



UNIVERSITÀ DEGLI STUDI DI PADOVA

Dipartimento di Psicologia Generale

Corso di Laurea Magistrale in Psicologia Cognitiva Applicata

TESI DI LAUREA MAGISTRALE

Il cobot come agente sociale: stile comunicativo dell'IA e collaborazione umano-robot in contesti industriali

The Cobot as a Social Agent: AI Communication Styles and Human-Robot
Collaboration in Industrial Settings

Relatore:

Prof. Gamberini Luciano

Correlatrice:

Prof.ssa Nenna Federica

Laureando:

De Marchi Mario

Matricola: 2087596

Anno Accademico 2025-2026

INDICE

1.INTRODUZIONE	1
1.1 Obiettivi della presente tesi	3
2.BACKGROUND TEORICO	4
2.1 Cobot e Industry 5.0	4
2.2 Agenti conversazionali come supporto al lavoro	7
<i>Definizione e tassonomia degli AC</i>	
<i>Stili comunicativi</i>	
<i>Effetti dello stile comunicativo</i>	
<i>Adozione degli AC nel lavoro</i>	
2.3 Dinamiche cognitive e psicologiche nella collaborazione	13
con tecnologie robotiche e conversazionali	
<i>Technology acceptance</i>	
<i>Fiducia</i>	
<i>Percezione di autonomia e controllo sulla tecnologia</i>	
<i>Percezione di sorveglianza e monitoraggio durante il lavoro</i>	
<i>Stress</i>	
<i>Workload</i>	
<i>Processi attentivi</i>	
3. DOMANDE DI RICERCA E IPOTESI	21
<i>RQ1 — Manipulation check (H1.1, H1.2)</i>	
<i>RQ2 — Percezione soggettiva del supporto dell'AC (H2.1, RQ2.2)</i>	
<i>RQ3 — Dinamiche cognitive di workload, attenzione e affaticamento (H3)</i>	
4. METODI	23
4.1 Campione sperimentale	23
4.2 Setup tecnico e strumenti	24
4.3 Design sperimentale	25
<i>Tipologia di supporto dell'AC</i>	
<i>Pressione temporale</i>	
4.4 Compito sperimentale	29
<i>Modalità di presentazione dei contenitori</i>	
<i>Compito di conteggio dei pezzi smontati</i>	
<i>Induzione della pressione temporale</i>	
4.5 Procedura sperimentale	33

<i>Preparazione</i>	
<i>Training</i>	
<i>Sessione sperimentale</i>	
<i>Intervista finale e debriefing</i>	
4.6 Misure.....	34
<i>Misure self-report</i>	
<i>Misure di eye-tracking</i>	
<i>Misure qualitative</i>	
4.7 Analisi dei dati.....	37
5. RISULTATI.....	39
5.1 RQ1 — Manipulation check.....	39
<i>Distress (DSSQ)</i>	
<i>Pressione temporale percepita (NASA-TLX)</i>	
<i>Percezione del sistema (Godspeed)</i>	
5.2 RQ2 — Percezione soggettiva del supporto dell'AC.....	43
<i>Accettazione (TAM)</i>	
<i>Fiducia (S-TiAS)</i>	
<i>Controllo, autonomia e sorveglianza</i>	
5.3 RQ3 — Dinamiche cognitive di workload, attenzione e affaticamento.....	45
<i>Carico cognitivo percepito (NASA-TLX)</i>	
<i>Attivazione fisiologica: pupillometria</i>	
<i>Frequenza dei blink</i>	
6. DISCUSSIONE.....	51
6.1 Domanda di ricerca 1.....	51
6.2 Domanda di ricerca 2.....	52
6.3 Domanda di ricerca 3.....	53
6.4 Limiti e direzioni future.....	55
7. CONCLUSIONI.....	57
BIBLIOGRAFIA.....	59

1. INTRODUZIONE

L'Industria 5.0 ha spostato l'attenzione del settore manifatturiero dall'ottimizzazione tecnologica alla centralità della persona, ponendo il benessere del lavoratore, fisico e mentale, al centro del processo produttivo (Lu et al., 2022; Nenna et al., 2023; Bassi et al., 2025). Una forza lavoro che opera con minore fatica, maggiore supporto e ridotto carico cognitivo rappresenta un vantaggio competitivo per l'organizzazione in termini di produttività, soddisfazione lavorativa e senso di appartenenza, con ricadute positive sulla prevenzione dello stress lavoro-correlato (Baltrusch et al., 2022; Mariscal et al., 2024).

Realizzare questa visione, tuttavia, non è automatico. Le implementazioni tecnologiche rischiano di fallire quando non tengono conto delle variabili personali del lavoratore (Faccio et al., 2023): il successo di un'innovazione dipende dalla percezione che ne hanno gli operatori, dalle modalità con cui viene introdotta e dalla capacità del sistema di adattarsi alla persona piuttosto che il contrario (Liao et al., 2023). La tecnologia deve essere concepita come strumento di supporto e non come fonte di complessità aggiuntiva.

I robot collaborativi (cobot) rappresentano una risposta tecnologica diretta a questi principi: progettati per operare a stretto contatto con l'operatore umano, non ne sostituiscono il ruolo ma ne ampliano le capacità, facendosi carico dei compiti fisicamente gravosi o ad alta ripetitività per liberare risorse cognitive destinate ad attività di supervisione e decisione (Faccio et al., 2023). A differenza dell'automazione tradizionale, il cobot condivide lo spazio di lavoro con l'essere umano in modo dinamico (Di Pasquale et al., 2022). La condivisione dello spazio fisico e del compito definisce la collaborazione persona-robot (*human-robot collaboration*, HRC) come una forma di interazione qualitativamente diversa rispetto alle precedenti modalità di automazione industriale. L'adozione dei cobot nel settore manifatturiero è in costante crescita a livello europeo e globale, trainata da benefici documentati in termini di produttività, riduzione della fatica fisica e qualità del lavoro (Baltrusch et al., 2022; Faccio et al., 2023). Progettare un cobot tecnicamente affidabile non è tuttavia condizione sufficiente: è la qualità della HRC a determinare se la tecnologia diventa risorsa o ostacolo.

Accanto ai cobot, un'ulteriore categoria di tecnologie sta trasformando l'interazione persona-macchina nel contesto lavorativo: gli agenti conversazionali (*Conversational Agents*, AC), sistemi in grado di comunicare con l'utente in linguaggio naturale attraverso modalità testuale, vocale o multimodale (Allouch et al., 2021). Con il termine AC si intende qualsiasi sistema di dialogo capace di comprendere e generare linguaggio naturale, includendo sia i *chatbot* propriamente detti sia le più ampie famiglie di intelligenze artificiali (*artificial intelligence*, IA) conversazionali; i due termini sono spesso usati in modo intercambiabile nella letteratura recente (Adamopoulou e Moussiades, 2020; Allouch et al., 2021).

Gli AC variano non solo per funzionalità tecniche, ma anche per lo stile comunicativo con cui si rivolgono all'utente. La terminologia usata in letteratura per descrivere questa dimensione non è uniforme (Van Pinxteren et al., 2020): gli elementi più ricorrenti sono la presenza o assenza di segnali sociali, di calore percepito (*warmth*, inteso come percezione delle intenzioni benevolenti dell'altro) e di tratti antropomorfi nella comunicazione (Nass et al., 1994; Klein, 2025). Due poli emergono come maggiormente presenti in letteratura (Shum et al., 2018; Wang et al., 2023). Il primo, che definiamo *Relational Conversational Agent* (R-CA), raggruppa etichette quali *relational*, *open-domain*, *human-like* e *warmth* (Klein, 2025; Kolomaznik et al., 2024): prevede *social cues*, messaggi di supporto emotivo e di incoraggiamento, e un registro amichevole che tende a essere percepito con maggiore *warmth* e maggiore antropomorfismo (Xu et al., 2022; Cai et al., 2024). Il secondo, che definiamo *Functional Conversational Agent* (F-CA), è definito in modo sottrattivo rispetto all'R-CA (Adamopoulou e Moussiades, 2020): più diretto e conciso, privo di *social cues* e di messaggi di supporto emotivo, non incorpora tratti volutamente umani e tende a essere percepito come più meccanico, con minore antropomorfismo attribuito (Shum et al., 2018). Gli studi che trattano questa distinzione raramente si riferiscono a contesti industriali: la maggior parte della letteratura sugli AC è stata sviluppata in ambiti di servizio, e-commerce o assistenza sanitaria, rendendo difficile trasferire direttamente le conclusioni a scenari di collaborazione umano-cobot. I pochi studi condotti in contesti lavorativi reali non offrono indicazioni univoche: alcuni rilevano che lavoratori industriali tendono a preferire sistemi privi di segnali emotivi e comportamenti sociali. Dotare un cobot di *social cues* come quelli veicolati da un AC (voce, linguaggio naturale, feedback adattivo) può generare negli operatori l'impressione di avere a che fare con un sistema integrato di natura sociale: un'entità con capacità di interazione, di supporto al lavoro e di riconoscimento dell'utente che vanno oltre la mera funzionalità operativa. Ne emerge una nuova concezione del cobot come agente sociale nell'ambiente di lavoro. Questo fenomeno è teoricamente coerente con il framework CASA (Computers Are Social Actors; Nass et al., 1994), secondo cui qualsiasi entità che emette segnali sociali attiva automaticamente schemi relazionali, indipendentemente dalla consapevolezza dell'utente circa la natura artificiale del sistema.

La letteratura offre qualche esempio di robot industriali dotati di segnali sociali in contesto manifatturiero. Nunnari et al. (2025) hanno integrato un avatar con funzioni di co-regolazione emotiva in un compito di assemblaggio; Freire et al. (2024) hanno sviluppato un'architettura cognitiva socialmente adattiva (DAC-HRC) per un cobot in contesto di riciclo elettronico, con alternanza dei turni e comunicazione gestuale; Kopp et al. (2023) hanno manipolato il framing linguistico del cobot (nome umano vs. etichetta funzionale), rilevando che la sola denominazione antropomorfica modifica le dinamiche di fiducia e diffidenza. Sui cobot specificamente, tuttavia, gli studi rimangono pochi e i loro esiti non uniformi: Nielsen et al. (2026) documentano una preferenza degli operatori industriali per cobot privi di display emotivi, segnalando che l'antropomorfismo non produce effetti lineari sull'accettazione in contesti produttivi. A nostra conoscenza, la letteratura offre ancora scarsa evidenza su come lo stile comunicativo di un AC integrato a un cobot industriale, e in particolare la distinzione tra R-CA e F-CA, influenzi in modo simultaneo accettazione tecnologica, fiducia e carico cognitivo degli operatori in condizioni controllate.

Inoltre, il contesto manifatturiero è caratterizzato da una naturale variabilità delle richieste operative. Per tale motivo, l'accettazione e la percezione di un AC o di un cobot sociale vanno lette anche alla luce del fatto che, in momenti diversi del processo produttivo, la stessa forma di supporto può risultare più o meno appropriata e ben accolta. In quanto tecnologie che introducono una componente sociale nell'interazione, il loro impatto non può essere considerato indipendente dal momento e dalle condizioni in cui vengono impiegate. Una di queste condizioni è la pressione temporale: l'urgenza di compiere un lavoro entro scadenze talvolta molto stringenti. Sotto elevata pressione temporale, le risorse cognitive disponibili si riducono, orientando l'elaborazione verso modalità euristiche (Hermanns e Teubner, 2025): in queste condizioni, i segnali socio-emotivi dell'R-CA possono competere con le richieste del compito primario, trasformandosi da risorsa a interferenza, meccanismo coerente con il modello della capacità limitata di elaborazione (Lang, 2000). Klein (2025), in una meta-analisi su 142 studi, documenta che gli effetti positivi della *human-likeness* si attenuano in contesti *task*-critici ad alta complessità cognitiva. Verhagen et al. (2022) mostrano che ridurre la frequenza comunicativa del robot sotto pressione temporale migliora sia la performance di team sia la fiducia nel robot. A nostra conoscenza, i sistemi che integrano AC e cobot in ambienti industriali sono stati studiati in sul piano tecnico (Li et al., 2023; Bousdekis et al., 2025), senza manipolare lo stile comunicativo dell'agente né misurarne gli effetti su accettazione, fiducia e carico cognitivo dell'operatore.

L'accettazione di una tecnologia dipende anche dal modo in cui le sue caratteristiche progettuali si combinano con le condizioni operative in cui l'utente si trova: uno stesso stile comunicativo può quindi essere percepito in modo diverso a seconda del contesto e del momento d'uso. Il presente studio si propone di verificare se, anche in un paradigma di lavoro con cobot, la preferenza soggettiva per lo stile comunicativo di un AC vari in funzione della pressione temporale, in linea con quanto osservato in altri domini.

1.1 Obiettivi della presente tesi

La presente tesi adotta un paradigma sperimentale controllato che confronta tre condizioni di stile comunicativo del AC integrato al cobot, all'interno di un contesto sperimentale che include diverse condizioni di pressione temporale. Le tre manipolazioni del AC sono:

- i. **No-CA (No Conversational Agent):** il sistema prevede cobot e interfaccia a schermo con le istruzioni lavorative, senza AC. Tale condizione funge da baseline.
- ii. **R-CA (Relational Conversational Agent):** il sistema di supporto prevede cobot e interfaccia con un AC multimodale (testo scritto a schermo e voce parlata), che comunica le istruzioni lavorative e fornisce messaggi di supporto e motivazione. Presenta modalità con maggiori social cues, come l'uso della prima persona plurale per i compiti che prevedono collaborazione persona-cobot.

iii. **F-CA (Functional Conversational Agent):** il sistema di supporto prevede cobot e interfaccia con un AC multimodale (testo scritto a schermo e voce parlata), che comunica le istruzioni lavorative con modalità neutre, senza aggiungere connotati sociali, di supporto o motivanti.

