

Komponen utama *Human Layer* meliputi:

1. *Domain Expertise*
2. *Contextual Judgment*
3. *Ethical and Regulatory Review*
4. *Strategic Alignment*
5. *Accept, Revise, or Override Decision*

Domain expertise memungkinkan pengguna menilai kewajaran rekomendasi berdasarkan pengalaman dan pengetahuan substantif. *Contextual judgment* menambahkan faktor eksternal

atau situasional yang tidak seluruhnya direpresentasikan oleh data. *Ethical and regulatory review* memastikan bahwa keputusan tidak hanya efisien, tetapi juga selaras dengan nilai, aturan, dan batas risiko organisasi. *Strategic alignment* menilai kesesuaian keputusan dengan tujuan jangka panjang. Berdasarkan evaluasi tersebut, pengguna berwenang menerima, merevisi, atau mengesampingkan rekomendasi AI.

Output Human Layer adalah keputusan akhir beserta alasan, tingkat keyakinan, dan tindakan tindak lanjut yang relevan. Informasi ini menjadi bagian dari *decision log* dan dapat digunakan untuk mengevaluasi kualitas keputusan serta memperbaiki sistem pada penggunaan berikutnya.

3.3.4. Alur Pengambilan Keputusan

Proses pengambilan keputusan dalam *framework* berlangsung melalui enam tahap berikut:

1. Data dikumpulkan dari sumber yang terdokumentasi dan diperiksa kualitas serta *provenance*-nya.
2. AI menganalisis data, menyusun alternatif, dan menghasilkan rekomendasi beserta sinyal ketidakpastian.
3. *Interaction Layer* menyajikan penjelasan, alternatif, asumsi, dan informasi yang diperlukan untuk menilai rekomendasi.
4. Pengguna mengevaluasi rekomendasi berdasarkan keahlian domain, konteks organisasi, risiko, dan konsekuensi etis.
5. Pengguna menerima, merevisi, atau mengesampingkan rekomendasi dan menetapkan keputusan akhir.
6. Keputusan, alasan, dan hasil tindak lanjut dicatat sebagai umpan balik untuk audit, pemantauan, dan kalibrasi terkendali.

Alur tersebut membentuk kolaborasi berkelanjutan antara manusia dan AI. Penting untuk dicatat bahwa umpan balik tidak otomatis mengubah model; perubahan model harus melalui prosedur evaluasi dan tata kelola yang sesuai agar sistem tidak memperkuat bias atau keputusan yang tidak tepat.

3.4. Demonstrasi Konseptual

Demonstrasi konseptual digunakan untuk menunjukkan penerapan *framework* pada keputusan organisasi yang berbeda. Demonstrasi bersifat *scenario walkthrough*, bukan eksperimen atau studi lapangan. Dengan demikian, tujuan bagian ini adalah menguji keterpakaian alur artefak secara logis dan memperjelas titik interaksi manusia-AI.

3.4.1. Skenario Seleksi Vendor

Dalam seleksi vendor, AI menganalisis harga, kualitas layanan, performa historis, risiko, dan kepatuhan kontrak. Sistem menyajikan peringkat vendor, alternatif yang berdekatan, faktor yang memengaruhi penilaian, serta sinyal ketidakpastian apabila data tidak lengkap atau kondisi vendor berubah.

Manajer pengadaan menilai keluaran tersebut dengan mempertimbangkan hubungan bisnis jangka panjang, kondisi pasar, risiko reputasi, dan kebutuhan strategis organisasi. Manajer dapat menerima, merevisi bobot, atau mengesampingkan rekomendasi; alasan keputusan dicatat untuk memudahkan evaluasi dan pembelajaran berikutnya.

3.4.2. Skenario Prioritas Investasi Proyek

Pada skenario ini, AI membandingkan alternatif proyek berdasarkan *Return on Investment* (ROI), tingkat risiko, kebutuhan sumber daya, dan potensi manfaat bisnis. Sistem menyajikan prioritas, asumsi utama, sensitivitas terhadap perubahan parameter, dan alternatif keputusan yang layak dipertimbangkan.

Manajemen kemudian menilai rekomendasi berdasarkan visi organisasi, strategi bisnis, keterbatasan anggaran, kondisi eksternal, dan toleransi risiko. Keputusan investasi tidak ditetapkan otomatis oleh model; keputusan akhir beserta alasannya disimpan sebagai bagian dari proses tata kelola.

