

Keberkesanan Interaksi Manusia-AI dalam Membuat Keputusan di Tempat Kerja: Satu Kajian Empirikal

Ainun Rafieza binti Ahmad Tajuddin¹, Nur Elyani binti Mohammad¹, Nur Nadia binti Musa^{1,*}

¹ Unit Teknologi Maklumat, Kolej Komuniti Arau, Perlis, Malaysia

Email: ¹ainunrafieza@staf.kkarau.edu.my, ²nurelyani@staf.kkarau.edu.my, ³nadiamusa@staf.kkarau.edu.my

Email Penulis Korespondensi: nadiamusa@staf.kkarau.edu.my

Abstrak-Kecerdasan Buatan (AI) kini menjadi komponen penting dalam menyokong proses membuat keputusan harian di tempat kerja, khususnya dalam persekitaran digital yang memerlukan ketepatan, kecekapan dan respons pantas terhadap maklumat. Walaupun AI mampu mengoptimumkan tugas rutin dan membantu dalam penganalisan data berskala besar, tahap kebergantungan, kefahaman, dan kepercayaan pengguna terhadap sistem AI masih belum difahami secara menyeluruh. Dalam banyak kes, pengguna menerima atau menolak cadangan AI tanpa benar-benar memahami rasional di sebalik keputusan tersebut. Fenomena ini menimbulkan persoalan tentang sejauh mana AI boleh diterima sebagai rakan kolaboratif dalam membuat keputusan yang lebih bermakna dan bertanggungjawab. Kertas konsep ini mencadangkan satu kajian kuantitatif yang bertujuan untuk: (i) Mengenal pasti bentuk interaksi dan corak kebergantungan pengguna terhadap sistem AI di tempat kerja, (ii) Mengukur tahap kepercayaan, kefahaman dan keberkesanan keputusan pengguna terhadap saranan AI, dan (iii) Menganalisis hubungan antara kepercayaan, kefahaman serta keberkesanan keputusan terhadap tahap penerimaan teknologi AI. Kajian ini akan menggunakan soal selidik berstruktur dan dapatannya dianalisis serta melibatkan sekurang-kurangnya 100 responden profesional dalam kalangan pensyarah Kolej Komuniti Wilayah Utara, Malaysia. Data akan dianalisis menggunakan statistik deskriptif, korelasi Pearson dan regresi linear. Hasil yang dijangka daripada kajian ini akan dapat memberi kepada pembangunan kerangka konseptual kolaborasi manusia-AI yang lebih telus, etikal dan praktikal, serta menjadi panduan kepada organisasi dan pembuat dasar dalam merangka strategi pelaksanaan teknologi AI yang seimbang dan berfokuskan manusia.

Kata Kunci: Kecerdasan Buatan; Interaksi Manusia-AI; Kepercayaan Pengguna; Keberkesanan Keputusan; Kolaborasi Digital Di Tempat Kerja.

1. PENDAHULUAN

Kecerdasan Buatan (AI) semakin memainkan peranan penting dalam menyokong proses membuat keputusan di tempat kerja, terutamanya dalam tugas rutin yang memerlukan ketepatan, kelajuan, dan kecekapan tinggi. AI kini digunakan secara meluas dalam pelbagai bentuk, termasuk pembantu maya, sistem cadangan, dan automasi pintar, yang direka bentuk bagi membantu individu dan organisasi menangani beban kerja yang semakin kompleks serta membuat keputusan secara lebih berasaskan data (Creswell & Creswell, 2018; Hair et al., 2020; Siau & Wang, 2018). Di era Revolusi Industri 4.0, peningkatan penggunaan teknologi digital termasuk AI dalam persekitaran kerja menandakan peralihan paradigma dalam cara kerja dijalankan, sekali gus menuntut pemahaman baharu terhadap dinamika interaksi manusia-mesin.

Namun, kejayaan integrasi AI bukan hanya bergantung kepada aspek teknikal semata-mata. Ia juga menuntut pemahaman yang mendalam terhadap interaksi manusia dengan sistem AI, khususnya bagaimana pengguna memahami, menilai, dan bertindak berdasarkan saranan atau maklumat yang disediakan oleh sistem tersebut (Glikson & Woolley, 2020; Joshi et al., 2015; Zhang et al., 2020). Kepercayaan pengguna terhadap AI adalah elemen penting yang mempengaruhi keberkesanan penggunaan teknologi ini. Kepercayaan ini dipengaruhi oleh beberapa faktor termasuk ketelusan algoritma, tahap explainability, pengalaman penggunaan, serta kefahaman terhadap kebolehan dan keterbatasan sistem (Eiband et al., 2021; Ribeiro et al., 2016). Walaupun banyak kajian menekankan keupayaan teknikal AI, masih kurang perhatian diberikan terhadap peranan pengguna sebagai rakan kolaboratif dalam proses membuat keputusan bersama AI (Field, 2018; Shrestha et al., 2019; Dellermann, Reck, & Ebel, 2021). Kajian terdahulu banyak tertumpu pada prestasi algoritma semata-mata tanpa mengambil kira dimensi psikososial dan konteks organisasi yang menjadi latar interaksi manusia-AI. Dalam konteks kerja harian yang melibatkan pelbagai pertimbangan seperti nilai kemanusiaan, etika, dan konteks sosial, sistem AI perlu direka bentuk dengan mempertimbangkan keperluan pengguna sebenar (Siau & Wang, 2018; Glikson & Woolley, 2020).