Le tre condizioni vengono confrontate in un compito ecologicamente valido di disassemblaggio con cobot UR10e, svolto in condizioni con e senza pressione temporale. Gli obiettivi procedono dal piano soggettivo a quello fisiologico:

- a. esaminare gli effetti dello stile comunicativo del AC sull'esperienza lavorativa: accettazione, fiducia percepita, senso di autonomia e percezione di sorveglianza;
- b. valutare in che misura la pressione temporale modula tali effetti;
- c. verificare se le diverse condizioni di stile comunicativo del AC producono differenze nel carico cognitivo e nello stress tramite misure fisiologiche (pupillometria, *blink rate*) e self-report.

L'integrazione di indicatori soggettivi e oggettivi consente di triangolare l'impatto dei diversi stili comunicativi sugli operatori, fornendo indicazioni per la progettazione di sistemi di supporto HRC centrati sul lavoratore.

2. BACKGROUND TEORICO

2.1 Cobot e Industria 5.0

I cobot costituiscono una risposta tecnologica diretta ai principi di Industria 5.0: la loro rilevanza non si esaurisce nella riduzione del carico fisico, ma investe una ridefinizione del contributo umano verso funzioni di supervisione e gestione dell'imprevisto. Per comprendere le implicazioni psicologiche di questa ridefinizione, è necessario partire dalle caratteristiche strutturali che distinguono l'HRC da ogni altra forma di automazione.

Sul piano tecnico, i cobot sono progettati per operare a contatto diretto con l'operatore, senza gabbie protettive; sistemi di forza-coppia rilevano contatti inattesi e attivano l'arresto immediato del braccio robotico, in conformità con la norma ISO/TS 15066:2016, che definisce i limiti di forza e pressione ammissibili per ciascuna modalità collaborativa (Faccio et al., 2023). A differenza degli impianti robotici tradizionali, i cobot moderni sono riprogrammabili senza competenze specialistiche avanzate, rendendoli accessibili alle piccole e medie imprese (PMI) e adattabili a lotti produttivi variabili (Faccio et al., 2023). Questa flessibilità ha accelerato la diffusione dei cobot nei settori manifatturieri a media intensità tecnologica, dove i compiti più diffusi (manipolazione, assemblaggio, disassemblaggio) richiedono precisione ripetitiva e generano fatica muscolare accumulata (Baltrusch et al., 2022).

La letteratura sull'HRC ha a lungo analizzato variabili fisiche e di sicurezza, come velocità del robot, prossimità e prevedibilità dei movimenti; un numero crescente di studi ha poi esteso l'analisi alla dimensione psicologica e comunicativa dell'interazione (Arai et al., 2010; Lu et al., 2022; Faccio et al., 2023). Gli approcci multimetodo, che integrano misure soggettive e fisiologiche, si sono affermati come strumento per rilevare carico mentale, fatica e livelli di risorse attentive in tempo reale durante l'HRC (Nenna et al., 2022, 2023; Gervasi et al., 2023; Orlando et al., 2025). Questo spostamento metodologico ha reso visibili alcune variabili che un approccio ingegneristico lasciava fuori campo: la risposta psicologica del lavoratore al cobot non è un residuo soggettivo secondario, ma un predittore dell'esito dell'adozione.

La ricerca recente ha esteso questa prospettiva, proponendo che il cobot possa elicitare risposte psicologiche di natura relazionale in virtù dei segnali sociali che emette, indipendentemente dalla consapevolezza dell'operatore di interagire con una macchina. Il paradigma CASA (Computers Are Social Actors; Nass et al., 1994) ha fornito il quadro teorico per anticipare questo meccanismo nel contesto della HCI, e la sua estensione al cobot fisico è corroborata da studi nel decennio recente. Kopp et al. (2022), in uno studio su operai di fabbrica tedeschi, mostrano che presentare il cobot con *framing* linguistico antropomorfo (nome umano, pronomi personali, voce attiva) aumenta significativamente la fiducia iniziale rispetto al *framing* tecnico-meccanico; l'effetto è però moderato dalla percezione di sostituzione tecnologica: i lavoratori che temono di essere sostituiti reagiscono negativamente all'antropomorfismo, invertendo il vantaggio. In un secondo studio dello

stesso gruppo (Kopp et al., 2023), il *framing* antropomorfo amplifica sia la fiducia iniziale sia la diffidenza conseguente agli errori del sistema, evidenziando che fiducia e diffidenza sono costruiti con traiettorie evolutive separate e che l'attivazione di schemi relazionali aumenta la sensibilità dell'operatore alle prestazioni del sistema in entrambe le direzioni.

Sviluppi paralleli hanno esplorato l'integrazione di segnali sociali nel comportamento del cobot. Freire et al. (2024) propongono la DAC-HRC (Distributed Adaptive Control for Human-Robot Collaboration), un'architettura cognitiva che integra nel cobot un insieme di comportamenti socialmente adattativi. Sul piano prossemico, il modulo SASE (Socially Adaptive Safety Engine) regola la distanza fisica del cobot in funzione del livello di fiducia stimato per ciascun operatore. Sul piano gestuale, un repertorio bidirezionale di gesti predefiniti consente all'operatore di impartire comandi al cobot, che ne conferma il riconoscimento tramite feedback visivo, una forma di turno comunicativo che presuppone attenzione reciproca e comprensione condivisa dei segnali. Sul piano della personalizzazione, il *Worker Model* memorizza le preferenze individuali e adatta nel tempo il comportamento del sistema al singolo operatore. Rispetto all'approccio del presente studio, che manipola lo stile comunicativo di un agente conversazionale verbale, Freire et al. (2024) operano sul versante non verbale della stessa ipotesi di fondo: che l'incorporazione di segnali sociali adattativi nel cobot modifichi la qualità dell'interazione umano-robot in contesti industriali.

Nunnari et al. (2025) hanno sviluppato il *covatar*, un cobot industriale (Fanuc CRX10iA/L) abbinato a un *avatar* virtuale con funzioni di co-regolazione emotiva, per supportare il *flow* e il benessere degli operatori durante un compito di assemblaggio ripetitivo. Il sistema si basa sul modello BASSF (Boredom, Anxiety, Self-efficacy, Self-compassion, Flow): sensori rilevano in tempo reale i valori PAD (Pleasure, Arousal, Dominance) dall'espressione facciale dell'operatore, il modello classifica lo stato emotivo e attiva uno dei 16 interventi pre-configurati, consistenti in feedback verbali dell'*avatar* e/o modulazioni della velocità del cobot. In due studi sequenziali (N = 20 e N = 12), i lavoratori riferiscono maggiore fluenza lavorativa e meno noia rispetto alla condizione di base; tuttavia, i partecipanti con atteggiamenti negativi verso i robot si sentono più tesi e percepiscono meno supporto e hanno percepito meno supporto nei momenti critici. Gli autori concludono che il sistema mostra potenziale per sostenere benessere e *flow* in HRC industriale, ma che la limitatezza dei campioni e la forte dipendenza delle risposte dal profilo individuale non consentono ancora di generalizzare i risultati. Lo stile comunicativo dell'*avatar* è implicitamente relazionale (messaggi brevi, empatici, orientati alla riformulazione cognitiva) ma non viene manipolato sperimentalmente né confrontato con uno stile alternativo; il sistema, inoltre, non è un AC in senso proprio, poiché non risponde a richieste dell'operatore né assiste nel compito, ma eroga interventi socio-emotivi in risposta allo stato affettivo rilevato.

I risultati di questo filone non sono però uniformi. Nielsen et al. (2026), in uno studio a metodi misti con 37 futuri lavoratori industriali (tirocinanti danesi trattati come utenti *proxy* del manifatturiero), esplorano percezioni, atteggiamenti ed esperienze di fronte a cobot, assistenti conversazionali AI e cobot *voice-enabled*. Il disegno è quasi-sperimentale controbalanciato: i partecipanti visitano in ordine invertito due contesti, un'esibizione artistica con robot emotivamente espressivi e uno Smart Lab industriale in cui interagiscono con un cobot Franka Emika Panda abbinato a un assistente

vocale IA per un compito di assemblaggio. I dati quantitativi provengono dalla NARS (*Negative Attitude toward Robots Scale*) somministrata prima e dopo le esposizioni; i dati qualitativi da appunti scritti, conversazioni di gruppo e osservazioni comportamentali.

Sul piano quantitativo, l'esposizione alle tecnologie produce un aumento significativo delle attitudini negative sulla subscala *Social Influence* della NARS (variazione media $M = 8.51$, $SD = 6.56$, $p \leq .001$), con effetto più pronunciato nei partecipanti più giovani, con bassa affinità tecnologica e senza esperienza pregressa con robot e IA; le sottoscale *Interaction* ed *Emotions* non cambiano significativamente. Il pattern indica che l'incontro ravvicinato con robot socialmente espressivi amplifica la preoccupazione di essere influenzati da sistemi artificiali. In altre parole, tale incontro non influenza la resistenza all'interazione in sé, ma la percezione di una pressione sociale indesiderata. Sul piano qualitativo emergono cinque temi che possono essere rilevanti per future ricerche: visione radicalmente strumentale del cobot (il robot come attrezzo, non come agente relazionale); rifiuto dei display emotivi robotici, percepiti come artificiali e controproducenti in fabbrica; accettazione dell'assistenza vocale se strettamente funzionale, concreta ed esplicita; fiducia pragmatica nell'esecuzione affidabile delle istruzioni, con riluttanza ad affidarsi a sistemi con capacità emotiva; apertura condizionale alla voce purché rimanga uno strumento operativo.

La trasferibilità diretta al presente studio rimane tuttavia parziale. Il contesto dell'esibizione artistica era provocatorio per design e gli autori stessi riconoscono questa asimmetria come limitazione; il sistema vocale adottato si limitava a comandi vocali, senza AC con funzione di supporto al compito. La questione di quali caratteristiche comunicative supportano l'operatore in modo efficace, e in quali condizioni operative, rimane aperta.

2.2 Agenti conversazionali come supporto al lavoro

Definizione e tassonomia degli AC

Il termine agente conversazionale (Allouch et al., 2021) indica qualsiasi sistema in grado di produrre e interpretare linguaggio naturale in forma dialogica, indipendentemente dalla modalità di *output* e dall'architettura tecnologica sottostante. All'interno di questa categoria si distinguono famiglie con caratteristiche percettive e funzionali differenti.

I *chatbot* tradizionali replicano le caratteristiche di un interlocutore umano all'interno di un'interfaccia testuale di tipo *chat*. Il sistema opera tramite corrispondenza di pattern su regole predefinite o recupero da database chiusi: le risposte sono prevedibili, il registro è uniforme e l'interazione tende a essere percepita come meccanica e impersonale. Questi sistemi mantengono una memoria del contesto limitata e trovano impiego prevalente nel *customer care* (Adamopoulou e Moussiades, 2020). A partire dai primi anni 2010, lo sviluppo di modelli linguistici di grandi dimensioni (*Large Language Model*, LLM) ha progressivamente esteso le potenzialità anche di questa categoria di AC (Adamopoulou e Moussiades, 2020; Shum et al., 2018).

I sistemi basati LLM generano risposte adattive al contesto, mantengono coerenza conversazionale su più turni e incorporano segnali sociali (variazioni di tono, riconoscimento dell'errore, adattamento al registro dell'utente) che li rendono percettivamente più fluidi e più antropomorfi (Klein, 2025).

I sistemi *voice-based* aggiungono alle caratteristiche dell'interazione testuale la dimensione paralinguistica: prosodia, ritmo e timbro contribuiscono alla percezione di presenza sociale indipendentemente dall'architettura sottostante (Allouch et al., 2021). Il riconoscimento vocale consente di ricevere linguaggio naturale parlato e di produrre risposte audio; questa modalità attiva schemi di risposta sociale analoghi a quelli evocati da un interlocutore umano, anche in assenza di un corpo fisico (Nass et al., 1994). Rispetto al canale testuale, il canale vocale introduce una variabile aggiuntiva nella progettazione del AC: le scelte paralinguistiche (tono caldo o neutro, ritmo rapido o pausato) diventano veicolo di segnali socio-emotivi indipendenti dal contenuto verbale.

Indipendentemente dalla famiglia tecnologica di appartenenza, gli AC variano lungo una seconda dimensione trasversale, lo stile comunicativo, che determina in quale misura il sistema incorpora segnali sociali, *warmth* e tratti antropomorfi nell'interazione. È questa dimensione, e non la distinzione tecnica tra chatbot e LLM, a costituire la variabile indipendente del presente studio.

Stili comunicativi

La letteratura sugli AC non ha adottato una terminologia univoca: etichette come *social*, *relational*, *task-focused* e *transactional* ricorrono con sovrapposizioni concettuali tra contributi diversi (Shum et al., 2018; Chattaraman et al., 2019; Van Pinxteren et al., 2020). Da questi contributi emergono due poli ricorrenti dello stile comunicativo. Al polo funzionale, il dialogo è circoscritto al compito: comprende richieste di informazioni puntuali e risposte appropriate al contesto operativo, senza componenti relazionali (Chattaraman et al., 2019; Adamopoulou e Moussiades, 2020). Al polo relazionale, l'AC integra strategie conversazionali informali (saluti, *small talk*, espressioni di supporto emotivo, segnali di ascolto attivo) orientate a obiettivi relazionali e al mantenimento del rapporto con l'utente (Chattaraman et al., 2019; Kolomaznik et al., 2024; Klein, 2025). Tra i due estremi si collocano registri ibridi che combinano struttura *funzionale* e segnali socio-emotivi attenuati (Xu et al., 2022). La Tabella 2.1 riporta le definizioni operative dei due poli negli studi esaminati.

Nel presente studio le etichette F-CA e R-CA sono adottate come operazionalizzazioni proprie che mappano rispettivamente il polo transazionale e quello relazionale. F-CA designa un AC limitato al trasferimento di informazioni operative, privo di componenti socio-emotive; R-CA designa un AC che affianca le istruzioni operative con elementi di supporto relazionale: incoraggiamento, uso della prima persona plurale nelle fasi di collaborazione umano-cobot, riconoscimento del lavoro svolto. La descrizione operativa di ciascuna condizione è riportata nella sezione Metodo (§4.3).