3.4.3. Skenario Evaluasi Risiko Bisnis

Dalam evaluasi risiko bisnis, AI mengidentifikasi pola risiko dari data historis dan indikator operasional. Sistem menyajikan tingkat risiko, indikator pemicu, strategi mitigasi yang mungkin, serta informasi mengenai batas keyakinan rekomendasi.

Pengguna menafsirkan hasil tersebut dengan mempertimbangkan risiko reputasi, perubahan regulasi, kondisi sosial ekonomi, dan dampak terhadap pemangku kepentingan. Manajemen memilih tindakan mitigasi, merekam alasan serta hasilnya, lalu menggunakan informasi itu untuk meninjau kebutuhan pemantauan atau kalibrasi model.

Ketiga skenario menunjukkan bagaimana *framework* menjaga alur yang sama pada konteks yang berbeda: data dan AI menghasilkan rekomendasi; interaksi menyajikan bukti, ketidakpastian, dan ruang kontestasi; manusia membuat keputusan akhir; dan hasilnya dicatat sebagai umpan balik. Demonstrasi ini mendukung keterpakaian konseptual *framework*, tetapi belum membuktikan keunggulan performa dibandingkan pendekatan lain.

4. PEMBAHASAN

4.1. Evaluasi Framework

Framework dievaluasi secara analitis untuk menilai apakah komponen artefak menanggapi masalah desain yang diidentifikasi. Evaluasi membandingkan fitur *framework* dengan prinsip *Human-Centered AI*, *XAI*, *trust in automation*, dan *Human-AI Collaboration* [10]-[12], [15]. Karena belum ada prototipe maupun data pengguna, evaluasi ini tidak digunakan untuk mengklaim peningkatan akurasi, kepercayaan, atau kualitas keputusan secara kausal.

Transparansi ditangani melalui *explanation and uncertainty module* pada *Interaction Layer*. Modul ini menyediakan alasan rekomendasi, faktor yang berpengaruh, asumsi, dan batas keyakinan sehingga pengguna dapat mengevaluasi apakah rekomendasi layak digunakan pada situasi tertentu. Pendekatan ini memperluas peran XAI dari sekadar interpretasi model menjadi dukungan untuk evaluasi keputusan [6], [7], [12].

Human control diwujudkan melalui kewenangan *accept*, *revise*, atau *override*. Penetapan keputusan akhir tetap berada pada pengguna yang memiliki keahlian dan tanggung jawab

organisasi. Pengaturan tersebut dapat mengurangi risiko automation bias, tetapi tidak menjamin bahwa keputusan manusia selalu benar; efektivitasnya tetap perlu diuji melalui studi empiris pada tugas dan pengguna nyata [11], [20].

Adaptabilitas *framework* berasal dari sifatnya yang independen terhadap satu jenis algoritma atau domain. Komponen yang sama dapat dipetakan ke sektor bisnis, pendidikan, kesehatan, pemerintahan, dan manufaktur, tetapi implementasinya harus menyesuaikan tingkat risiko, kualitas data, regulasi, dan siapa yang berwenang mengambil keputusan. Karena itu, adaptabilitas diposisikan sebagai potensi desain yang perlu divalidasi lintas domain, bukan sebagai klaim generalisasi yang sudah terbukti.

Mekanisme *feedback* dan *decision log* mendukung *continuous improvement* melalui pencatatan keputusan, hasil, dan alasan *override*. Namun, penggunaan umpan balik untuk pembaruan model harus disertai pemantauan kualitas data, evaluasi bias, dan tata kelola perubahan agar sistem tidak secara otomatis memperkuat keputusan masa lalu yang keliru. Prinsip ini mempertahankan hubungan antara pembelajaran sistem dan akuntabilitas organisasi [1], [2], [15].

4.2. Validasi Konseptual Framework

Validasi konseptual dilakukan melalui evaluasi artefak pada tahap awal DSR dengan mengacu pada relevansi, kegunaan, konsistensi konseptual, transparansi, kontrol manusia, dan adaptabilitas [17], [18]. Setiap kriteria diperiksa melalui bukti yang tampak dalam artefak dan bukan melalui pengukuran pengguna. Oleh sebab itu, status pada Tabel 5 perlu dibaca sebagai "ditangani secara konseptual", bukan sebagai bukti bahwa *framework* telah efektif dalam praktik.