Tambahan pula, fenomena seperti "automation bias" (Dietvorst et al., 2015) dan "algorithm aversion" menunjukkan bahawa pengguna mungkin cenderung untuk menerima atau menolak saranan AI secara membuta tuli, bergantung kepada pengalaman interaksi mereka sebelum ini. Kajian oleh Binns et al. (2018), Ribeiro et al. (2016), dan Jussupow et al. (2022) turut menegaskan bahawa persepsi terhadap keadilan dan kebolehpercayaan sistem AI amat bergantung kepada keupayaannya untuk memberikan justifikasi yang telus dan mudah difahami. Budaya organisasi juga memainkan peranan dalam mempengaruhi cara AI diterima dan digunakan oleh pekerja. Jika sesebuah organisasi mempunyai pendekatan terbuka terhadap inovasi dan menyediakan sokongan latihan serta polisi etika yang jelas, pengguna lebih cenderung untuk mempercayai dan mengadaptasi teknologi AI secara positif (Langer et al., 2021).

Akhir sekali, integrasi AI yang berkesan memerlukan sinergi antara pembangunan teknologi, kemahiran digital pengguna, dan dasar organisasi yang menyokong penggunaan teknologi secara etikal dan manusiawi. Oleh itu, kajian ini memberi tumpuan kepada pemahaman mendalam terhadap bagaimana AI digunakan dalam konteks kerja harian, bagaimana pengguna membina dan mengekalkan kepercayaan terhadap teknologi tersebut, serta bagaimana reka bentuk sistem boleh ditambah baik untuk menyokong keputusan bersama yang lebih berkualiti, telus dan beretika. Dalam era transformasi digital, AI semakin digunakan secara meluas untuk menyokong keputusan harian pekerja dalam pelbagai

sektor. Namun, sebahagian besar penyelidikan sedia ada masih memfokuskan kepada AI sebagai sistem autonomi sepenuhnya tanpa meneliti secara mendalam bentuk interaksi manusia dengan saranan atau keputusan yang dicadangkan oleh AI (Glikson & Woolley, 2020). Dalam konteks kerja harian, keputusan bukan hanya bergantung kepada data semata-mata, tetapi turut melibatkan nilai, empati, dan pertimbangan sosial yang kompleks. Oleh itu, pengabaian terhadap interaksi kolaboratif manusia-AI boleh menjejaskan nilai tambah manusia dalam proses membuat keputusan.

Masalah utama yang dikenalpasti ialah kurangnya pemahaman tentang bagaimana pengguna membina kepercayaan terhadap cadangan AI serta bagaimana sistem ini boleh direka bentuk bagi menyokong keputusan bersama yang lebih telus dan berkualiti. Pengguna akhir sering berdepan dengan kesukaran untuk memahami rasional di sebalik keputusan AI akibat kurangnya elemen explainability dan ketelusan sistem (Zhang et al., 2020; Eiband et al., 2021). Ketiadaan justifikasi yang jelas turut menyebabkan wujudnya jurang kepercayaan dan potensi salah faham terhadap fungsi sebenar AI. Di samping itu, tiada garis panduan yang kukuh atau kerangka kerja yang menyeluruh dalam organisasi bagi membimbing pelaksanaan AI yang berfokuskan pengguna. Tanpa struktur sokongan yang jelas, organisasi berisiko berdepan dengan kebergantungan melampau atau penolakan teknologi sepenuhnya, yang kedua-duanya boleh mengurangkan kualiti keputusan dan produktiviti (Dellermann et al., 2021). Tambahan pula, budaya organisasi yang tidak bersedia untuk menyepadukan AI secara strategik boleh mencetuskan konflik antara pekerja dan teknologi (Langer et al., 2021).