	Polo funzionale	Polo relazionale
Shum et al. (2018)	Task-completion system ▶ Efficienza · dominio chiuso · conclusione rapida · output informativo	Social chatbot ▶ Connessione emotiva · senso di appartenenza · intelligenza emotiva (EQ)
Chattaraman et al. (2019)	Task-oriented interaction style ▶ Formale · dialogo esclusivamente sul compito · obiettivi funzionali	Social-oriented interaction style ▶ Informale · relazionale · saluti · small talk · supporto emotivo · obiettivi socio-emotivi
Adamopoulou e Moussiades (2020)	Task-based ▶ Funzioni specifiche · richieste puntuali · risposta appropriata	Chat-based / Conversational ▶ Conversazione naturale · simile a interlocutore umano · open-domain
Van Pinxteren et al. (2020)	(polo opposto non definito) ▶ Comunicazione impersonale · unidirezionale · solo contenuto informativo	Human-like communicative behaviors ▶ Somiglianza umana · reattività ai bisogni · modalità verbale, non verbale, aspetto
Allouch et al. (2021)	Goal-oriented CA (taskbot) ▶ Compiti multi-step · customer service · assistente personale · dominio chiuso	Social-supporting agent ▶ Supporto emotivo e psicologico · benessere · apprendimento · nessun dominio fisso
Xu et al. (2022)	Task-oriented communication style ▶ Orientato agli obiettivi · efficiente · formale · dialogo solo sul compito	Social-oriented communication style ▶ Personale · bisogni emotivi · informale · relazionale · small talk · elogi sociali
Kopp et al. (2022, 2023)	Machine-like framing ▶ Nome tecnico · pronomi impersonale · linguaggio meccanico · nessun costrutto umano	Human-like / Anthropomorphic framing ▶ Nome umano · pronomi personale · voce attiva · autonomia · linguaggio antropomorfo
Wang et al. (2023)	Task-oriented ▶ Orientato agli obiettivi · strutturato · conciso · senza ridondanze	Social-oriented ▶ Parole empatiche · supporto emotivo · costruzione della relazione
Kolomaznik et al. (2024)	Technical proficiency (polo implicito) ▶ Efficienza tecnica · competenza funzionale · nessuna dimensione emotiva	Socio-emotional intelligence ▶ Fiducia · empatia · rapport · coinvolgimento · antropomorfizzazione
Cai et al. (2024)	Task-oriented communication style ▶ Formale · dialogo sul compito · orientato agli obiettivi · chiarisce · informa	Social-oriented communication style ▶ Comunicazione educata e relazionale · saluti · small talk · assistenza emotiva · bisogni emotivi

Klein (2025)	Less human-like / machine-like ▶ <i>Privo di social cues · assenza di caratteristiche umane · registro neutro</i>	Human-like / Relational ▶ <i>Stile colloquiale · emoji · ascolto attivo · avatar sorridente · presenza sociale</i>
Chen et al. (2025)	Task oriented ▶ <i>Focalizzato sugli obiettivi · forte senso di scopo</i>	Social-oriented (interaction-oriented) ▶ <i>Relazioni sociali forti · enfasi sull'individualità</i>
Wang et al. (2026)	Perceived reliability and system competence ▶ <i>Accuratezza · reattività · consistenza · trasparenza · sicurezza dati → fiducia cognitiva</i>	Human-like interaction and emotional connection ▶ <i>Naturalezza · empatia · personalizzazione · presenza sociale → fiducia affettiva</i>

Tabella 2.1 Definizioni operative dei due poli negli studi esaminati. In blu la terminologia usata negli studi per indicare lo stile conversazionale dell'agente. Il simbolo ▶ indica la sintesi degli attributi in italiano. Kopp et al. (2022, 2023) riportano due studi sequenziali con la medesima manipolazione sperimentale.

Come già anticipato, gli esseri umani tendono a rispondere socialmente a qualsiasi sistema che presenti segnali sociali (voce, linguaggio naturale, tono emotivo) attivando automaticamente schemi propri dell'interazione interpersonale, indipendentemente dalla consapevolezza di interagire con una macchina (Nass et al., 1994): i computer, se dotati di caratteristiche sociali minime, elicitano le stesse euristiche valutative che gli esseri umani applicano agli interlocutori umani. Ricerche successive hanno mostrato che tale risposta si estende all'attribuzione di tratti di personalità (calore, competenza, dominanza) sulla base dei segnali linguistici prodotti dal sistema (Nass et al., 1995), e che il *matching* tra lo stile comunicativo dell'AC e le aspettative dell'utente rafforza questa attribuzione, seguendo il principio della *similarity-attraction* documentato nell'interazione interpersonale (Moon e Nass, 1996).

Lo stile comunicativo di un AC orienta quindi il tipo di schema valutativo attivato dall'utente. Un R-CA che produce segnali di calore e disponibilità sollecita la dimensione affettiva dell'attribuzione interpersonale, con effetti attesi su variabili relazionali come fiducia e accettazione; un F-CA che si limita all'informazione procedurale attiva la dimensione della competenza percepita, mantenendo l'interazione in un registro strumentale. Se e in quale misura questi effetti si manifestino in un contesto di HRC industriale, e come vengono moderati dalla pressione temporale, è la domanda a cui il presente studio intende rispondere.

Effetti dello stile comunicativo

Xu et al. (2022), in due studi sperimentali su *chatbot*, mostrano che lo stile relazionale aumenta la soddisfazione dell'utente attraverso la percezione di calore come mediatore; l'ansia di attaccamento del consumatore modera l'intensità dell'effetto, con risposte più pronunciate negli utenti con maggiore sensibilità relazionale. Cai et al. (2024) replicano il meccanismo in un contesto di *service*

failure: la comunicazione a orientamento sociale aumenta soddisfazione e fiducia rispetto allo stile funzionale, ancora tramite calore percepito, suggerendo una certa robustezza del *pattern* anche sotto condizioni di aspettativa delusa. Wang et al. (2023), in uno studio a scenari su gestione del fallimento del servizio in *e-commerce*, identificano fiducia cognitiva e fiducia affettiva come doppio meccanismo mediatore: entrambe mediano il legame tra stile e soddisfazione, ma con dinamiche diverse al variare della complessità del compito; sotto alta complessità, lo stile *relational* viene preferito in maniera minore, indicando che la componente relazionale non inibisce necessariamente la valutazione basata su competenza percepita, ma può potenziarla in situazioni impegnative. Chen et al. (2025), in due esperimenti su robot di servizio (N = 272 ciascuno), documentano che lo stile *relational* supera quello *funzionale* sulla disponibilità continuata all'uso (M = 5.87 vs M = 5.11, $F(1,271) = 40.87$, $p < .001$); le emozioni positive fungono da mediatore e il contesto di servizio da moderatore: l'effetto è più pronunciato nei contesti edonici ($\beta = 1.67$) rispetto a quelli utilitaristici ($\beta = 0.97$), indicando che l'obiettivo perseguito dall'utente influenza l'entità della risposta allo stile relazionale. Klein (2025), in una meta-analisi su 142 paper (N = 41.642, 800 *effect size*), quantifica l'effetto complessivo della *human-likeness* sulle risposte sociali in $g = 0.36$ [0.27, 0.44]. L'eterogeneità tra *outcome* è moderata (fiducia, soddisfazione e comportamento mostrano traiettorie distinte) e l'effetto si attenua in contesti *task*-critici ad alta complessità cognitiva; le *cues* verbali (stile colloquiale, espressioni di ascolto attivo) risultano almeno altrettanto efficaci di quelle visive (*avatar*, *emoji*), con implicazioni dirette per il design di AC esclusivamente testuale o vocale.

Il quadro non è tuttavia univocamente favorevole ad uno stile piuttosto che all'altro. Chattaraman et al. (2019), in un campione di 121 adulti anziani che interagivano con un assistente digitale per lo *shopping online*, mostrano che lo stile funzionale supera quello relazionale su variabili cognitive e prestazionali (*perceived ease of use*, *self-efficacy*, riduzione del carico informativo percepito) con effetti moderati dalla competenza digitale dell'utente: gli utenti a bassa competenza traggono maggiore beneficio funzionale dallo stile transazionale, mentre quelli ad alta competenza rispondono meglio allo stile relazionale su *outcome* sociali (fiducia affettiva, *engagement*). Il contesto limita la generalizzabilità diretta al manifatturiero, ma il *pattern* indica che la dicotomia *social* vs. *task* non rappresenta una gerarchia assoluta di efficacia: il profilo dell'utente e il tipo di *outcome* considerato sono variabili moderatrici da non trascurare.

In ambienti operativi con forte pressione temporale, i segnali socio-emotivi possono competere con il compito primario per le risorse attentive disponibili. Gilles et al. (2025), in uno studio qualitativo su piloti di aerei (N = 10) per la progettazione di assistenti virtuali di nuova generazione, rilevano che in condizioni di elevata criticità operativa l'AC deve essere conciso e funzionale, mentre nelle fasi non critiche un registro meno formale risulta accettabile. La preferenza per lo stile comunicativo non è stabile: lo stesso utente si orienta verso stili diversi nel corso della medesima sessione, in funzione del livello di criticità percepita. Il campione ridotto e il contesto aeronautico limitano la generalizzabilità; il *pattern* è tuttavia coerente con le dinamiche di carico descritte dal modello della capacità limitata (Lang, 2000), e la trasferibilità al manifatturiero, pur non diretta, potrebbe essere teoricamente plausibile.

Verhagen et al. (2022), in un esperimento *online* su *teamwork* persona-robot ($N = 72$), confrontano quattro stili comunicativi del robot: silente (*baseline*), trasparente, esplicativo e adattivo. Qualsiasi forma di comunicazione rispetto al silenzio aumenta la fiducia senza incrementare il *workload*. Sotto alta pressione temporale, lo stile adattivo (che riduce la frequenza comunicativa limitando l'*output* alle informazioni più rilevanti) migliora simultaneamente performance di team e fiducia rispetto allo stile esplicativo; lo stile meramente trasparente non produce guadagni di *situation awareness*, che emerge esclusivamente nella condizione esplicativa sotto alta interdipendenza. La manipolazione è di natura procedurale (cosa fa il robot, perché lo fa), non socio-emotiva: la domanda su come messaggi di supporto operativo e relazionale si comportino rispetto a quelli puramente funzionali sotto pressione temporale rimane aperta.

Hermanns e Teubner (2025) conducono un esperimento *online* ($N = 228$; 8.208 osservazioni) in cui i partecipanti valutano risposte AI su compiti di conoscenza a complessità variabile, con o senza vincolo di tempo. Il *framework* teorico è il modello euristico-sistematico (HSM; Chaiken, 1980, come citato in Hermanns e Teubner, 2025), secondo cui le risorse cognitive disponibili determinano se l'elaborazione avviene in modo analitico o per via di scorciatoie cognitive. Sotto pressione temporale la *performance* si deteriora significativamente ($\beta = -0.693, p < .001$), con un'asimmetria diagnostica: la capacità di rilevare risposte AI errate cala più di quella di riconoscere le risposte corrette, coerentemente con l'attivazione di processamento euristico. Un risultato controintuitivo riguarda l'interazione pressione $\times$ complessità: la pressione temporale attenua il danno della complessità di circa il 75% ($\beta = +0.374, p < .001$), interpretato dagli autori come adozione di strategie *satisficing* (strategia decisionale per cui, invece di cercare la soluzione ottimale, ci si ferma alla prima opzione "sufficientemente buona" che soddisfa una soglia minima accettabile). Sul piano teorico, l'HSM implica che le risorse cognitive disponibili possano modulare la tipologia di supporto AI preferita: in condizioni di bassa pressione il processamento sistematico consente di elaborare anche elementi socio-emotivi; sotto pressione, il processamento euristico renderebbe preferibili messaggi funzionalmente più diretti, ipotesi direttamente rilevante per il presente studio.

Kolomaznik et al. (2024), in una revisione sistematica su 108 *paper* (2018–2023) condotti in *healthcare*, *istruzione* e *customer service*, identificano cinque attributi socio-emotivi critici per la collaborazione umano-AI (*rapport*, fiducia, *user engagement*, empatia e antropomorfizzazione) e segnalano come *gap* esplicito l'assenza di studi che esaminino il contesto operativo (pressione temporale, carico cognitivo) come moderatore di tali attributi. Nessuno studio incluso nella review riguarda HRC manifatturiero o compiti fisico-cognitivi. Il presente studio si colloca in questo spazio: verifica se e come lo stile comunicativo dell'AC (R-CA vs F-CA vs assenza di AC) influenzi accettazione, fiducia e carico cognitivo in un compito di disassemblaggio con cobot, in condizioni di bassa e alta pressione temporale.

Adozione degli AC nel lavoro

La documentazione empirica sull'impatto degli AC in contesti lavorativi è ancora limitata ma in crescita. Brynjolfsson et al. (2024), in uno studio quasi-sperimentale su 5.172 agenti di supporto clienti, rilevano un aumento medio della produttività del 15,2%, con benefici proporzionalmente

maggiori per i lavoratori meno esperti (+36% nel quintile più basso): il supporto conversazionale compensa asimmetrie di competenza avvicinando le prestazioni dei meno esperti a quelle dei più esperti (*leveling effect*). Questo *pattern* si applica al manifatturiero, dove molti operatori si confrontano per la prima volta con sistemi collaborativi avanzati. Nel contesto industriale specifico, Bousdekis et al. (2025) documentano un sistema di assistente vocale testato in un ambiente produttivo reale (stabilimento Whirlpool; $N = 6$). L'usabilità vocale risultava elevata ($VUS > 80$), il carico cognitivo basso sulle dimensioni mentale e temporale del NASA-TLX, e l'accuratezza nel riconoscimento dell'intenzione pari al 95.3% su 402 turni conversazionali; una delle poche evidenze empiriche pubblicate su un AC vocale in produzione industriale reale, che ne conferma la fattibilità tecnica senza incremento di carico cognitivo. In ambito HRC, Nielsen et al. (2026) rilevano che la presenza di capacità conversazionale in un cobot *voice-enabled* modifica le percezioni degli operatori senza produrre effetti univocamente positivi: i partecipanti esprimono una preferenza per robot privi di *display* emotivi e comportamenti sociali.