Dari aspek relevansi, *framework* mengakomodasi *opacity*, keterbatasan konteks, bias algoritmik, serta risiko kepercayaan yang tidak proporsional melalui komponen penjelasan, komunikasi ketidakpastian, kontestabilitas, dan pencatatan keputusan [4]-[7], [15], [16]. Hubungan ini menunjukkan bahwa masalah desain dan fitur artefak memiliki keterkaitan yang eksplisit.

Dari aspek kegunaan, *framework* membagi peran sesuai keunggulan masing-masing: AI mendukung analisis data dan penyusunan alternatif, sedangkan manusia melakukan interpretasi, penilaian nilai, dan keputusan akhir. Pembagian ini konsisten dengan prinsip *Human-Centered AI*, tetapi kemudahan penggunaan, beban kognitif, dan efektivitas keputusan tetap memerlukan pengujian pengguna [11].

Dari aspek konsistensi, *AI and Data Layer* menghasilkan informasi keputusan; *Interaction Layer* menerjemahkan informasi tersebut menjadi rekomendasi yang dapat dipahami dan ditantang; *Human Layer* menetapkan keputusan; dan mekanisme lintas lapisan mencatat serta memantau hasil keputusan. Keterkaitan tersebut menyediakan rantai penalaran dari data hingga keputusan akhir dan kembali ke proses pembelajaran yang terkendali.

Tabel 5 menyajikan hasil evaluasi konseptual *framework* berdasarkan kriteria DSR.

Tabel 5. Hasil Validasi Konseptual Framework

Kriteria	Bukti dalam artefak	Status
Relevansi	Menangani <i>opacity</i> , bias, konteks, dan trust melalui penjelasan, <i>override</i> , dan log keputusan.	Ditangani secara konseptual
Kegunaan	Membagi peran AI untuk analisis dan manusia untuk interpretasi serta keputusan.	Ditangani secara konseptual
Konsistensi	Alur AI/data - interaksi - keputusan manusia - feedback terhubung. Penjelasan, alternatif, asumsi, dan ketidakpastian disajikan pada <i>Interaction Layer</i> .	Ditangani secara konseptual
Transparansi	Pengguna dapat menerima, merevisi, atau mengesampingkan rekomendasi.	Ditangani secara konseptual
Kontrol manusia	Komponen tidak bergantung pada satu algoritma atau domain.	Perlu uji lintas domain
Adaptabilitas		

Hasil evaluasi konseptual menunjukkan bahwa *framework* secara desain telah menangani prinsip-prinsip penting *Human-Centered AI*, *XAI*, *trust calibration*, dan *Human-AI Collaboration*. Namun, hasil ini bersifat argumentatif dan berbasis literatur. Validasi berikutnya perlu melibatkan ahli, pengguna potensial, prototipe, dan pengukuran keputusan untuk menilai apakah mekanisme yang diusulkan benar-benar meningkatkan pemahaman, *appropriate reliance*, dan kualitas keputusan.

4.3. Perbandingan dengan Penelitian Sebelumnya

Penelitian terdahulu membahas *Human-AI Collaboration* dari berbagai perspektif, antara lain perilaku pengguna, desain interaksi, transparansi model, dan kepercayaan terhadap otomatisasi. Perbandingan berikut digunakan untuk menjelaskan posisi desain *framework*, bukan untuk menyatakan bahwa *framework* secara empiris lebih unggul dibandingkan setiap penelitian sebelumnya [10], [12], [14].

Cai et al. [10] menunjukkan pentingnya kolaborasi manusia dan AI pada konteks pengambilan keputusan yang memerlukan kombinasi kemampuan analitis dan pengalaman praktis. Kontribusi tersebut menegaskan pentingnya peran pengguna dan proses *onboarding*, tetapi tidak secara khusus menyusun blueprint DSS yang memetakan alur rekomendasi, penjelasan, keputusan, dan umpan balik lintas domain.

Shneiderman [11] menempatkan kontrol manusia dan keandalan sistem sebagai prinsip *Human-Centered AI*. *Framework* ini menerjemahkan prinsip tersebut ke dalam desain DSS melalui *Human Layer* serta mekanisme *accept*, *revise*, dan *override*. Dengan demikian, prinsip abstrak mengenai kontrol manusia dihubungkan dengan titik keputusan dan jejak alasan yang lebih operasional.

Arrieta et al. [12] menekankan pentingnya XAI untuk meningkatkan interpretabilitas dan akuntabilitas sistem AI.