Justeru, terdapat keperluan mendesak untuk mengkaji bagaimana AI dapat dimanfaatkan secara kolaboratif dalam persekitaran kerja dengan menekankan aspek kepercayaan, kefahaman, dan keberkesanan keputusan. Penyelidikan ini bertujuan untuk mengisi jurang sedia ada dengan menyelidik bentuk interaksi, persepsi dan penerimaan pengguna terhadap AI, serta membina cadangan untuk reka bentuk sistem yang lebih mesra pengguna, telus dan menyokong agensi manusia dalam membuat keputusan. Kajian ini dijalankan bagi mencapai objektif-objektif berikut: 1) Mengenal pasti bentuk interaksi dan corak kebergantungan pengguna terhadap sistem kecerdasan buatan (AI) dalam membuat keputusan harian di tempat kerja; 2) Mengukur tahap kepercayaan, kefahaman dan keberkesanan keputusan pengguna terhadap cadangan AI melalui soal selidik dan analisis deskriptif; 3) Menganalisis hubungan antara kepercayaan, kefahaman dan keberkesanan keputusan terhadap tahap penerimaan pengguna terhadap sistem AI.

Literatur terkini menunjukkan pertumbuhan pesat dalam kajian berkaitan penggunaan AI dalam persekitaran kerja. Tumpuan utama ialah bagaimana AI boleh menjadi rakan kolaboratif dalam proses membuat keputusan melalui pendekatan hybrid intelligence (Glikson & Woolley, 2020; Shrestha et al., 2019). Pendekatan ini menekankan sinergi antara kebolehan pemprosesan AI yang pantas dan keupayaan pertimbangan manusia yang bersifat kompleks dan kontekstual. Sistem AI yang menyokong keputusan pengguna terbukti meningkatkan ketepatan, kelajuan dan keberkesanan proses kerja, terutamanya apabila AI digunakan dalam situasi berisiko tinggi seperti kewangan dan kesihatan (Siau & Wang, 2018; Dellermann et al., 2021).

Walau bagaimanapun, tahap keberkesanan sistem ini sangat bergantung kepada persepsi dan tahap kepercayaan pengguna. Kajian oleh Eiband et al. (2021) dan Zhang et al. (2020) menunjukkan bahawa sistem AI yang mampu memberikan justifikasi yang jelas terhadap saranannya lebih mudah diterima pengguna. Malah, konsep explainability dan user-centred design telah menjadi tumpuan utama dalam pembangunan AI yang dipercayai dan mudah digunakan. Kegagalan AI untuk menjelaskan keputusan atau menunjukkan alasan logik telah terbukti mengurangkan keyakinan pengguna, serta menyebabkan ketidakcekapan dalam penerimaan teknologi (Ribeiro et al., 2016; Binns et al., 2018). Tambahan pula, terdapat perbezaan pandangan dalam kalangan penyelidik berkenaan tahap autonomi AI yang ideal. Sesetengah kajian menyokong automasi penuh manakala yang lain menekankan pentingnya campur tangan manusia untuk mengelakkan automation bias (Dietvorst et al., 2015; Jussupow et al., 2022). Dalam banyak organisasi, integrasi AI juga dipengaruhi oleh struktur organisasi, budaya kerja, dan tahap literasi digital pekerja (Langer et al., 2021). Justeru, penyelidikan mengenai interaksi manusia-AI perlu diperluaskan kepada konteks kerja yang lebih umum, termasuk pekerjaan digital dan jarak jauh.

Oleh itu, kajian ini mengisi jurang dalam literatur dengan menyelidik bukan sahaja bagaimana AI digunakan, tetapi juga bagaimana pengguna memahami dan menerima teknologi tersebut dalam keputusan harian mereka. Dengan menggabungkan sintesis literatur dan dapatan empirikal, penyelidikan ini akan mencadangkan satu model konseptual baharu yang boleh dijadikan panduan dalam pembangunan dan pelaksanaan sistem AI yang lebih beretika, mesra pengguna dan menyokong kolaborasi manusia-teknologi secara mampan.

2. METODE PENELITIAN

2.1 Kerangka Konseptual

Berdasarkan sorotan literatur dan pernyataan masalah yang telah dibincangkan, kajian ini membangunkan satu kerangka konseptual yang menunjukkan hubungan antara bentuk interaksi manusia-AI, tahap kepercayaan, kefahaman, keberkesanan keputusan, dan penerimaan pengguna terhadap sistem AI di tempat kerja. Kerangka ini disesuaikan daripada pendekatan *hybrid intelligence* (Glikson & Woolley, 2020), prinsip *explainable AI* (Dellermann et al., 2021), dan sokongan daripada *Technology Acceptance Model* (Davis, 1989). Berikut merupakan jadual yang merumuskan hubungan antara pemboleh ubah utama dalam kajian ini:

Tabel 1. Hubungan Antara Pemboleh Ubah

Komponen	Hubungan Antara Pemboleh Ubah
Interaksi AI-Manusia	Memberi kesan langsung kepada tahap kepercayaan pengguna terhadap sistem AI
Kepercayaan terhadap AI	Mempengaruhi tahap kefahaman terhadap sistem serta kesediaan pengguna menerima teknologi
Kefahaman terhadap AI	Meningkatkan keberkesanan dalam membuat keputusan apabila berinteraksi dengan sistem AI
Keberkesanan Keputusan	Menjadi asas kepada pembentukan penerimaan yang positif terhadap teknologi AI
Penerimaan AI	Merupakan hasil akhir yang dipengaruhi oleh kepercayaan, kefahaman, dan keberkesanan sistem