Questi studi non affrontano la variabile comunicativa: nessuno manipola lo stile dell'AC, né esaminano se gli effetti documentati in contesti di servizio si trasferiscono, si invertano o si attenuano quando il lavoratore gestisce simultaneamente un'interazione fisica con un *cobot*. L'integrazione dei due sistemi in un unico ambiente operativo genera una configurazione che in letteratura è stata testata limitatamente, nella quale le dinamiche psicologiche possono comportarsi in modo qualitativamente diverso rispetto ai contesti in cui ciascun sistema è studiato separatamente. Come discusso nel paragrafo precedente, sotto alta pressione temporale uno stile comunicativo ridotto alle sole informazioni rilevanti per il compito risulta vantaggioso per performance di team e fiducia (Verhagen et al., 2022); la domanda su come messaggi di supporto relazionale si comportino rispetto a quelli puramente funzionale in queste condizioni rimane aperta.

2.3 Dinamiche cognitive e psicologiche nella collaborazione con tecnologie robotiche e conversazionali

Technology Acceptance

La *technology acceptance*, in italiano accettazione della tecnologia (Davis, 1989), viene definita come l'intenzione comportamentale di utilizzare un sistema tecnologico, determinata dalla PU (*Perceived Usefulness*) e dalla PEOU (*Perceived Ease of Use*). In ambito HRC, il costrutto è multifattoriale: oltre agli aspetti tecnici, include fattori organizzativi e individuali come ansia, formazione ricevuta ed esperienze pregresse con la tecnologia (Puspanathan et al., 2024). Il rischio fisico percepito, lo stress e la paura del cambiamento possono ridurre l'accettazione in ambienti manifatturieri (Jacob et al., 2023). La fiducia, fattore strettamente connesso all'accettazione in contesti di collaborazione fisica, è trattata nel paragrafo successivo.

Per la componente cobot, i segnali sociali producono un effetto ridotto sull'accettazione nei task critici e ad alta complessità (Klein, 2025). Per la componente AC, invece, lo stile comunicativo

modula direttamente l'accettazione. Chattaraman et al. (2019), pur in un contesto di assistenti digitali per adulti anziani, mostrano che lo stile a orientamento sociale produce maggiore fiducia affettiva ed engagement, mentre lo stile a orientamento funzionale aumenta la PEOU e la *self-efficacy*, con differenze legate al livello di competenza dell'utente. In contesti di gestione del fallimento di servizio, la comunicazione a orientamento sociale aumenta soddisfazione e fiducia attraverso la percezione di calore come mediatore (Cai et al., 2024); sebbene tale evidenza provenga dal customer service, il meccanismo è coerente con quanto osservato in interazioni umano-AC in altri domini. Una meta-analisi (Klein, 2025) stima un effetto medio dei *cues human-like* sull'accettazione degli AC, con attenuazione in condizioni di alto carico e task critici. In sistemi umano-AI, il costrutto di accettazione si allarga includendo attributi socio-emotivi quali fiducia, empatia, rapport, engagement e antropomorfizzazione (Kolomaznik et al., 2024). L'accettazione del sistema integrato cobot + AC costituisce uno degli *outcome* primari di interesse.

Fiducia

La fiducia è definita come l'attitudine di un agente a sostenere gli obiettivi individuali in una situazione caratterizzata da incertezza e vulnerabilità (Kohn et al., 2021). In ambito HRC, Hoff e Bashir (2015) ne propongono una struttura tripartita: la fiducia disposizionale riflette la propensione individuale stabile a fidarsi della tecnologia; la fiducia situazionale si forma in base alle caratteristiche specifiche del sistema e del contesto; la fiducia appresa emerge dall'esperienza diretta con il robot. Queste tre componenti si combinano dinamicamente nel corso dell'interazione e determinano il livello di supervisione che l'operatore esercita sul sistema.

La dinamica temporale della fiducia in automazione è asimmetrica: si costruisce lentamente e si perde rapidamente dopo un errore. Rieger et al. (2023) mostrano che gli utenti tollerano meno errori da sistemi automatizzati rispetto a operatori umani (un fenomeno definito *imperfect automation schema*) con conseguente erosione più rapida della fiducia. Il *framing* antropomorfo del cobot (es. attribuire un nome al robot) aumenta la fiducia iniziale ma amplifica la diffidenza post-errore; questa dinamica è documentata in una serie di studi su lavoratori di fabbrica (cfr. supra), trattando fiducia e diffidenza come costrutti separati con traiettorie evolutive distinte. Kopp (2024) propone per questi costrutti una concettualizzazione multidimensionale che integra dimensioni cognitive, affettive e comportamentali nel contesto specifico del cobot.

Per la componente AC, Wang et al. (2026) distinguono tra fiducia affettiva, generata da *cues human-like*, e fiducia cognitiva, basata sulla competenza percepita del sistema, indicando che entrambe le componenti contribuiscono alla fiducia complessiva dell'utente. Sul piano comunicativo, Verhagen et al. (2022), in un campione di N = 72 partecipanti, mostrano che uno stile trasparente ed esplicativo dell'AC aumenta la fiducia senza incrementare il *workload*; uno stile adattivo, che riduce la frequenza delle comunicazioni sotto pressione temporale, migliora ulteriormente la performance di team e la fiducia rispetto allo stile meramente esplicativo.

Percezione di autonomia e controllo sulla tecnologia

L'autonomia rappresenta un bisogno psicologico di base nel lavoro: secondo la SDT (*Self-Determination Theory*; Deci e Ryan, 1985, 2000), la percezione di agire secondo la propria volontà è condizione per la motivazione autonoma e la soddisfazione lavorativa. L'intensità di questo bisogno non è tuttavia uniforme: varia in funzione delle differenze individuali e della complessità del compito. Compiti con richieste cognitive molto basse riducono la varietà e il significato percepito del lavoro (Berretta et al., 2023); compiti di complessità estrema possono invece saturare le risorse attentive disponibili (Kahneman, 1973), comprimendo lo spazio per scelte deliberate.

Nelle organizzazioni ad alta automazione, l'autonomia percepita è sistematicamente a rischio. Quando il sistema decide l'allocazione dei compiti in modo unilaterale, i lavoratori riportano riduzioni nell'autonomia percepita e nell'identità del compito, indipendentemente dall'efficienza operativa del sistema (Tausch e Kluge, 2020). In contesti HRC, insicurezza e mancanza di controllo figurano tra gli esiti negativi più ricorrenti, anche quando il cobot opera in modo funzionalmente corretto (Baltrusch et al., 2022). L'autonomia percepita del robot amplifica queste dinamiche: all'aumentare dell'autonomia robotica percepita crescono la minaccia all'identità, alla sicurezza e alle risorse del lavoratore, con effetti più marcati in chi presenta atteggiamenti negativi preesistenti verso la tecnologia (Złotowski et al., 2017).

Quando al cobot si affianca un AC, questi costrutti diventano variabili di esito direttamente misurabili. L'introduzione di un interlocutore AI capace di prendere decisioni interattive e allocare mansioni può amplificare la percezione di riduzione dell'autonomia (Berretta et al., 2023); la modalità comunicativa con cui il sistema segnala i propri interventi può determinare se il monitoraggio viene percepito come controllante o informativo (Deci e Ryan, 2000). Mantenere autonomia e senso di controllo non è quindi una conseguenza automatica dell'efficienza tecnica del sistema, ma un requisito di progettazione dell'interazione, e una delle dimensioni su cui lo stile comunicativo dell'AC può esercitare un effetto differenziale.

Percezione di sorveglianza e monitoraggio durante il lavoro

La percezione di sorveglianza si riferisce alla consapevolezza di essere osservati e valutati nel corso dell'attività lavorativa. I suoi effetti sulla *performance* e sul benessere non sono uniformemente negativi: il *framework* CHTA (*Challenge-Hindrance-Threat Appraisal*; Chen et al., 2024) distingue tra stressori percepiti come sfida, ostacolo o minaccia. A livelli contenuti, la sorveglianza può attivare un *appraisal* di *challenge*: la consapevolezza di essere osservati incrementa la cura nell'esecuzione e la motivazione prestazionale, effetto documentato nella letteratura sulla *social facilitation* (Zajonc, 1965) e, più recentemente, negli studi sull'*electronic performance monitoring* (Bhave, 2014; Ravid et al., 2020). A livelli elevati, o quando le risorse disponibili sono percepite come insufficienti, lo stesso fattore produce *appraisal* di *hindrance* o minaccia (Lazarus e Folkman, 1984): riduce il senso di controllo, amplifica la paura dell'errore e incrementa l'arousal generale, con ricadute negative sulla qualità del lavoro.

La direzione di questo *appraisal* dipende tuttavia non solo dall'intensità della sorveglianza, ma dalle modalità con cui essa viene mediata comunicativamente. La SDT distingue il feedback controllante, che segnala al lavoratore di essere valutato in funzione di standard esterni, dal feedback informativo, che supporta l'autoregolazione senza minacciare l'autonomia (Deci e Ryan, 2000). Nel lavoro con sistemi AI, un agente che monitora la performance e ne restituisce l'esito introduce per definizione un elemento valutante: la percezione di sorveglianza dipende non solo dalla presenza del monitoraggio, ma dalle modalità con cui viene comunicato e dal grado di agentività residua che il sistema lascia al lavoratore (Tausch e Kluge, 2020).

Una distinzione operativamente rilevante separa due forme di monitoraggio. Il monitoraggio della performance, centrato sulla valutazione dell'esito come corretto o errato, è più facilmente percepito come feedback controllante nel senso della SDT (Deci e Ryan, 2000): segnala che l'individuo è oggetto di giudizio esterno e mina la motivazione autonoma. Il monitoraggio dell'avanzamento, che restituisce informazioni sullo stato del lavoro senza formulare un giudizio finale, tende a essere vissuto come supporto informativo, più adatto a preservare l'autonomia e a ridurre l'incertezza sul proprio operato. Quest'ultima forma risulta utile per lavoratori alle prime armi, che ne traggono beneficio in termini di orientamento e correzione precoce degli errori, mentre è percepita come più intrusiva da lavoratori esperti, che la valutano ridondante rispetto alle competenze già acquisite.

In presenza di un AC che esercita funzioni di monitoraggio, queste dinamiche si intersecano con l'ansia da IA (*AI Anxiety Scale*, AIAS; Wang e Wang, 2022): individui con livelli elevati di ansia tecnologica o timori legati alla sostituzione lavorativa tendono a interpretare il monitoraggio come minaccia, amplificando la risposta di *hindrance*. Sul fronte della fiducia, gli errori dell'IA producono effetti non lineari: Rieger et al. (2023) documentano che il calo di fiducia successivo a un errore automatizzato è più marcato e duraturo rispetto allo stesso errore commesso da un operatore umano, poiché dai sistemi automatizzati ci si aspetta per *default* un'affidabilità superiore. Al tempo stesso, l'errore del sistema può ridurre la percezione di onnipotenza dell'IA, attenuando i timori di sostituzione e aumentando l'accettazione in individui che percepivano la tecnologia come minacciosa.

Uno stile comunicativo adattivo, che riduce la frequenza dei messaggi sotto pressione temporale limitandosi alle informazioni più rilevanti, ottimizza *performance* e fiducia rispetto agli stili fissi (Verhagen et al., 2022). L'effetto del monitoraggio sull'esperienza del lavoratore dipende quindi non dalla quantità di *feedback* erogato, ma dalla sua qualità, pertinenza e dalla modalità comunicativa con cui viene restituito.

Stress

Lo stress è definito come la risposta dell'organismo a richieste percepite che eccedono le risorse disponibili per fronteggiarle. Gli *stressor* lavorativi si distinguono in due categorie: *challenge stressors*, percepiti come impegnativi ma superabili, che tendono ad attivare motivazione e apprendimento; e *hindrance stressors*, percepiti come ostacoli al raggiungimento degli obiettivi, associati a *distress* e riduzione delle performance (Zielonka e Rothlauf, 2021). Nel contesto della

HRC, i due profili coesistono: la novità del compito e la coordinazione con il cobot possono configurarsi come sfida (*challenge*), fungendo da fonte di motivazione; la percezione di perdita di controllo, il rischio fisico e la pressione temporale come ostacolo (*hindrance*), ostacolando la fluidità della collaborazione (Lu et al., 2022; Babamiri e Heidarimoghadam, 2024).

Lo stress non ha solo un profilo soggettivo, ma produce risposte fisiologiche documentabili: attivazione dell'asse HPA, rilascio di glucocorticoidi e modificazioni della risposta autonoma (Sandi e Pinelo-Nava, 2007). Arai et al. (2010) hanno misurato la *Skin Potential Response* (SPR) come indicatore psicofisiologico dello stress durante HRC, mostrando che parametri fisici del robot (velocità e prossimità) producono risposte di attivazione anche in assenza di eventi critici. Gervasi et al. (2024) confermano che dopo sessioni prolungate di HRC ripetitivo, gli operatori accumulano fatica psicofisica rilevabile tramite biosensori non invasivi.

Il *technostress* è la forma di stress associata all'uso o all'anticipazione dell'uso di tecnologie digitali e fisiche (Lampi et al., 2023; Srivastava et al., 2015). Vänni et al. (2019) introducono il concetto di *robostress*, *technostress* specifico per robot fisici, distinguendolo dal *technostress* da tecnologie digitali per la componente di prossimità corporea e il rischio fisico percepito. La fiducia nella tecnologia modera questa risposta: lavoratori con alta fiducia tecnologica valutano i *technostressor* come meno minacciosi e riportano minore *distress* (Zielonka, 2022).