Framework mengadopsi konsep tersebut dalam *Interaction Layer*, tetapi memperluasnya dengan komunikasi ketidakpastian, alternatif keputusan, kontestabilitas, dan pencatatan alasan pengguna. Perluasan ini penting karena pengguna DSS memerlukan dukungan untuk mengambil keputusan, bukan hanya penjelasan tentang model.

Dietvorst et al. [16] memperlihatkan *algorithm aversion* setelah pengguna menyaksikan kesalahan sistem. Bersama literatur *trust in automation* [15], temuan tersebut mendukung penggunaan *trust calibration* yang mendorong evaluasi proporsional, bukan penerimaan atau penolakan ekstrem. *Framework* menempatkan mekanisme ini pada *Interaction Layer* agar pengguna dapat memeriksa bukti dan menggunakan hak *override* ketika diperlukan.

Secara keseluruhan, *framework* menggabungkan kontribusi terdahulu menjadi arsitektur keputusan yang terhubung: *AI and Data Layer* menghasilkan rekomendasi; *Interaction Layer* mendukung pemahaman dan kontestasi; *Human Layer* memegang keputusan akhir; serta mekanisme tata kelola menyimpan jejak dan memantau perubahan. Nilai tambahnya terletak pada konfigurasi hubungan tersebut dan bukan pada klaim bahwa setiap komponennya merupakan konsep baru.

Tabel 6 membandingkan fokus desain *framework* dengan penekanan umum pada penelitian terdahulu.

Tabel 6. Perbandingan fokus desain dengan penelitian terdahulu

Aspek	Penekanan umum penelitian terdahulu	Framework yang diusulkan
Fokus AI	Performa model atau prediksi.	Peran AI dihubungkan dengan peran dan kewenangan manusia.
Explainability	Penjelasan dibahas sebagai fitur model.	Penjelasan, alternatif, dan ketidakpastian mendukung keputusan.
Human control	Tidak selalu operasional dalam DSS.	Accept, revise, dan override menjadi titik keputusan eksplisit.
DSS hibrida	Sering dibahas pada kasus atau domain tertentu.	Blueprint konseptual untuk lintas domain.
Trust calibration	Dibahas sebagai sikap atau perilaku pengguna.	Bukti, batas sistem, dan contestability diposisikan dalam interaksi.
Feedback	Mekanisme pembelajaran umum.	Decision log, alasan override, review, dan kalibrasi terkontrol.

4.4. Implikasi Teoretis

Secara teoretis, penelitian ini memperluas diskusi *Human-AI Collaboration* dalam DSS dengan menggeser fokus dari keberadaan manusia dalam sistem ke pengaturan hubungan yang

dapat ditelusuri antara rekomendasi, penjelasan, kewenangan, dan umpan balik. *Framework* menyatukan dimensi teknis, interaksional, dan pengambilan keputusan dalam satu artefak konseptual [10], [14].

Penelitian juga menghubungkan *Human-Centered AI* dan XAI sebagai konsep yang saling melengkapi. *Human-Centered AI* mengatur siapa yang memiliki kontrol dan tanggung jawab keputusan, sedangkan XAI mendukung pengguna untuk memahami serta menilai dasar rekomendasi. Trust calibration dan contestability melengkapi keduanya dengan menjelaskan bagaimana pengguna seharusnya menggunakan rekomendasi secara proporsional [11], [12], [15].

Bagi kajian DSR dalam sistem informasi, penelitian ini menunjukkan bagaimana sintesis literatur dapat diterjemahkan menjadi kebutuhan desain, struktur artefak, demonstrasi skenario, dan evaluasi analitis. Artefak ini dapat menjadi titik awal untuk prototipe dan pengujian berikutnya, bukan pengganti validasi empiris [17], [18].

4.5. Implikasi Praktis

Secara praktis, *framework* dapat digunakan sebagai *blueprint* awal untuk mengembangkan DSS berbasis *Human-AI Collaboration* pada berbagai sektor organisasi. Penerapan sebaiknya dimulai dari keputusan yang memiliki manfaat analitis jelas, risiko yang dapat dikelola, dan pengambil keputusan yang dapat terlibat aktif dalam perancangan serta evaluasi sistem.

Bagi organisasi, *framework* dapat membantu merancang alur kerja yang memanfaatkan kemampuan analitis AI tanpa menghilangkan akuntabilitas manajerial. Organisasi perlu menetapkan siapa yang memiliki otoritas keputusan, kondisi ketika *override* dapat dilakukan, informasi apa yang harus ditampilkan kepada pengguna, serta bagaimana keputusan dan hasilnya dicatat.