2.2 Justifikasi Reka Bentuk Kajian

Kajian ini menggunakan reka bentuk penyelidikan kuantitatif berbentuk tinjauan (survey-based research design) kerana ia paling sesuai untuk mengukur dan mengenal pasti hubungan antara pemboleh ubah seperti tahap penggunaan AI, persepsi, dan keberkesanan keputusan pengguna. Reka bentuk ini membolehkan pengumpulan data daripada sampel yang besar secara sistematik dan objektif dalam tempoh masa yang terhad, serta menyediakan data kuantitatif yang boleh dianalisis secara statistik (Creswell & Creswell, 2018).

Tambahan pula, pendekatan kuantitatif membolehkan penyelidik menguji hipotesis dan membuat inferens berdasarkan hubungan antara pemboleh ubah, berbanding pendekatan kualitatif yang lebih menumpukan pada pemerolehan mendalam. Oleh kerana kajian ini bertujuan membina kerangka konseptual berasaskan dapatan empirikal tentang interaksi manusia-AI, reka bentuk tinjauan amat sesuai kerana ia dapat menilai tahap kepercayaan, kefahaman dan keberkesanan keputusan secara terukur dalam kalangan pengguna AI (Zhang et al., 2020; Dellermann et al., 2021).

2.3 Populasi dan Strategi Persampelan

Populasi sasaran bagi kajian ini merangkumi golongan profesional yang menggunakan aplikasi atau sistem kecerdasan buatan (AI) dalam tugas harian mereka. Ini termasuklah pengguna sistem cadangan, pembantu maya, analitik data automatik serta platform ramalan yang berperanan dalam menyokong proses membuat keputusan. Antara sektor utama yang terlibat ialah teknologi, kewangan, pendidikan dan perkhidmatan digital, di mana penggunaan AI semakin meluas sebagai alat sokongan kerja. Bagi memastikan keberkesanan pemilihan responden, kaedah pensampelan bertujuan (purposive sampling) telah dipilih kerana ia membolehkan penyelidik menumpukan kepada individu yang mempunyai pengalaman relevan dalam interaksi manusia-AI.

Responden dipilih berdasarkan beberapa kriteria utama iaitu: pertama, individu berumur antara 22 hingga 50 tahun; kedua, mempunyai pengalaman sekurang-kurangnya enam bulan dalam menggunakan sistem AI berkaitan kerja; dan ketiga, terlibat secara langsung dalam membuat keputusan yang dipengaruhi atau disokong oleh sistem AI. Jumlah sampel yang disasarkan adalah sekurang-kurangnya 100 orang responden. Angka ini dianggap mencukupi bagi memenuhi keperluan analisis statistik seperti korelasi dan regresi, yang memerlukan sampel bersaiz sederhana bagi memastikan tahap kebolehppercayaan dan ketepatan dapatan. Ini juga sejajar dengan saranan Hair et al. (2020) berkenaan jumlah minimum responden dalam kajian kuantitatif berasaskan model hubungan antara pemboleh ubah.

2.4 Kaedah dan Instrumen Pengumpulan Data

Pengumpulan data dijalankan menggunakan soal selidik berstruktur yang dibina berdasarkan konstruk teori dan penemuan daripada literatur terdahulu. Soal selidik ini dirancang bagi mengukur dimensi utama dalam kajian dan mengandungi lima bahagian utama yang saling berkaitan. Bahagian pertama merangkumi profil demografi responden, termasuk umur, jantina, sektor pekerjaan, dan pengalaman mereka dalam menggunakan AI. Bahagian kedua menilai tahap penggunaan AI dalam kerja harian, manakala bahagian ketiga memberi tumpuan kepada persepsi pengguna terhadap kecekapan sistem AI yang digunakan. Seterusnya, bahagian keempat mengukur tahap kepercayaan pengguna terhadap saranan atau keputusan yang dihasilkan oleh sistem AI.

Bahagian kelima pula menilai keberkesanan keputusan yang diambil hasil daripada kolaborasi antara manusia dan sistem AI. Setiap item dalam soal selidik akan menggunakan skala Likert lima mata, iaitu dari 1 (Sangat Tidak Setuju) hingga 5 (Sangat Setuju), bagi membolehkan pengukuran dilakukan secara lebih konsisten dan sistematik (Joshi et al., 2015). Soal selidik ini akan diedarkan melalui platform dalam talian seperti Google Forms. Responden akan diberikan surat kebenaran penyertaan yang menjelaskan tujuan kajian serta menjamin kerahsiaan data yang dikumpulkan. Semua maklumat yang diperolehi akan digunakan semata-mata untuk tujuan penyelidikan akademik dan dianalisis secara agregat bagi melindungi identiti individu yang terlibat.