La pressione temporale amplifica gli effetti degli *stressor* lavorativi: sotto tale condizione, le risorse cognitive si riducono e la dipendenza da supporti esterni aumenta (Lang, 2000; Rieger e Manzey, 2022), con conseguenze sulla qualità decisionale che variano in funzione dell'affidabilità del sistema di supporto. In HRC, Shirakura et al. (2022) mostrano che al variare del livello di pressione temporale variano sia il carico percepito sia la produttività dell'operatore, con un rapporto non lineare tra l'intensità del vincolo temporale e la qualità della *performance*. Rieger e Manzey (2022) precisano che l'effetto della pressione temporale sulla *performance* dipende dalle caratteristiche del sistema di supporto: un sistema ad alta affidabilità attenua il calo prestazionale, mentre uno a bassa affidabilità lo aggrava, poiché in entrambi i casi la pressione riduce la capacità di valutare criticamente il suggerimento automatico.

La relazione tra stress e performance è curvilinea: livelli moderati di attivazione ottimizzano la performance (Yerkes e Dodson, 1908, come citato in Chen et al., 2024), mentre livelli elevati, in particolare in presenza di *appraisal* di tipo *threat* e *hindrance* (Cavanaugh et al., 2000; Crawford et al., 2010, come citati in Chen et al., 2024), la deteriorano. Lu et al. (2022), in una *review* su 25 studi HRC, mostrano che stress mentale elevato si associa a ridotta consapevolezza della sicurezza e aumento della probabilità di errore. Babamiri et al. (2024) osservano che le scale generiche di stress lavorativo non catturano la specificità del costrutto in HRC, dove la componente di interazione fisica con un sistema autonomo introduce dimensioni di valutazione assenti nei contesti lavorativi tradizionali. Lo stile comunicativo dell'AC affiancato al cobot rappresenta una variabile progettuale che può modulare l'*appraisal* degli *stressor* descritti: un R-CA che segnala disponibilità e incoraggiamento può attenuare la risposta di *hindrance*, mentre un F-CA che riduce il carico informativo può preservare le risorse cognitive sotto pressione temporale.

Workload

Il carico cognitivo (*mental workload*, MWL) è un costrutto multidimensionale che riflette la fatica mentale derivante dall'esecuzione di un compito sotto determinate condizioni di domanda (Carissoli et al., 2024). Panchetti et al. (2023) lo definiscono come la proporzione tra le risorse cognitive disponibili dell'operatore e le richieste cognitive del compito: quando la domanda supera la capacità disponibile si produce *cognitive overload* (sovraccarico cognitivo), condizione che deteriora la qualità della performance e aumenta la probabilità di errore procedurale (Lu et al., 2022).

In HRC, i principali contributori al MWL sono la velocità del cobot, l'imprevedibilità delle traiettorie, l'organizzazione dei compiti e le modalità comunicative (Carissoli et al., 2024). Il processo opposto, il *cognitive offload*, si verifica quando l'operatore trasferisce parte del carico elaborativo su risorse esterne, incluso il cobot stesso, riducendo la domanda sulla memoria di lavoro. Hmedan et al. (2021) dimostrano empiricamente questa dinamica: un sistema HRC che adatta la sequenza operativa all'ordine scelto dal lavoratore aumenta la performance complessiva senza incrementare il carico cognitivo, perché il cobot assorbe la componente di pianificazione e libera risorse per l'esecuzione.

Il *workload* elevato compromette la memoria di lavoro: risorse cognitive impegnate nell'elaborazione del compito e nel monitoraggio del *cobot* riducono la capacità di *encoding* e recupero delle informazioni (Sandi e Pinelo-Nava, 2007). In condizioni di *dual-task* nell'HRC (Eesee et al., 2024), le risorse attentive dell'operatore, per definizione limitate (Kahneman, 1973), vengono saturate, esponendo la memoria procedurale a interferenze che generano errori nella sequenza operativa (Pluchino et al., 2023). La correzione dell'errore consuma ulteriori risorse cognitive già ridotte, innescando la spirale di perdita descritta dal modello *Conservation of Resources* (COR; Hobfoll, 1989): quando le risorse percepite risultano inferiori alle domande ambientali, l'operatore attiva un *appraisal* di minaccia (Lazarus e Folkman, 1984) che si traduce in un aumento del rischio percepito nell'ambiente di collaborazione (Lu et al., 2022; Molan et al., 2025).

Il MWL è stato misurato estensivamente nei contesti di HRC industriale (Prati et al., 2023; Bassi et al., 2025) e, separatamente, negli studi sugli agenti conversazionali (Schmidhuber et al., 2021; Kontogiorgos e Gustafson, 2021), dove emerge che lo stile della risposta dell'agente ne modula l'entità. A nostra conoscenza, nessuno studio ha ancora esaminato come lo stile comunicativo di un agente AI di supporto influenzi il *workload* del lavoratore all'interno di un contesto reale di collaborazione umano-cobot.

Processi attentivi

Il costrutto di *engagement* attentivo si distingue dal *workload* per il referente teorico e per la direzione della relazione con la domanda cognitiva. Mentre il *workload* quantifica il divario tra richieste del compito e risorse disponibili, l'*engagement* attentivo descrive l'orientamento e il mantenimento dell'attenzione verso il compito (Kahneman, 1973), mediato dall'attività

dopaminergica frontostriatale (Maffei e Angrilli, 2019). I due costrutti possono coesistere o dissociarsi: un *task* ad alta domanda può produrre *workload* elevato e al tempo stesso *engagement* alto se l'operatore è pienamente orientato verso il compito; un *task* a domanda moderata ma a bassa salienza percepita può produrre basso *engagement* con *workload* contenuto.

Il *blink rate* spontaneo (frequenza di ammiccamento, ammiccamenti/min) è un indice fisiologico non invasivo dei processi attentivi dell'operatore (Maffei e Angrilli, 2018; Ranti et al., 2020). Il sistema dopaminergico inibisce il riflesso di ammiccamento in modo inconscio e adattivo quando l'attenzione è orientata verso uno stimolo saliente, minimizzando la perdita di informazione visiva (Maffei e Angrilli, 2018); la soppressione è proporzionale alla salienza percepita dello stimolo (Ranti et al., 2020). Maffei e Angrilli (2018), con il *Mackworth Clock Test* a tre livelli di difficoltà, hanno mostrato che i due effetti agiscono in modo indipendente: la difficoltà produce una soppressione sostenuta del *blink rate*, mentre il tempo trascorso genera un incremento nelle condizioni a domanda moderata e bassa, pattern replicato su materiale filmico (Maffei e Angrilli, 2019). In generale, una riduzione del *blink rate* è associata a un maggiore orientamento attentivo verso il compito, mentre un incremento progressivo è più frequentemente interpretato come segnale di fatica attentiva cumulativa. La misura è soggetta a variabilità interindividuale (range basale: 10–25 ammiccamenti/min negli adulti; Maffei e Angrilli, 2018), ma risulta utile per cogliere cambiamenti dinamici nello stato dell'operatore durante l'interazione (Urasaki e Niitsuma, 2021).

Nel contesto dell'HRC con sistemi di supporto AI, i processi attentivi mostrano un'ulteriore fonte di variabilità. In HRC prolungata, l'orientamento attentivo non è stabile ma si modifica per effetto di pratica, fatica e adattamento strategico nel corso dei blocchi (Saleem et al., 2026); il *blink rate* mostra in tali contesti evidenza preliminare di validità convergente con indici EEG di orientamento al compito, in linea con la sua applicabilità in ambienti HRI realistici (Saleem et al., 2026). L'incremento del *blink rate post-task* rispetto al *pre-task* è stato rilevato in condizioni di HRC con *timing* robotico non ottimale, a indicare fatica mentale indotta dall'interazione (Urasaki e Niitsuma, 2021). Nel presente studio, l'orientamento attentivo dell'operatore è suscettibile di essere influenzato sia dalla componente operativa (il disassemblaggio con il cobot in condizioni di pressione temporale variabile) sia dalla presenza e dallo stile del sistema di supporto conversazionale; a nostra conoscenza, nessuno studio ha esaminato tale effetto nel contesto dell'HRC industriale, rendendo il *blink rate* un *outcome* fisiologico di interesse per RQ3.

3. DOMANDE DI RICERCA E IPOTESI

RQ1: Manipulation check

RQ1.1: In che misura la manipolazione della pressione temporale (alta vs. bassa) influenza i livelli di **distress** e di **domanda temporale** percepita durante il compito?

H1.1. La condizione di alta pressione temporale dovrebbe associarsi, in media, a livelli più elevati di distress (DSSQ) e di domanda temporale percepita (NASA-TLX, sottoscala temporal demand) rispetto alla condizione di bassa pressione temporale.

RQ1.2: In che modo le condizioni di supporto (R-CA, F-CA, No-CA) influenzano la percezione delle **caratteristiche antropomorfe e socio-affettive del sistema**?

H1.2. Le condizioni con agente conversazionale (R-CA e F-CA) dovrebbero associarsi, in media, a livelli più elevati di antropomorfismo, animazione e simpatia (Godspeed) rispetto alla condizione No-CA. In particolare, la condizione R-CA dovrebbe essere percepita come più antropomorfa, animata e simpatica rispetto alla F-CA, riflettendo il maggiore grado di umanizzazione del suo stile comunicativo.

RQ2: Percezione soggettiva del supporto del CA

RQ2.1: In che modo le diverse condizioni di supporto (R-CA, F-CA, No-CA) e la pressione temporale influenzano l'esperienza soggettiva del compito in termini di **accettazione e fiducia**?

H2.1. Le condizioni con supporto CA dovrebbero associarsi, in media, a livelli più elevati di accettazione e fiducia rispetto alla condizione No-CA. In particolare, in condizioni di bassa pressione temporale ci si attende che la condizione R-CA sia percepita più favorevolmente della F-CA, con la condizione No-CA come riferimento più basso (Klein, 2025; Chen et al., 2025). In condizioni di alta pressione temporale, ci si attende che F-CA risulti più adeguata di R-CA, con No-CA come riferimento meno favorevole (Hermanns e Teubner, 2025).

RQ2.2: Come si associano le diverse condizioni di supporto AC e la pressione temporale alla **percezione di controllo** sul compito, di **autonomia decisionale** e di **sorveglianza**?

Questa domanda viene affrontata in chiave esplorativa, senza formulare ipotesi direzionali a priori. La scelta di includere queste variabili risponde alla loro rilevanza teorica e applicativa: la letteratura indica che l'introduzione di un sistema AI in grado di monitorare la performance e fornire feedback valutativi può modificare la percezione di controllo e autonomia del lavoratore e amplificare la sensazione di essere osservato (Tausch e Kluge, 2020; Berretta et al., 2023; Stock-Homburg e Nguyen, 2023), con effetti che possono variare in funzione dello stile comunicativo del sistema e del livello di pressione temporale.

RQ3: Dinamiche cognitive di workload, attenzione e affaticamento

RQ3: In che modo le condizioni di supporto AC e la pressione temporale influenzano le dinamiche cognitive del compito — in termini di **carico cognitivo percepito** (NASA-TLX), **attivazione fisiologica** (pupillometria) e **indici oculari di attenzione e affaticamento** (blink rate)?

H3. Ci si attende che il carico cognitivo risulti più elevato in condizioni di alta pressione temporale rispetto alla bassa pressione, sia sul piano soggettivo (NASA-TLX) sia su quello fisiologico (dilatazione pupillare). Le differenze tra le condizioni di stile comunicativo dell'AC potrebbero invece variare in funzione del contesto. In particolare, sotto alta pressione temporale, la condizione R-CA potrebbe associarsi, rispetto a F-CA e No-CA, a un carico cognitivo più contenuto e a pattern oculari compatibili con una maggiore stabilità attentiva e un minore affaticamento (riflessi nel *blink rate*); in assenza di pressione, tali differenze potrebbero attenuarsi o assumere configurazioni diverse.

4. METODI

4.1 Campione sperimentale

Il campione era composto da 24 partecipanti (12 uomini, 11 donne; età: $M = 29.3$, $DS = 5.0$, range = 20–40). Le caratteristiche socio-demografiche complete (titolo di studio, occupazione e familiarità con AI e cobot) sono riportate nella Tabella 4.1. Un partecipante è stato escluso dalle analisi per perdita dei dati fisiologici in una condizione sperimentale (N analizzato = 23). La quasi totalità del campione aveva sentito parlare di sistemi di intelligenza artificiale (95.8%) e la maggioranza li aveva utilizzati (83.3%), mentre la familiarità con i robot collaborativi era più limitata (29.2% a conoscenza dei cobot; 8.3% con esperienza diretta). Lo studio è stato approvato dal Comitato Etico per la Ricerca in Psicologia dell'Università di Padova (protocollo n. 2026_287R1).

Variabile	Categoria	n	%
Genere			
	Uomini	12	50.0
	Donne	11	45.8
	Preferisce non specificare	1	4.2
Istruzione			
	Laurea magistrale / ciclo unico	7	29.2
	Laurea triennale	7	29.2
	Scuola secondaria II grado	7	29.2
	Scuola secondaria I grado	3	12.5
Occupazione			
	Lavoratore	13	54.2
	Studente-lavoratore	6	25.0
	Studente	3	12.5
	Altro	2	8.3
Familiarità con AI			
	Sentito parlare di sistemi AI	23	95.8
	Ha utilizzato sistemi AI	20	83.3
Familiarità con cobot			
	Sentito parlare di cobot	7	29.2
	Esperienza diretta	2	8.3

Tabella 4.1. Caratteristiche socio demografiche del campione

4.2 Setup tecnico e strumenti

Il setup sperimentale è stato progettato per garantire il controllo delle variabili ambientali e la replicabilità della procedura, integrando una *workstation* fisica di interazione persona-robot e un'infrastruttura digitale per la gestione del compito e la raccolta dati. L'esperimento è stato condotto in una stanza dei laboratori dell'Università di Padova, in Galleria Spagna 28. Lo spazio è stato organizzato per consentire un'interazione libera con gli strumenti, senza vincoli fisici rilevanti. La postazione di lavoro includeva un tavolo operativo su cui era collocato un braccio robotico collaborativo UR10e, con 6 gradi di libertà e capacità di carico fino a 12,5 kg, controllato tramite software Polyscope (versione 5.11). Il cobot era integrato in un tavolo di lavoro ad altezza regolabile. Alla destra del tavolo era posizionato uno schermo collegato via HDMI a un computer *MacBook Pro* 13-inch (2020, Apple M1, 16 GB RAM, *macOS Sonoma* 14.6.1), utilizzato per l'esecuzione dell'interfaccia sperimentale. Un microfono a condensatore, collegato a una scheda audio *Behringer U-Phoria* UMC202HD, era posizionato sopra l'area di lavoro per la registrazione dei comandi vocali. L'output audio veniva gestito da due sistemi distinti: una cassa *Sony SRS-XB 32* per i suoni dell'interfaccia e una *Bose SoundLink Revolve* per la riproduzione di suoni ambientali industriali in loop. L'area di lavoro era suddivisa in quattro zone funzionali: un'area del cobot con i contenitori per i copriviti, un'area centrale dedicata al partecipante con un contenitore delle viti, un'area di smontaggio delle placche ed un'area a destra dove erano posizionate le placche difettose da smontare durante il turno di lavoro.