Bagi pengembang sistem, *framework* menyediakan struktur integrasi antara kualitas data, model AI, antarmuka keputusan, penjelasan, dan log keputusan. Struktur ini dapat digunakan sebagai checklist desain sebelum memilih algoritma atau mengembangkan dashboard, sehingga kebutuhan pengguna dan tata kelola tidak tertinggal di belakang keputusan teknis.

Bagi peneliti, *framework* menyediakan proposisi desain yang dapat diuji melalui prototipe, studi laboratorium, simulasi keputusan, atau penelitian lapangan. Penelitian berikutnya dapat menguji apakah fitur seperti komunikasi ketidakpastian, alasan *override*, dan alternatif keputusan meningkatkan pemahaman serta appropriate reliance pengguna.

4.6. Implikasi Manajerial

Human-AI Collaboration Framework memiliki implikasi manajerial bagi organisasi yang akan mengadopsi AI dalam proses pengambilan keputusan. Nilai manajerialnya bukan hanya pada percepatan analisis, tetapi juga pada kemampuan menjaga akuntabilitas, pengawasan, dan kesesuaian keputusan dengan tujuan organisasi ketika rekomendasi AI mulai digunakan pada proses strategis, taktis, maupun operasional.

Dari perspektif manajemen, *framework* membantu organisasi mengintegrasikan AI ke dalam DSS tanpa menciptakan ketergantungan berlebihan terhadap otomatisasi. Dengan menempatkan manusia sebagai pemegang keputusan akhir, organisasi dapat menetapkan batas penggunaan AI, mekanisme eskalasi, serta persyaratan peninjauan manusia sesuai tingkat risiko keputusan.

Explanation and uncertainty module memiliki implikasi bagi *AI governance* karena membantu manajer memahami dasar rekomendasi, data yang digunakan, asumsi penting, dan batas penerapan model. Informasi tersebut mendukung audit, evaluasi keputusan, dan komunikasi lintas pemangku kepentingan, terutama pada sektor dengan tingkat regulasi atau dampak yang tinggi.

Trust calibration membantu organisasi membangun penggunaan AI yang proporsional. Tujuannya bukan membuat pengguna selalu menerima rekomendasi, melainkan membantu mereka mengenali kapan bukti AI kuat, kapan ketidakpastian tinggi, dan kapan pertimbangan manusia perlu lebih dominan. Pendekatan ini dapat mengurangi risiko automation bias maupun *algorithm aversion* [13], [21].

Feedback loop dan *decision log* mendukung *continuous improvement* dengan menyediakan informasi tentang keputusan, alasan *override*, dan hasil tindak lanjut. Agar bermanfaat, informasi tersebut perlu ditinjau secara periodik sebagai bagian dari tata kelola perubahan model, bukan langsung digunakan sebagai data pelatihan tanpa pemeriksaan kualitas dan bias.

Tabel 7 merangkum implikasi manajerial *framework* yang diusulkan.

Tabel 7. Implikasi Manajerial *Human-AI Collaboration Framework*

Aspek manajerial	Implikasi
Kualitas keputusan	Menggabungkan analisis AI dengan penilaian konteks, etika, dan strategi oleh manusia.
Transparansi dan audit	Menyediakan penjelasan, ketidakpastian, alasan keputusan, dan jejak audit.
Kepercayaan pengguna	Mendorong <i>appropriate reliance</i> , bukan penerimaan otomatis terhadap AI.
Pengendalian risiko	Memungkinkan eskalasi, <i>override</i> , dan peninjauan manusia sesuai tingkat risiko.
Efisiensi operasional	Mempercepat sintesis data dan penyajian alternatif tanpa mengotomatisasi keputusan akhir.
Continuous improvement	Memfaatkan log keputusan dan umpan balik untuk pemantauan serta kalibrasi terkontrol.
AI governance	Mendukung batas kewenangan, dokumentasi, evaluasi, dan akuntabilitas penggunaan AI.

Secara keseluruhan, keberhasilan implementasi AI dalam organisasi tidak hanya bergantung pada kinerja algoritma. Keberhasilan juga dipengaruhi oleh kualitas data, desain interaksi, kejelasan otoritas keputusan, kemampuan pengguna memahami batas sistem, serta tata kelola umpan balik.

<https://doi.org/10.25077/TEKNOSI.v12i2.2026.327-339>

Framework ini menawarkan struktur untuk menyeimbangkan faktor-faktor tersebut dalam rancangan DSS.