2.5 Prosedur Analisis Data

Data yang dikumpulkan melalui soal selidik akan dianalisis menggunakan perisian Statistical Package for the Social Sciences (SPSS) versi terkini. Prosedur analisis akan dimulakan dengan analisis deskriptif bagi menggambarkan profil demografi responden serta skor min, sisihan piawai dan taburan jawapan bagi setiap konstruk utama iaitu tahap

penggunaan AI, kepercayaan, kefahaman, dan keberkesanan keputusan. Analisis ini bertujuan menyokong objektif pertama dan kedua kajian iaitu untuk mengenal pasti bentuk interaksi serta tahap persepsi pengguna terhadap sistem AI. Seterusnya, bagi memenuhi objektif ketiga kajian yang menumpukan kepada hubungan antara kepercayaan, kefahaman dan keberkesanan terhadap penerimaan AI, analisis korelasi Pearson akan digunakan untuk mengenal pasti kekuatan dan arah hubungan antara pemboleh ubah-pemboleh ubah tersebut. Analisis ini penting untuk menentukan sama ada terdapat perkaitan yang signifikan secara statistik dalam konteks interaksi manusia-AI.

Selain itu, analisis regresi linear akan dijalankan bagi menilai sejauh mana pemboleh ubah bebas seperti kepercayaan dan kefahaman dapat meramalkan keberkesanan keputusan atau tahap penerimaan pengguna terhadap sistem AI. Analisis ini membolehkan penyelidik menguji model empirikal yang dicadangkan serta membuat inferens yang menyokong kerangka konseptual kajian. Keseluruhan prosedur ini dirancang bagi memastikan penemuan kajian adalah selari dengan objektif yang ditetapkan dan dapat memberikan gambaran yang menyeluruh terhadap pola penggunaan dan persepsi terhadap AI di tempat kerja (Field, 2018).

2.6 Langkah Kesahan dan Kebolehpercayaan

Bagi menjamin kesahan kandungan (content validity), item-item dalam soal selidik ini telah dibangunkan berdasarkan konstruk yang diadaptasi daripada kajian terdahulu seperti Zhang et al. (2020), Glikson dan Woolley (2020), serta Eiband et al. (2021) yang membincangkan aspek kepercayaan, kefahaman dan keberkesanan penggunaan AI. Soal selidik ini kemudiannya akan dirujuk dan disemak oleh tiga orang pakar dalam bidang berkaitan iaitu kecerdasan buatan (AI), tingkah laku organisasi dan penyelidikan kuantitatif untuk memastikan kandungannya sesuai dan relevan dengan konteks kajian. Kajian rintis (pilot study) pula akan dijalankan ke atas seramai 30 responden, selaras dengan cadangan Krejcie dan Morgan (1970) yang mencadangkan sampel minimum bersamaan 10% daripada saiz sampel sebenar untuk menguji kefahaman, kejelasan dan konsistensi soalan.

Kajian rintis ini juga bertujuan mengesan sebarang isu semantik atau struktur ayat sebelum edaran sebenar. Kebolehpercayaan instrumen akan diuji menggunakan analisis Cronbach's Alpha, di mana nilai $\alpha \geq 0.70$ dianggap mencukupi untuk menunjukkan tahap kebolehpercayaan yang baik (Hair et al., 2020). Sekiranya terdapat mana-mana item yang menurunkan nilai Cronbach's Alpha secara keseluruhan, penambahbaikan atau penguguran item akan dilaksanakan bagi memastikan integriti data kajian sebenar.

2.7 Pertimbangan Etika

Kajian ini akan mematuhi sepenuhnya garis panduan etika penyelidikan seperti yang digariskan oleh Lembaga Etika Institusi. Semua prosedur yang melibatkan manusia sebagai responden akan mendapat kelulusan terlebih dahulu daripada jawatankuasa etika institusi bagi memastikan pematuhan terhadap prinsip keadilan, kebebasan bersuara, dan perlindungan hak asasi peserta. Setiap responden akan diberikan borang persetujuan dimaklumkan (informed consent form) yang menerangkan secara jelas tujuan kajian, prosedur pengumpulan data, manfaat dan risiko yang terlibat, jaminan kerahsiaan maklumat, serta hak untuk menarik diri pada bila-bila masa tanpa sebarang penalti atau kesan negatif terhadap mereka.