Figura 4.1 Postazione collaborativa e schermo con interfaccia.

Per la rilevazione delle misure fisiologiche, i partecipanti hanno indossato un dispositivo non invasivo: un eye tracker binoculare Pupil Labs Neon per la registrazione del movimento oculare, del

diametro pupillare e del blink rate. I flussi di dati fisiologici, i marker dell'interfaccia sperimentale e gli eventi comportamentali erano sincronizzati tramite il protocollo Lab Streaming Layer (LSL), che garantisce l'allineamento temporale di stream eterogenei su una timeline comune.

4.3 Design sperimentale

Lo studio adotta un disegno *within-subjects* con due variabili indipendenti manipolate entro i partecipanti: il tipo di supporto del sistema (No-CA, F-CA, R-CA) e il livello di pressione temporale (bassa, alta). Il disegno $3 \text{ (tipologia di supporto)} \times 2 \text{ (pressione temporale)}$ produce sei condizioni sperimentali, ciascuna corrispondente a un turno di lavoro. Ogni partecipante ha completato tutte e sei le condizioni. Per controllare gli effetti d'ordine, le sei condizioni sono state somministrate secondo dodici sequenze costruite con un quadrato latino, che combinano l'ordine dei tre tipi di supporto con i due livelli di pressione temporale. L'ordine dei livelli di pressione (bassa $\rightarrow$ alta o alta $\rightarrow$ bassa) era controbalanciato tra i partecipanti, ma mantenuto fisso all'interno della sessione, per evitare effetti di *carry-over* dell'induzione di stress (Keppel e Wickens, 2004) sulle condizioni di bassa pressione temporale.

Tipologia di supporto AC

La prima variabile indipendente è il tipo di supporto fornito dal sistema durante il turno di lavoro, implementato attraverso tre modalità distinte (Tabella 4.2):

- **No-CA**: il partecipante riceve esclusivamente le istruzioni operative necessarie allo svolgimento del compito (una per coppia di placche) e il feedback sulla prestazione al termine di ogni blocco, senza alcuna forma di supporto aggiuntivo. Questa condizione funge da baseline per il confronto con le condizioni sperimentali.
- **Agente Conversazionale Funzionale (F-CA)**: il sistema (denominato Tk.2) fornisce supporto di carattere funzionale e informativo. I messaggi sono formulati in modo impersonale e orientato al compito, con uso della seconda persona singolare riferita alle azioni del partecipante, al fine di mantenere un registro esclusivamente procedurale. I suggerimenti intermedi riportano il tempo trascorso nelle condizioni di bassa pressione e il tempo residuo in quelle di alta pressione. Questa caratterizzazione si basa sulla letteratura sui sistemi funzionali, che evidenzia uno stile comunicativo formale, impersonale e orientato all'efficienza e al raggiungimento dell'obiettivo, con focus su contenuti informativi e dominio chiuso (Shum et al., 2018; Xu et al., 2022; Wang et al., 2023; Cai et al., 2024; Klein, 2025).
- **Agente Conversazionale Relazionale (R-CA)**: il sistema (denominato Helen) fornisce supporto di carattere relazionale e socio-emotivo. I messaggi utilizzano la prima persona plurale per rimarcare la dimensione collaborativa del compito e adattano il registro linguistico alla condizione di pressione temporale: nelle condizioni di bassa pressione enfatizzano il ritmo sostenibile del lavoro, mentre nelle condizioni di alta pressione riconoscono l'urgenza operativa e offrono incoraggiamento. Questa caratterizzazione deriva

dalla letteratura sugli agenti *relational* e *human-like*, che sottolinea l'importanza di un linguaggio relazionale, empatico e collaborativo, capace di supportare i bisogni socio-emotivi e favorire la connessione con l'utente (Chattaraman et al., 2019; Allouch et al., 2021; Wang et al., 2023; Kolomaznik et al., 2024; Chen et al., 2025; Wang et al., 2026).

Condizione	Esempi di messaggi (blocco 1, bassa pressione)	Sistema (nome)	Registro comunicativo e voce TTS	Msg per blocco (n)
No-CA	Istruzione «Smonta le placche 1 e 2, conta il numero di coprivi blu e rossi. Alla fine, pronuncia il numero e il rispettivo colore ad alta voce.» Esito «Esito: corretto.»	—	— (testo a schermo) —	2
F-CA	Intro (× 1 a sessione) «Sistema di supporto alla produzione Tk.2: attivato.» Istruzione «Smonta le placche 1 e 2. Conta i coprivi blu; io conterò quelli rossi. Alla fine, comunica il numero di coprivi blu ad alta voce.» Suggerimento tempo (min. 2 dopo istruzione) «Tempo trascorso: circa 2 minuti.» Feedback «Risposta corretta. Rimossi 8 coprivi blu e 6 coprivi rossi.»	Tk.2	Impersonale, procedurale 2 ^a persona singolare <i>it-IT-IsabellaNeural</i>	3
R-CA	Intro (× 1 a sessione) «Ciao, sono Helen! Sono qui per lavorare insieme a te nello smontaggio delle placche difettose. Durante il turno, ti darò supporto tenendo traccia di alcuni conteggi, così potrai concentrarti sul lavoro senza appesantirti troppo.» Istruzione «Iniziamo smontando le placche 1 e 2. Tu conta i coprivi blu, io mi occuperò dei rossi. Quando hai finito, dimmi a voce quanti coprivi blu hai trovato.» Suggerimento (min. 2 dopo istruzione) «Continua così: individua e mantieni il ritmo che per te è più sostenibile.» Feedback «Risposta corretta! Da queste due placche difettose, abbiamo recuperato 8 coprivi blu e 6 coprivi rossi. Mi sembra di lavorare bene insieme.»	Helen	Relazionale, collaborativo 1 ^a persona plurale <i>it-IT-ElsaNeural</i>	3

Tabella 4.2. Caratteristiche operative delle tre condizioni di supporto AC.

Entrambe le condizioni AC presentavano lo stesso numero di interventi per blocco, secondo una struttura fissa: istruzione di apertura, suggerimento in corso d'opera, *feedback* sul conteggio in chiusura. Tutte le condizioni sono implementate tramite un paradigma *Wizard-of-Oz* (WoZ; Dahlbäck et al., 1993). I partecipanti interagiscono con quello che credono essere un sistema AI automatizzato, mentre le risposte dell'agente (file audio pre-registrati) vengono attivate in tempo reale dal ricercatore tramite l'interfaccia sperimentale, in risposta agli input vocali del partecipante rilevati da un sistema di riconoscimento vocale locale (VOSK, modello italiano). Il numero di copriviti dei colori bersaglio era predefinito e costante tra partecipanti per ciascun turno, in modo che il ricercatore conoscesse a priori la risposta corretta. In caso di mancato o errato riconoscimento vocale del conteggio dichiarato dal partecipante, il ricercatore inseriva manualmente l'esito (corretto/errato) tramite l'interfaccia di controllo. Questo approccio ha permesso di standardizzare il contenuto delle risposte tra i partecipanti, mantenendo al tempo stesso la verosimiglianza dell'interazione. I partecipanti vengono informati del paradigma WoZ al termine dell'ultima sessione, durante il *debriefing*.

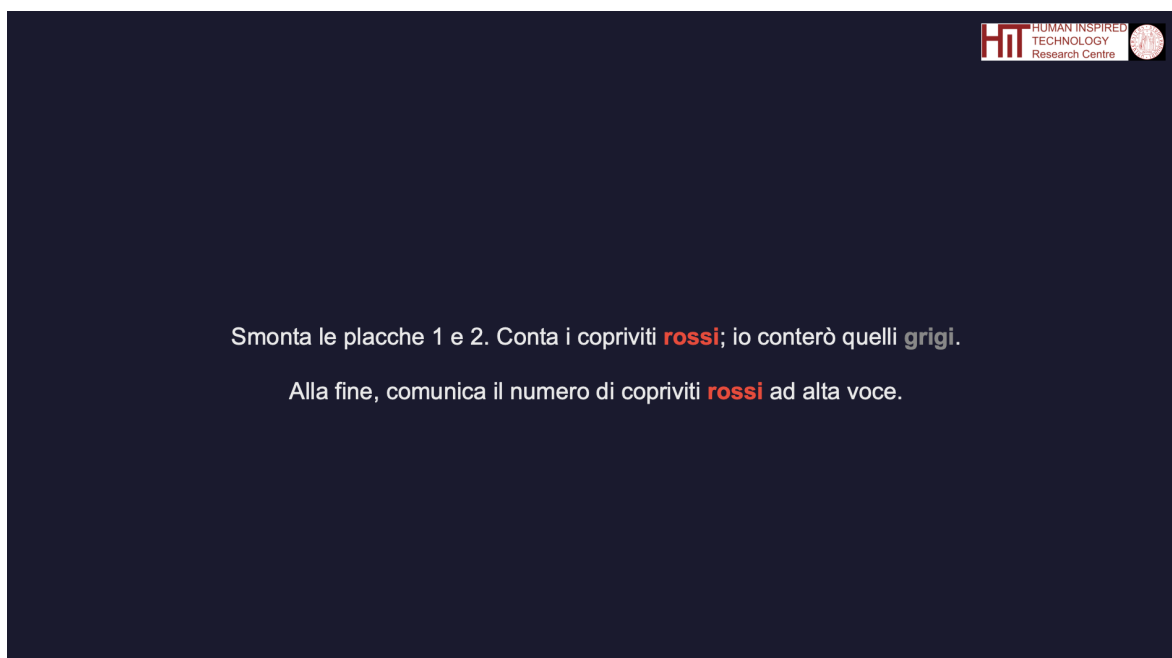


Figura 4.2. Indicazione a schermo dell'interfaccia per blocco 1; placche 1 e 2, condizione F-CA, bassa pressione temporale.

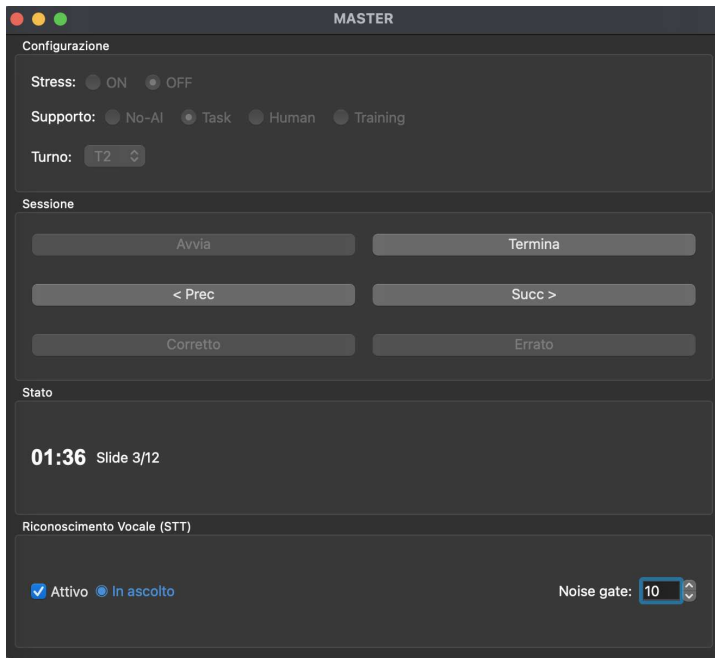


Figura 4.3. *Interfaccia sperimentatore: i tasti ‘Sessione’ servivano all’occorrenza quando riconoscimento verbale STT non riconosceva l’input.*

Pressione temporale

La seconda variabile indipendente è il livello di pressione temporale, manipolato attraverso un vincolo di urgenza operativa. La letteratura identifica la pressione temporale come *stressor* operativo primario nei contesti di lavoro con automazione (Lazarus e Folkman, 1984; Rieger e Manzey, 2022; Shirakura et al., 2022): la percezione di un vincolo temporale visivamente esplicito è sufficiente a produrre aumenti misurabili di *distress* e degrado della funzione esecutiva anche in assenza di sanzioni oggettive (Sussman e Sekuler, 2022).

Nelle condizioni di alta pressione temporale, all'inizio del turno compare un *banner* con la notifica di un ritardo produttivo; per tutta la durata del turno sono presenti un *timer* circolare a conto alla rovescia (un arco rosso si consumava in senso orario, con il tempo rimanente in formato MM:SS al centro) e un segnale acustico ripetuto a carattere percussivo (50 BPM). Queste componenti mirano a simulare le condizioni di urgenza tipiche dell'ambiente manifatturiero, rendendo la pressione temporale saliente su canali visivi e uditivi e ancorandola a un contesto operativo ecologicamente plausibile. Nelle condizioni di bassa pressione temporale nessun vincolo visivo o uditivo è presente: i partecipanti completano il turno senza indicazioni di urgenza operativa.