4.7. Keterbatasan Penelitian

Penelitian ini memiliki beberapa keterbatasan. Pertama, artefak masih berada pada tahap konseptual dan belum diimplementasikan sebagai sistem nyata maupun diuji melalui eksperimen. Karena itu, penelitian tidak dapat menyimpulkan bahwa *framework* meningkatkan akurasi, kualitas keputusan, kepercayaan, atau efisiensi secara kuantitatif.

Kedua, evaluasi berbasis literatur dan *PRISMA-guided selection* sangat dipengaruhi oleh cakupan basis data, strategi penelusuran, dan kualitas publikasi yang terpilih. Walaupun alur seleksi dan kriteria inklusi telah dijelaskan, reproduksibilitas penuh memerlukan arsip tanggal pencarian, string final tiap basis data, batas bahasa/periode, dan daftar artikel yang diekspor.

Ketiga, *framework* disengaja berada pada tingkat abstraksi yang tidak mengikat satu algoritma, dataset, atau domain tertentu. Konsekuensinya, implementasi teknis, jenis penjelasan, prosedur tata kelola, serta bentuk antarmuka harus disesuaikan dengan risiko keputusan dan konteks organisasi yang sebenarnya.

Keempat, penelitian belum melibatkan *expert review*, *Delphi study*, *focus group discussion* (FGD), atau pengujian pengguna. Penerimaan praktis, beban kognitif, kualitas penjelasan, dan dampak mekanisme *override* terhadap keputusan belum dapat dinilai dari artefak konseptual ini.

4.8. Agenda Penelitian Mendatang

Penelitian selanjutnya dapat memulai dengan *expert validation* menggunakan *Delphi study* atau FGD yang melibatkan akademisi, praktisi sistem informasi, pengembang AI, dan pengambil keputusan organisasi. Validasi tersebut dapat menguji kelengkapan komponen, keterpakaian, kemudahan implementasi, dan kesesuaian *framework* pada domain berisiko rendah hingga tinggi.

Tahap berikutnya adalah mengembangkan prototipe DSS berdasarkan *framework* dan melakukan evaluasi pengguna. Pengukuran dapat mencakup kualitas keputusan, pemahaman terhadap penjelasan, *appropriate reliance*, waktu keputusan, beban kognitif, *usability*, frekuensi serta alasan *override*, dan kualitas audit trail. Desain evaluasi perlu membandingkan kondisi tanpa penjelasan, dengan penjelasan, dan dengan mekanisme kontestabilitas untuk mengetahui kontribusi masing-masing fitur.

Penelitian lintas domain diperlukan untuk menguji adaptabilitas *framework* pada kesehatan, pendidikan, keuangan, pemerintahan, dan manufaktur. Setiap domain perlu menetapkan ukuran risiko, aturan akuntabilitas, jenis keputusan, serta standar kualitas data yang berbeda. Studi lintas domain akan membantu membedakan elemen *framework* yang bersifat umum dari elemen yang perlu disesuaikan.

Penelitian berikutnya juga dapat mengembangkan mekanisme adaptif yang menyesuaikan tingkat keterlibatan AI dan manusia berdasarkan tingkat risiko keputusan, kualitas data, dan ketidakpastian model. Dalam konteks ini, pengembangan XAI interaktif, evaluasi fairness, dan tata kelola pembaruan model

perlu dipertimbangkan secara terpadu agar kolaborasi manusia-AI tetap efektif dan bertanggung jawab.

5. KESIMPULAN

Penelitian ini merancang *Human-AI Collaboration Framework* untuk *Decision Support Systems* (DSS) menggunakan *Design Science Research* (DSR). Melalui identifikasi masalah, sintesis literatur sebagai design knowledge base, perumusan kebutuhan desain, pengembangan artefak, demonstrasi konseptual, dan evaluasi analitis, penelitian menghasilkan blueprint bagi pengambilan keputusan hibrida yang tetap menempatkan manusia sebagai pemegang kewenangan akhir.

Framework terdiri atas *AI and Data Layer*, *Interaction Layer*, dan *Human Layer*. *AI and Data Layer* menghasilkan rekomendasi, alternatif, *provenance data*, dan sinyal ketidakpastian. *Interaction Layer* menyajikan penjelasan, komunikasi ketidakpastian, *trust calibration*, *contestability*, alasan *override*, umpan balik, dan *decision log*. *Human Layer* mengintegrasikan keahlian domain, penilaian konteks, pertimbangan etis dan strategis, lalu menerima, merevisi, atau mengesampingkan rekomendasi AI.