Kajian ini tidak akan mengumpul sebarang data peribadi sensitif seperti nombor kad pengenalan, alamat, atau maklumat kewangan. Semua jawapan daripada responden akan dikodkan dan dianalisis secara agregat bagi memastikan identiti individu tidak dapat dikenal pasti dalam sebarang laporan atau penerbitan. Selain itu, data yang dikumpul akan disimpan dalam sistem fail yang dilindungi kata laluan dan hanya boleh diakses oleh penyelidik utama. Prosedur keselamatan data akan dipatuhi sepenuhnya, termasuk penggunaan penyulitan fail jika perlu, bagi menjamin integriti dan privasi maklumat. Sebarang penggunaan data di luar skop kajian ini tidak akan dibenarkan tanpa kebenaran lanjut daripada responden. Langkah-langkah ini diambil bagi memastikan bahawa kajian ini dilaksanakan secara beretika, bertanggungjawab dan selaras dengan prinsip-prinsip penyelidikan sosial yang menghormati hak dan martabat peserta kajian.

2.8 Keterbatasan Kajian

Walaupun pendekatan kuantitatif memberikan kelebihan dari segi keupayaan untuk menghasilkan dapatan umum yang bersifat statistik, ia mempunyai keterbatasan dalam memahami secara mendalam sebab-sebab yang mendasari corak persepsi atau tingkah laku responden. Ini kerana data kuantitatif yang dikumpul melalui soal selidik berstruktur lebih tertumpu kepada pengukuran tahap atau kekerapan sesuatu fenomena, berbanding penerokaan naratif atau pengalaman subjektif yang boleh diperoleh melalui kaedah kualitatif. Oleh itu, dimensi konteks atau nuansa dalam interaksi manusia dan AI yang mungkin mempengaruhi penerimaan atau kepercayaan terhadap teknologi ini berpotensi tidak sepenuhnya terungkap melalui pendekatan ini.

Selain itu, penggunaan kaedah pensampelan bertujuan (purposive sampling) walaupun sesuai dengan objektif kajian yang memerlukan responden berpengalaman menggunakan AI dalam tugas harian, masih menimbulkan isu berkenaan kebolehulangan (replicability) dan kebolehumuman (generalizability) dapatan kepada populasi yang lebih luas. Oleh kerana sampel tidak dipilih secara rawak, kemungkinan terdapat bias pemilihan yang boleh mempengaruhi kesahan luaran kajian. Tambahan pula, kajian ini banyak bergantung kepada data sendiri (self-reported data) yang diperoleh melalui soal selidik, yang berpotensi terdedah kepada pelbagai bentuk bias, seperti bias sosial (social desirability bias), di mana responden mungkin memberikan jawapan yang dianggap lebih boleh diterima secara sosial, dan recall bias, di mana mereka mungkin tidak mengingati pengalaman interaksi mereka dengan AI secara tepat.

Kedua-dua jenis bias ini boleh menjejaskan ketepatan data. Namun demikian, beberapa langkah kawalan telah diambil bagi meminimumkan kesan keterbatasan ini. Antaranya ialah reka bentuk soal selidik yang telah disemak oleh pakar dalam bidang berkaitan bagi menjamin kejelasan dan ketelusan item, serta penggunaan skala pengukuran yang neutral dan tidak memihak. Selain itu, kajian rintis turut dijalankan bagi mengenal pasti sebarang kelemahan dalam instrumen sebelum ia digunakan dalam kajian sebenar. Diharapkan langkah-langkah ini dapat meningkatkan kesahan dalaman dan kebolehpercayaan dapatan kajian secara keseluruhan.

3. HASIL DAN PEMBAHASAN

Kajian ini dijangka menghasilkan pelbagai penemuan yang dapat memperkukuh kefahaman tentang bagaimana manusia dan kecerdasan buatan (AI) berinteraksi dalam membuat keputusan harian di tempat kerja. Salah satu hasil utama yang diunjurkan adalah pengenalan corak interaksi pengguna dengan sistem AI yang digunakan dalam persekitaran kerja digital. Ini termasuk pemahaman tentang sejauh mana pengguna bergantung kepada cadangan AI, serta bagaimana persepsi mereka terhadap kebolehpercayaan dan kecekapan sistem mempengaruhi keputusan yang diambil. Dapatan ini dijangka memperkukuh teori sedia ada dan pada masa sama mencadangkan satu kerangka konseptual awal yang boleh dijadikan panduan kepada organisasi dan pembangun sistem dalam membentuk model kerjasama AI-manusia yang lebih seimbang dan mesra pengguna (Dellermann et al., 2021).

Dari segi dapatan empirikal, penyelidik menjangkakan bahawa terdapat hubungan yang signifikan antara tahap kepercayaan dan kefahaman pengguna terhadap AI dengan keberkesanan keputusan harian. Dapatan ini disokong oleh kajian terdahulu yang menunjukkan bahawa pengguna lebih cenderung menerima saranan AI apabila mereka memahami bagaimana keputusan itu dihasilkan dan dapat melihat justifikasi yang telus (Zhang et al., 2020). Selain itu, kajian oleh Eiband et al. (2021) juga mencadangkan bahawa explainability merupakan faktor penting dalam membina kepercayaan dan penggunaan AI secara beretika. Justeru, pengguna yang mempunyai literasi AI yang lebih tinggi berkemungkinan membuat keputusan yang lebih tepat dan berkeyakinan apabila berinteraksi dengan sistem AI.