4.4 Compito sperimentale

Il compito sperimentale simula un'operazione di disassemblaggio di placche difettose in un contesto manifatturiero. I partecipanti smontano sei placche difettose per turno, organizzate in tre coppie (blocchi). Ogni placca era un parallelepipedo a base quadrata con 25 o 36 fori, su cui erano montate viti con copriviti colorati (blu, rossi, verdi, grigi) posizionati sulla testa di ciascuna vite. Il compito di smontaggio consiste nello sfilare i copriviti dalla testa della vite, senza l'uso di strumenti, e nel riporli nei contenitori appositi (un contenitore per colore). Dopo aver fatto ciò, deve quindi smontare le viti della placca, sempre sfilandole a mano, e mettere la placca ormai vuota in un'area specifica alla sua sinistra. Ai partecipanti viene comunicato che l'azienda stava testando due nuovi sistemi intelligenti di supporto ai lavoratori (Helen e Tk.2); in alcuni turni avrebbero potuto ricevere messaggi vocali dal sistema. Questa informazione è funzionale al mantenimento della verosimiglianza del paradigma WoZ senza rivelare il vero obiettivo dello studio.

Modalità di presentazione dei contenitori

Il cobot UR10e presenta i contenitori secondo un ciclo che copriva tutti e quattro i colori, uno per volta, in ordine randomizzato (es. rosso → blu → verde → grigio; grigio → verde → blu → rosso; ecc.). Il partecipante attiva il ciclo del cobot, con una pressione sull'effettore a pinza. A quel punto il cobot preleva il primo contenitore di copriviti, lo sposta davanti al partecipante, tenendolo sollevato dal banco di lavoro di circa 10 centimetri e leggermente inclinato per facilitare l'inserimento dei copriviti.

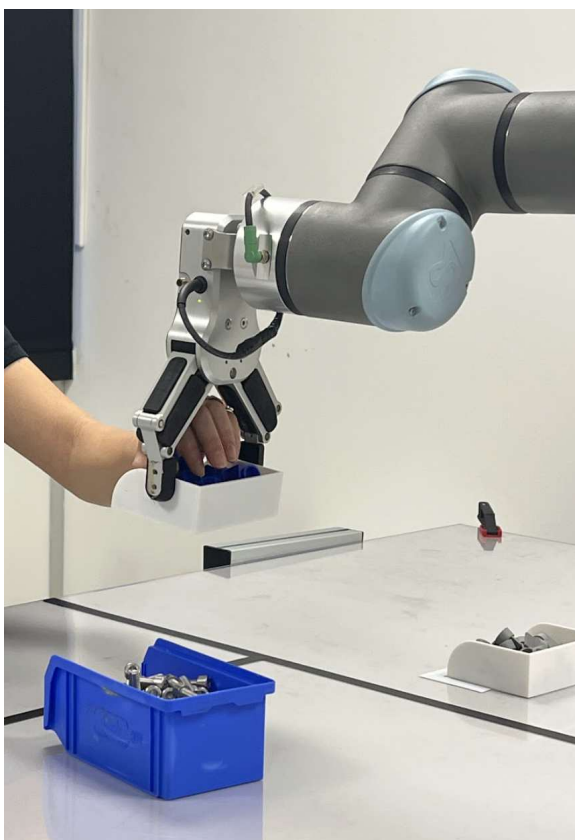
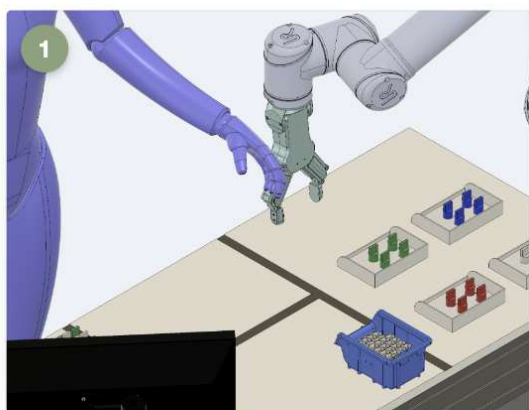


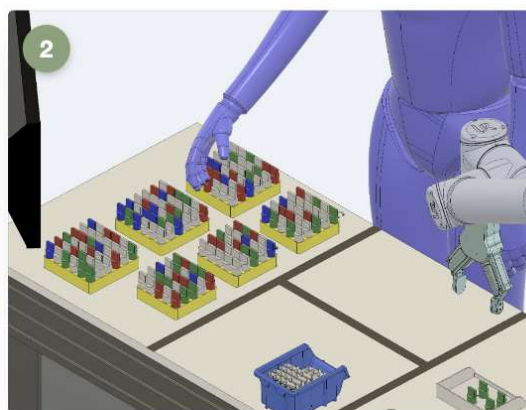
Foto 4.1. *Il cobot UR10e presenta il contenitore dei copriviti al partecipante, mantenendolo sollevato e inclinato per facilitare l'inserimento.*

A questo punto il cobot rimane fermo, aspettando una pressione del partecipante sull'effettore per adagiare il contenitore al suo posto e provvedere a prendere un altro da porgere al partecipante. Questa sequenza si ripete per 4 volte, una per colore, dopo di che il ciclo termina con il cobot che ritorna in posizione di partenza, ad attendere la pressione sulla pinza per ripartire. Ogni contenitore segnala il colore di competenza poiché conteneva già alcuni copriviti di quel colore; il partecipante riceve dalla placca tutti i copriviti del colore corrispondente e li ripone nel contenitore, quindi esercita una pressione sull'effettore per ricevere il contenitore successivo. Il ciclo si ripete per tutte le placche da smontare.



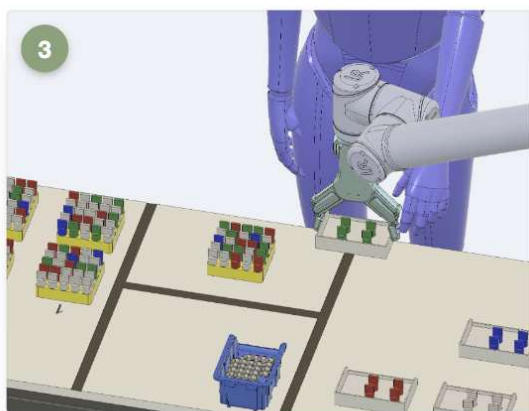
Avvio del ciclo

Il partecipante preme la pinza del cobot per avviare il ciclo.



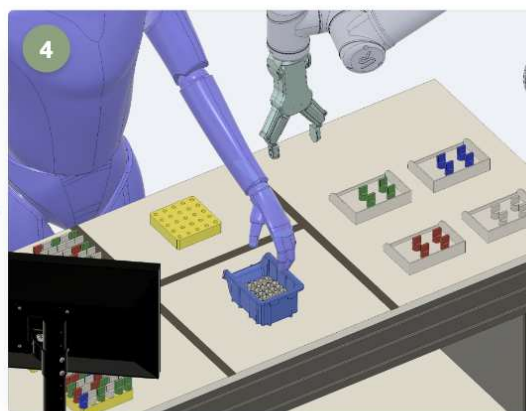
Posizionamento placca

Il partecipante posiziona la placca sul piano di lavoro e smonta i copriviti del colore indicato dal contenitore portato dal cobot.



Cambio contenitore

Dopo aver riposto i copriviti, il partecipante preme l'effettore: il cobot ripone il contenitore e ne preleva uno nuovo.



Completamento ciclo

Il partecipante completa il ciclo rimuovendo le viti dalla placca, dopodiché prende la placca successiva e ricomincia dall'avvio.

Figura 4.4. Ciclo lavorativo. L'immagine mostra l'interazione del partecipante con il cobot.

Compito di conteggio dei pezzi smontati

Per rendere il compito più ecologicamente valido e introdurre una componente cognitiva accessoria, è stato affiancato allo smontaggio un breve compito di conteggio. Per ogni blocco (due placche) sono assegnati due colori target da rilevare e contare durante la rimozione dei copriviti. Questi colori target cambiano ad ogni blocco. Al termine del blocco, il partecipante comunica verbalmente i propri conteggi (es. "dieci blu, sei rossi"). L'interfaccia rileva la risposta tramite riconoscimento vocale locale (VOSK, modello italiano; vocabolario vincolato a numeri e colori rilevanti); nei casi

di mancato riconoscimento, il ricercatore inseriva il valore manualmente tramite l'interfaccia WoZ. Il sistema conferma la correttezza del conteggio, fornendo un feedback immediato sull'esito (corretto o errato).

La struttura del compito variava in funzione della condizione sperimentale. In condizione No-CA, il partecipante è responsabile di entrambi i colori target e comunica entrambi i valori al termine del blocco. Nelle condizioni F-CA e R-CA, l'AC (attraverso WoZ) gestiva il conteggio di uno dei due colori; il partecipante contava il secondo e comunica unicamente quel valore.

Induzione della pressione temporale

Nelle condizioni di bassa pressione temporale non è presente alcun timer. Nelle condizioni di alta pressione temporale, all'inizio del turno compare un *banner* con la notifica di un ritardo produttivo; per tutta la durata del turno è visibile sullo schermo un timer circolare con sfondo scuro, accompagnato da un segnale acustico ripetuto a carattere percussivo: un arco rosso si consumava in senso orario a partire dalla sommità del cerchio proporzionalmente al tempo trascorso, mentre il tempo rimanente in formato MM:SS era mostrato al centro in cifre rosse (Figura 4.5).

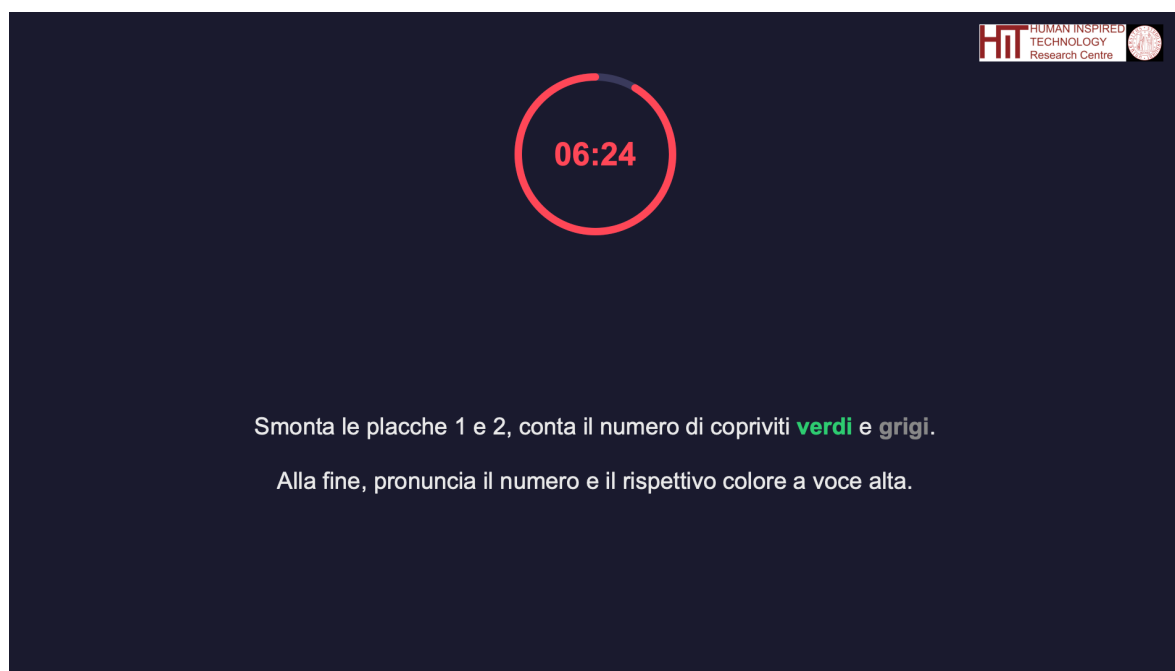


Figura 4.5. Indicazione a schermo interfaccia per blocco 1 (placche 1 e 2) condizione No-CA, alta pressione temporale. Si noti il timer in rosso.

4.5 Procedura sperimentale

La procedura sperimentale si è articolata in quattro fasi principali: preparazione, training e sessione sperimentale, intervista finale e *debriefing*.

Preparazione

All'arrivo in laboratorio, i partecipanti hanno firmato il consenso informato e indossato il dispositivo di misurazione fisiologica (Eye tracker Pupil Labs Neon per registrazione oculare e pupillometria). È stata fornita una breve introduzione generale allo studio e al funzionamento della postazione.

Training

Successivamente, i partecipanti hanno completato una fase di training strutturata, durante la quale venivano illustrati il compito di disassemblaggio, il funzionamento del cobot e le modalità di interazione con l'interfaccia sperimentale. Il training includeva una prova pratica per familiarizzare con la postazione prima dell'avvio della sessione vera e propria. Alla fine del training hanno risposto ad una domanda qualitativa, posta a voce, sulle loro opinioni riguardo ai sistemi di supporto intelligenti.

Sessione sperimentale

La sessione sperimentale era composta da sei turni di lavoro, ciascuno corrispondente a una condizione sperimentale. Prima del primo turno, i partecipanti compilano il questionario *Dundee Stress State Questionnaire (DSSQ)* (Matthews et al., 2002) in versione pre-sessione. Ogni turno consisteva in un compito di disassemblaggio collaborativo con il cobot UR10e. Nelle condizioni F-CA e R-CA, l'agente forniva istruzioni vocali e *feedback* in tempo reale tramite l'interfaccia WoZ; nella condizione No-CA il partecipante operava autonomamente senza supporto aggiuntivo, se non le indicazioni operative a schermo dell'interfaccia. La pressione temporale era indotta tramite un timer circolare a conto alla rovescia ed un segnale sonoro associato. Al termine di ogni turno i partecipanti compilano i questionari post-condizione: *DSSQ* (Matthews et al., 2002), *TAM* (Davis, 1989), *S-TIAS* (Körber, 2019), *Godspeed scale* (Bartneck, 2009) e la scala ad hoc su autonomia, controllo e sorveglianza.

Intervista finale e debriefing

Al termine dell'ultimo turno viene svolta un'intervista semistrutturata finale. La procedura si conclude con la rimozione dell'*eye tracker* e un *debriefing* completo sull'obiettivo dello studio e sul paradigma Wizard-of-Oz utilizzato.

4.6 Misure

Le misure utilizzate nello studio sono state selezionate sulla base della letteratura sull'interazione umano-IA e sui fattori che influenzano benessere, performance e accettazione nei contesti di Human-Robot Collaboration.