Evaluasi konseptual menunjukkan bahwa *framework* secara desain telah menanggapi kebutuhan transparansi, kontrol manusia, akuntabilitas, dan pembelajaran terkendali. Namun, evaluasi ini tidak membuktikan efektivitas empiris. Klaim mengenai kualitas keputusan, tingkat kepercayaan, beban kognitif, atau performa sistem harus diuji melalui *expert validation*, prototipe, dan penelitian pengguna pada konteks keputusan nyata.

Kontribusi utama penelitian terletak pada konfigurasi artefak yang menghubungkan kemampuan analitis AI, interaksi keputusan yang dapat dijelaskan dan ditantang, serta kewenangan manusia dalam satu arsitektur DSS. *Framework* dapat digunakan sebagai landasan konseptual bagi organisasi dan pengembang sistem untuk merancang *AI-assisted decision making* yang lebih transparan, dapat diaudit, dan berpusat pada manusia.

Penelitian selanjutnya disarankan untuk memvalidasi *framework* melalui *Delphi study* atau FGD, membangun prototipe, dan menjalankan evaluasi empiris terhadap *appropriate reliance*, pemahaman penjelasan, *usability*, beban kognitif, kualitas keputusan, serta efektivitas mekanisme *override* dan umpan balik. Pengujian lintas domain diperlukan untuk menilai batas generalisasi dan kebutuhan adaptasi *framework* pada karakteristik risiko serta tata kelola yang berbeda.

DAFTAR PUSTAKA

- [1] D. J. Power, "A Brief History of Decision Support Systems," DSSResources.com, ver. 4.0, Mar. 10, 2007. [Online]. Available: <https://dssresources.com/history/dsshistory.html>. Accessed: Jun. 19, 2026.
- [2] E. Turban, C. Pollard, and G. Wood, Information Technology for Management: Driving Digital Transformation to Increase Local and Global Performance, Growth and Sustainability, 12th ed. Hoboken, NJ, USA: Wiley, 2021.
- [3] T. H. Davenport and J. G. Harris, Competing on Analytics, Updated, with a New Introduction: The New Science of Winning. Boston, MA, USA: Harvard Business Review Press, 2017.
- [4] N. Mehrabi, F. Morstatter, N. Saxena, K. Lerman, and A. Galstyan, "A survey on bias and fairness in machine learning," ACM Comput. Surv., vol. 54, no. 6, pp. 1–35, 2021.
- [5] S. Barocas, M. Hardt, and A. Narayanan, Fairness and Machine Learning: Limitations and Opportunities. Cambridge, MA, USA: MIT Press, 2023.
- [6] D. Gunning and D. Aha, "DARPA's Explainable Artificial Intelligence (XAI) program," AI Mag., vol. 40, no. 2, pp. 44–58, 2019.
- [7] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you?: Explaining the predictions of any classifier," in Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining, 2016, pp. 1135–1144.
- [8] D. Kahneman, O. Sibony, and C. R. Sunstein, Noise: A Flaw in Human Judgment. New York, NY, USA: Little, Brown Spark, 2021.
- [9] A. Tversky and D. Kahneman, "Judgment under uncertainty: Heuristics and biases," Science, vol. 185, no. 4157, pp. 1124–1131, 1974.
- [10] C. J. Cai et al., "Hello AI: Uncovering the onboarding needs of medical practitioners for human-AI collaborative decision-making," Proc. ACM Hum.-Comput. Interact., vol. 3, no. CSCW, pp. 1–24, 2019.
- [11] B. Shneiderman, Human-Centered AI. New York, NY, USA: Oxford University Press, 2022.
- [12] A. B. Arrieta et al., "Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI," Inf. Fusion, vol. 58, pp. 82–115, 2020.
- [13] Q. Yang, A. Steinfeld, C. Rosé, and J. Zimmerman, "Re-examining whether, why, and how human-AI interaction is uniquely difficult to design," in Proc. CHI Conf. Human Factors in Computing Systems, 2020.
- [14] V. Lai, C. Chen, Q. V. Liao, A. Smith-Renner, and C. Tan, "Towards a science of human-AI decision making: An overview of design space in empirical human-subject studies," in Proc. 2023 ACM Conf. Fairness, Accountability, and Transparency, 2023.
- [15] J. D. Lee and K. A. See, "Trust in automation: Designing for appropriate reliance," Hum. Factors, vol. 46, no. 1, pp. 50–80, 2004.
- [16] B. J. Dietvorst, J. P. Simmons, and C. Massey, "Algorithm aversion: People erroneously avoid algorithms after seeing them err," J. Exp. Psychol. Gen., vol. 144, no. 1, pp. 114–126, 2015.
- [17] A. R. Hevner, S. T. March, J. Park, and S. Ram, "Design science in information systems research," MIS Q., vol. 28, no. 1, pp. 75–105, 2004.
- [18] K. Peffers, T. Tuunanen, M. A. Rothenberger, and S. Chatterjee, "A design science research methodology for information systems research," J. Manag. Inf. Syst., vol. 24, no. 3, pp. 45–77, 2007.