3.1 Sumbangan Akademik / Teori

Kajian ini berpotensi menyumbang secara signifikan kepada pembangunan teori dalam bidang interaksi manusia-komputer dan tingkah laku organisasi dengan memperkenalkan kerangka konseptual baharu mengenai kolaborasi manusia-AI. Berbeza dengan pendekatan sebelumnya yang lebih menumpukan pada keberkesanan teknikal sistem AI, kajian ini menekankan aspek manusiawi seperti kepercayaan, persepsi, dan pemahaman pengguna dalam konteks harian. Ini mengisi jurang dalam literatur yang banyak membincangkan AI sebagai sistem autonomi, tetapi kurang memberi perhatian kepada pendekatan hybrid intelligence yang menggabungkan kekuatan AI dan intuisi manusia (Dellermann et al., 2021; Glikson & Woolley, 2020). Penemuan kajian ini juga boleh memperkukuh teori tingkah laku teknologi seperti Technology Acceptance Model (TAM) dan Human-AI Teaming Framework dalam konteks kerja digital moden.

3.2 Aplikasi Praktikal

Dari segi aplikasi praktikal, kajian ini dapat membantu organisasi merancang strategi pelaksanaan AI yang lebih berfokuskan manusia (human-centric). Penemuan mengenai persepsi dan kepercayaan terhadap AI boleh digunakan untuk merangka program latihan yang meningkatkan literasi AI dalam kalangan pekerja, sekaligus menggalakkan penggunaan AI secara lebih berkesan dalam tugas harian. Hasil dapatan juga dapat dimanfaatkan oleh pembangun sistem dalam mereka bentuk antara muka pengguna yang lebih intuitif, telus dan boleh disesuaikan mengikut keperluan pengguna sebenar (Zhang et al., 2020; Eiband et al., 2021). Ini akan memastikan pengguna dapat memahami dan menilai saranan AI dengan lebih baik, sekali gus meningkatkan kualiti keputusan harian.

3.3 Manfaat Sosial

Kajian ini membawa potensi manfaat sosial yang besar, terutamanya dalam memperkukuh upaya individu untuk menyesuaikan diri dalam persekitaran kerja digital yang semakin kompleks. Dengan membantu pengguna memahami dan bekerjasama dengan AI secara lebih berkesan, kajian ini menyumbang kepada pembangunan masyarakat berpengetahuan digital (digitally literate society) dan meningkatkan tahap keyakinan terhadap teknologi dalam kalangan rakyat. Hal ini penting dalam konteks Malaysia yang sedang menuju ke arah Revolusi Industri 4.0, di mana AI menjadi komponen penting dalam pelbagai sektor pekerjaan (Glikson & Woolley, 2020). Selain itu, penggunaan AI yang lebih etikal dan bertanggungjawab seperti yang dicadangkan dalam kajian ini juga berupaya mengurangkan ketidaksamaan digital dan kebimbangan etika berkaitan automasi.

3.4 Implikasi Terhadap Industri

Bagi pihak industri, kajian ini memberikan panduan tentang bagaimana AI boleh diintegrasikan ke dalam aliran kerja organisasi tanpa mengurangkan peranan pekerja manusia. Dengan memahami faktor-faktor yang mempengaruhi kepercayaan dan penerimaan AI, organisasi dapat membina sistem yang lebih selamat, boleh dipercayai, dan mesra pengguna. Langkah ini dapat meningkatkan produktiviti dan kecekapan operasi, serta mengurangkan risiko penolakan teknologi oleh tenaga kerja sedia ada (Dellermann et al., 2021). Industri juga boleh memanfaatkan model kolaboratif yang dicadangkan untuk membentuk dasar dalaman berkaitan penggunaan AI secara seimbang dan menyeluruh.

3.5 Cadangan Dasar

Akhir sekali, kajian ini menyokong cadangan dasar awam yang menekankan kepentingan garis panduan jelas dalam penggunaan AI di tempat kerja. Dapatan kajian ini boleh dijadikan asas oleh pembuat dasar untuk merangka polisi kerja digital yang seimbang—yang bukan sahaja mengutamakan automasi, tetapi juga menekankan kesejahteraan psikososial pekerja. Polisi sedemikian perlu meliputi aspek kebolehhajelasan algoritma (algorithmic transparency), pemerkasaan pengguna melalui pendidikan AI, serta pelaksanaan prinsip human-in-the-loop dalam proses membuat keputusan penting (Eiband et al., 2021; Zhang et al., 2020). Dengan cara ini, pelaksanaan AI dapat dijalankan secara bertanggungjawab, progresif dan inklusif dalam jangka panjang.