- [19] M. J. Page et al., "The PRISMA 2020 statement: An updated guideline for reporting systematic reviews," *BMJ*, vol. 372, p. n71, 2021.
- [20] Z. Buçinca, M. B. Malaya, and K. Z. Gajos, "To trust or to think: Cognitive forcing functions can reduce overreliance on AI in AI-assisted decision making," *Proc. ACM Hum.-Comput. Interact.*, vol. 5, no. CSCW1, 2021.

BIODATA PENULIS



Penulis Pertama

Muhammad Damas Fatih, S.Kom., M.Kom. lahir di Jambi pada 5 Januari 1994 dan merupakan dosen tetap Program Studi Sistem Informasi, Fakultas Sains dan Teknologi, Universitas Jambi.

Penulis menyelesaikan pendidikan S-1 Teknik Informatika dan S-2 Informatika di Universitas Islam Indonesia. Bidang keilmuan yang ditekuni meliputi sistem informasi manajemen, pemodelan proses bisnis, perancangan sistem, interaksi manusia dan komputer, serta digital marketing. Penulis dapat dihubungi melalui email: muhammaddamasfatih@unja.ac.id



Penulis Kedua

Aulia Rachmawati, S.Kom., M.Eng. lahir pada tanggal 20 Juli 1998 dan merupakan dosen pada Program Studi Sistem Informasi, Fakultas Sains dan Teknologi, Universitas Jambi. Ia menyelesaikan pendidikan Sarjana (S1) di Universitas Muhammadiyah Surakarta dan Magister (S2) di Universitas Gadjah Mada. Bidang minat penelitian dan pengajarannya berfokus pada jaringan komputer (computer network) dan keamanan siber (cyber security), serta mencakup topik dalam disiplin sistem informasi, seperti pengembangan sistem informasi. Penulis aktif dalam kegiatan pendidikan, penelitian, dan pengabdian kepada masyarakat, serta berkomitmen pada pengembangan ilmu pengetahuan dan penerapan teknologi informasi. Penulis dapat dihubungi melalui email: auliarachmawati@unja.ac.id



Penulis Ketiga

Hamzah Alghifari, S.Tr.Kom., M.Kom. lahir di Bandung pada 4 Juni 1998. Penulis menyelesaikan Pendidikan S1 Teknologi Rekayasa Multimedia di Telkom University Bandung, S2 Sistem Informasi di Universitas

Dinamika Bangsa Jambi. Saat ini penulis merupakan Dosen Tetap Program Studi S1 Sistem Informasi, Fakultas Sains dan Teknologi, Universitas Jambi. Bidang keilmuan yang ditekuni meliputi Sistem Informasi Manajemen, Antarmuka Manusia dan Komputer, Desain Grafis, Digital Marketing dan Komunikasi. Penulis dapat dihubungi melalui email: hamzah.alghifari@unja.ac.id.



Penulis Keempat

Andicho Haryus Wirasapta, A.Md.T., S.Tr.T., M.Eng. lahir di Jakarta pada 15 September 1997. Penulis menyelesaikan pendidikan D-III Teknik Komputer di Politeknik Negeri Sriwijaya, D-IV Teknofisika Nuklir konsentrasi Elektronika Instrumentasi di STTN-BATAN,

serta Magister Teknik Elektro di Universitas Gadjah Mada. Saat ini penulis merupakan Dosen Tetap Program Studi Teknik Elektro, Fakultas Sains dan Teknologi, Universitas Jambi. Bidang keilmuan yang ditekuni meliputi elektronika, sistem kendali, Internet of Things, telekomunikasi, dan ketenaganukliran. Penulis dapat dihubungi melalui email: andicho@unja.ac.id.