4. KESIMPULAN

Kajian ini meneliti secara mendalam keberkesanan interaksi antara manusia dan sistem Kecerdasan Buatan (AI) dalam membuat keputusan harian di tempat kerja. Berdasarkan pendekatan kuantitatif melalui soal selidik yang dibangunkan secara sah dan boleh dipercayai, dapatan kajian ini diharap dapat memberikan gambaran jelas mengenai tahap kepercayaan, kefahaman serta keberkesanan pengguna terhadap AI dalam konteks tugasan profesional. Penemuan daripada kajian ini bukan sahaja menyumbang kepada pengembangan teori dalam bidang interaksi manusia-komputer dan tingkah laku organisasi, malah mampu membantu organisasi dan pembuat dasar dalam membentuk strategi pelaksanaan teknologi AI yang lebih inklusif, telus dan beretika (Glikson & Woolley, 2020; Dellermann et al., 2021). Pendekatan hybrid intelligence yang dibincangkan menekankan keperluan keseimbangan antara keupayaan teknikal AI dan pertimbangan manusia, sekaligus memperkukuh prinsip human-in-the-loop dalam setiap keputusan penting (Eiband et al., 2021). Tambahan pula, kajian ini membuka ruang untuk penyelidikan masa hadapan yang lebih mendalam tentang bentuk interaksi manusia-AI dalam pelbagai sektor pekerjaan, terutamanya dalam konteks pekerjaan jarak jauh dan automasi digital yang kian meluas. Literasi AI, kebolehhajelasan algoritma dan tahap penerimaan pengguna perlu terus diberi perhatian bagi memastikan AI benar-benar berfungsi sebagai rakan kolaboratif dan bukan sebagai pengganti mutlak kepada manusia (Zhang et al., 2020; Jussupow et al., 2022). Secara keseluruhannya, kajian ini memberi sumbangan penting kepada pembangunan kerangka konseptual interaksi manusia-AI yang lebih mampan dan mesra pengguna, serta menyediakan asas kukuh untuk pembentukan polisi teknologi yang lebih seimbang dalam era Revolusi Industri 4.0.

REFERENCES

- Binns, R., Veale, M., Van Kleek, M., & Shadbolt, N. (2018). 'It's reducing a human being to a percentage': Perceptions of justice in algorithmic decisions. *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*, 1–14. <https://doi.org/10.1145/3173574.3173951>
- Brynjolfsson, E., & McElheran, K. (2016). *Data in action: Data-driven decision making in US manufacturing* (CESifo Working Paper No. 5759).
- Creswell, J. W., & Creswell, J. D. (2018). *Research design: Qualitative, quantitative, and mixed methods approaches* (5th ed.). SAGE Publications.
- Dellermann, D., Reck, F., & Ebel, P. (2021). Breaking the decision-making black box: A taxonomy for design knowledge on explainable artificial intelligence. *Journal of Decision Systems*, 30(1), 1–20. <https://doi.org/10.1080/12460125.2020.1846486>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Eiband, M., Schneider, H., Bilandzic, M., Fazekas-Con, C., & Hussmann, H. (2021). Explaining with interactive examples: A user-centric strategy for AI explanations. *ACM Transactions on Interactive Intelligent Systems (TiiS)*, 11(1), 1–39. <https://doi.org/10.1145/3453170>
- Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). SAGE Publications.
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2020). *Multivariate data analysis* (8th ed.). Cengage Learning.
- Joshi, A., Kale, S., Chandel, S., & Pal, D. K. (2015). Likert scale: Explored and explained. *British Journal of Applied Science & Technology*, 7(4), 396–403. <https://doi.org/10.9734/BJAST/2015/14975>
- Jussupow, E., Heinzl, A., & Spohrer, K. (2022). The duality of algorithmic decision-making in knowledge work: Recurrent patterns of human–AI interaction. *Journal of the Association for Information Systems*, 23(2), 295–325.
- Krejcie, R. V., & Morgan, D. W. (1970). Determining sample size for research activities. *Educational and Psychological Measurement*, 30(3), 607–610.
- Langer, M., König, C. J., & Fitili, A. (2021). Information systems for the age of automation bias: A new challenge to system design. *Journal of Business Research*, 124, 518–530.
- Ribeiro, M. T., Singh, S., & Guestin, C. (2016). “Why should I trust you?” Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144. <https://doi.org/10.1145/2939672.2939778>
- Shrestha, Y. R., Ben-Menahem, S. M., & von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83.
- Siau, K., & Wang, W. (2018). Building trust in artificial intelligence, machine learning, and robotics. *Cutter Business Technology Journal*, 31(2), 47–53.

Zhang, Y., Liao, Q. V., & Bellamy, R. K. E. (2020). Effect of confidence and explanation on accuracy and trust calibration in AI-assisted decision making. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, 1–13. <https://doi.org/10.1145/3313831.3376207>