Misure self-report

Manipulation check: stress e percezione di antropomorfismo

La manipolazione della pressione temporale è stata verificata tramite due misure. Il *Dundee Stress State Questionnaire* (DSSQ; Matthews et al., 2002), somministrato prima e dopo ogni turno, ha rilevato le variazioni nel vissuto soggettivo di distress e worry indotte dalla manipolazione. La sottoscala *temporal demand* del NASA-TLX ha inoltre fornito un indice soggettivo aggiuntivo della percezione di urgenza temporale durante il turno.

La manipolazione del tipo di supporto AC è stata verificata tramite la *Godspeed Scale*: la condizione R-CA avrebbe dovuto produrre punteggi più elevati su antropomorfismo e simpatia rispetto a F-CA.

Fiducia

La fiducia verso il sistema è stata rilevata tramite la *Short Scale of Trust in Automated Systems* (S-TiAS; Körber, 2019), strumento sviluppato per misurare la tendenza dell'utente a fidarsi di sistemi automatizzati e intelligenti durante l'interazione con essi. Nello studio sono stati somministrati gli item della scala, valutati su scala Likert a 7 punti (1 = Non sono affatto d'accordo, 7 = Sono completamente d'accordo); il punteggio è calcolato come media degli item, con valori più alti indicativi di maggiore fiducia percepita.

Accettazione

L'accettazione è stata misurata tramite il questionario della *Technology Acceptance Model* (TAM; Davis, 1989), nelle sue tre dimensioni: utilità percepita (*perceived usefulness*), facilità d'uso percepita (*perceived ease of use*) e atteggiamento verso l'uso (*attitude toward using*). Ciascuna sottoscala comprende 4 item valutati su scala Likert a 7 punti (1 = Non sono affatto d'accordo, 7 = Sono completamente d'accordo); il punteggio di ciascuna dimensione è calcolato come media degli item corrispondenti. La scala era somministrata al termine di ogni turno sperimentale.

Controllo percepito, autonomia e sorveglianza percepita

La percezione di controllo sul compito, di autonomia decisionale e di sorveglianza da parte del sistema è stata rilevata tramite tre item singoli formulati ad hoc, somministrati al termine di ogni

turno contestualmente alle scale TAM e S-TiAS. Gli item erano preceduti dalla consegna: "Ripensando alla collaborazione con questo sistema intelligente..." e valutati su scala a 7 punti (1 = Poco, 7 = Molto):

1. Ho percepito di avere controllo della situazione.
2. Ho percepito di avere autonomia nel lavoro.
3. Mi sono sentito sorvegliato.

In assenza di strumenti validati specifici per questo contesto, gli item sono stati formulati ad hoc per rilevare tre dimensioni teoricamente rilevanti: controllo situazionale, autonomia percepita nel lavoro e percezione di monitoraggio. Data la natura esplorativa della rilevazione, i risultati sono interpretati come indicativi di *pattern* e non come test inferenziali.

Workload

Il carico cognitivo è stato operazionalizzato su due livelli di rilevazione. Il carico esplicito è stato misurato al termine di ogni turno tramite il NASA Task Load Index (NASA-TLX; Hart e Staveland, 1988), strumento multidimensionale che valuta il carico percepito su sei sottoscale: domanda mentale, domanda fisica, domanda temporale, performance, sforzo e frustrazione. Il carico implicito è stato rilevato in continuo durante il compito tramite il diametro pupillare registrato con l'eye tracker Pupil Labs Neon: la misura estratta è la variazione media del diametro pupillare rispetto alla baseline individuale (mean pupil, in mm), calcolata per blocco dopo correzione della baseline e time-warping del segnale (cfr. §4.7 per la pipeline di preprocessing). La dilatazione pupillare corretta per baseline è considerata indice fisiologico non invasivo del workload mentale (vedi paragrafo in "Misure di eye-tracking").

Misure di eye-tracking

Blink

Il *blink rate* spontaneo (frequenza di ammiccamento, ammiccamenti/min) è stato utilizzato come indice fisiologico non invasivo delle dinamiche di carico attentivo e affaticamento durante il compito. La letteratura mostra che questo parametro è sensibile sia all'attenzione sostenuta sia alla fatica mentale legata al *time-on-task*, e viene interpretato come indicatore di variazioni nell'allocazione attentiva (Maffei e Angrilli, 2018; Urasaki e Niitsuma, 2021). Il segnale è stato estratto in continuo dal segnale oculometrico del Pupil Labs Neon tramite il canale di apertura palpebrale (*eyelid aperture*, mm), registrato binocularmente. Un evento *blink* è stato definito come qualsiasi campione in cui l'apertura palpebrale di almeno un occhio scendeva sotto la soglia di 2.0 mm, derivata dalla documentazione tecnica del sensore (Pupil Labs, 2023). Sono stati applicati due filtri di pulizia del segnale: un filtro sulla durata dell'evento (50–500 ms, corrispondente al range fisiologico del *blink* spontaneo negli adulti; Schiffman, 2001) e un filtro sull'intervallo *inter-blink* (IBI $\geq$ 150 ms), criterio derivato dalla durata minima fisiologica del *blink* spontaneo (Schiffman,

2001), per escludere doppi rilevamenti di un unico evento. Il *blink rate* è stato calcolato come rapporto tra il numero di *blink* validi e la durata del blocco espressa in minuti, ottenendo un valore per partecipante $\times$ condizione $\times$ blocco.

Pupillometria

Il diametro pupillare è stato utilizzato come indice fisiologico non invasivo del carico cognitivo implicito durante il compito. La pupilla si dilata in risposta all'aumento della domanda mentale e dello sforzo attentivo allocato al compito (Kahneman, 1973), e la misura si è rivelata sensibile alle variazioni di carico in contesti di collaborazione tra umano e sistema intelligente (Kontogiorgos e Gustafson, 2021). Il segnale è stato registrato in continuo con il Pupil Labs Neon tramite il canale di diametro pupillare binoculare (mm).

Il segnale grezzo è stato sottoposto a una pipeline di pulizia sequenziale: rimozione degli artefatti da ammiccamento, esclusione degli outlier (± 3 DS), interpolazione lineare dei gap fino a 500 ms, smoothing con media mobile su finestra di 100 ms, e correzione della baseline per sottrazione (Mathôt et al., 2018; Kret e Sjak-Shie, 2019). I blocchi con meno di 10 campioni validi dopo il preprocessing sono stati esclusi dall'analisi. Per consentire il confronto visivo tra condizioni di durata variabile, il segnale è stato ricampionato su una griglia temporale uniforme di 100 punti per blocco tramite interpolazione lineare (cfr. Nenna et al., 2022). Le traiettorie medie per condizione (con deviazione standard) sono presentate a scopo esplorativo.

La baseline è stata calcolata separatamente per ciascuna condizione (partecipante $\times$ stile AC $\times$ livello di pressione) come media del segnale pupillare nei primi 500 ms del primo blocco di quella condizione. La misura estratta è la variazione media del diametro pupillare rispetto a questa baseline (mean pupil_bc, mm), calcolata separatamente per ciascun blocco. Valori positivi indicano dilatazione rispetto alla baseline, compatibile con maggiore carico cognitivo; valori negativi indicano costrizione.

Misure qualitative

Intervista finale

Al termine della sessione è stata condotta un'intervista semistrutturata finalizzata a integrare i dati quantitativi con una valutazione qualitativa dell'esperienza. L'intervista era composta da sette domande aperte, raggruppabili in quattro aree tematiche.

La prima area riguardava la discriminazione e la preferenza tra le condizioni: ai partecipanti veniva chiesto se avessero notato differenze tra i due sistemi di supporto (Helen e Tk.2) e se avessero preferito lavorare in autonomia o con il supporto dell'agente. La seconda area esplorava la preferenza per lo stile di supporto e la sua utilità percepita sotto pressione temporale: i partecipanti indicavano quale tipologia di AC avrebbero preferito in condizioni di lavoro stressante e motivano

la scelta. La terza area indagava il cambiamento di atteggiamento verso i sistemi intelligenti in contesti lavorativi a seguito dell'esperienza. La quarta area riguardava l'antropomorfizzazione e le preferenze per le modalità di interazione: ai partecipanti veniva chiesto se si fossero immaginati un aspetto fisico per l'AI e se avrebbero desiderato modalità di interazione diverse da quella vocale (ad esempio, un'interfaccia animata).

Le risposte sono state raccolte verbalmente al termine dell'ultimo turno sperimentale, prima del *debriefing*. I dati sono trattati come indicativi di *pattern* qualitativi e non sottoposti ad analisi inferenziale.

4.7 Analisi dei dati

Le analisi statistiche sono state condotte in Python 3.12 tramite i pacchetti statsmodels (Seabold e Perktold, 2010) e pingouin (Vallat, 2018).

Il *delta* del DSSQ (differenza POST – PRE per sessione sulle sottoscale *distress* e *worry*) è stato analizzato tramite *t*-test appaiato anziché LMM: poiché il confronto si riduce a una singola variabile binaria (bassa vs. alta pressione temporale), per la quale il *t*-test appaiato è il metodo appropriato.

Data la struttura *within-subjects* del disegno, tutti i dati ripetuti sono stati analizzati mediante modelli lineari misti (LMM), che gestiscono la non-indipendenza delle osservazioni includendo un'intercetta casuale per partecipante. Il modello fisso specificato per ciascuna variabile dipendente comprende il tipo di supporto (R-CA, F-CA, No-CA), la pressione temporale (assente, presente) e la loro interazione, con No-CA e assenza di pressione come categorie di riferimento. La significatività degli effetti fissi è stata valutata tramite test di Wald χ^2 ; i confronti *post-hoc* a coppie sono stati corretti con il metodo di Bonferroni e le dimensioni dell'effetto riportate come *g* di Hedges per misure appaiate.

I dati pupillometrici sono stati preprocessati tramite una pipeline sequenziale di cinque step (cfr. §4.6). La baseline è stata calcolata separatamente per ciascuna condizione (partecipante $\times$ stile AC $\times$ livello di pressione) come media del segnale pupillare nei primi 500 ms del primo blocco di quella condizione. La misura estratta, variazione media del diametro pupillare rispetto alla baseline (mean pupil_bc, in mm), calcolata separatamente per ciascun blocco, è trattata in modo esplorativo tramite statistiche descrittive e visualizzazione grafica delle traiettorie medie per condizione (con deviazione standard); non è stato condotto alcun test inferenziale formale sulla pupilla.

La frequenza di ammiccamento spontaneo (*blink rate*, ammiccamenti/min) è stata estratta tramite il segnale di apertura palpebrale del Pupil Labs Neon (Pupil Labs, 2023). Un evento *blink* veniva identificato quando l'apertura palpebrale scendeva sotto la soglia di 2.0 mm; erano inclusi solo gli eventi con durata compresa tra 50 e 500 ms (range fisiologico del *blink* spontaneo; Schiffman, 2001) e separati da un *inter-blink interval* ≥ 150 ms per escludere doppi conteggi dello stesso evento. Diversamente dagli altri *outcome*, il *blink rate* è stato analizzato con LMM su dati *trial-level* (tre blocchi per condizione), includendo, oltre all'intercetta casuale, *random slopes* per il

tipo di supporto e la pressione temporale; in caso di mancata convergenza del modello completo era previsto un *fallback* al solo *random intercept*.

5. RISULTATI

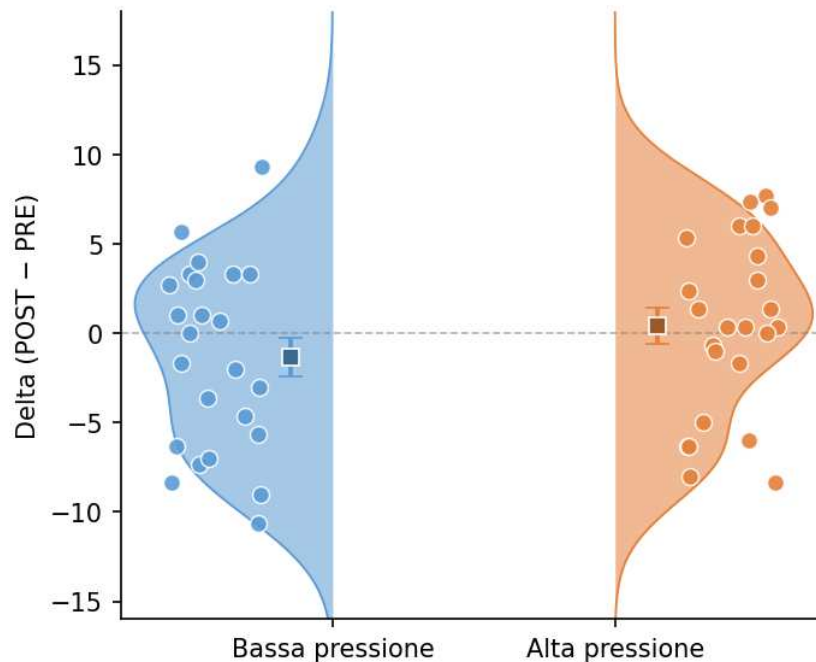
5.1 RQ1: Manipulation check

DSSQ per stress

La risposta soggettiva allo stress è stata misurata tramite il *delta* del Dundee Stress State Questionnaire (DSSQ; Matthews et al., 2002), calcolato come differenza POST – PRE per ciascuna sessione sulle sottoscale *distress* e *worry*. I *delta* sono stati analizzati tramite *t-test* appaiato confrontando le condizioni di bassa e alta pressione temporale, aggregando le condizioni di supporto ($N = 24$, 144 osservazioni).

Delta distress. La pressione temporale ha prodotto un effetto significativo, $t(23) = 2.24$, $p = .035$, $d = 0.34$. In condizione di bassa pressione i partecipanti mostravano una riduzione media del *distress* rispetto alla *baseline* ($M = -1.33$, $DS = 5.25$); in condizione di alta pressione la riduzione era pressoché assente ($M = 0.39$, $DS = 5.00$), con livelli di *distress* significativamente più elevati rispetto alla bassa pressione.

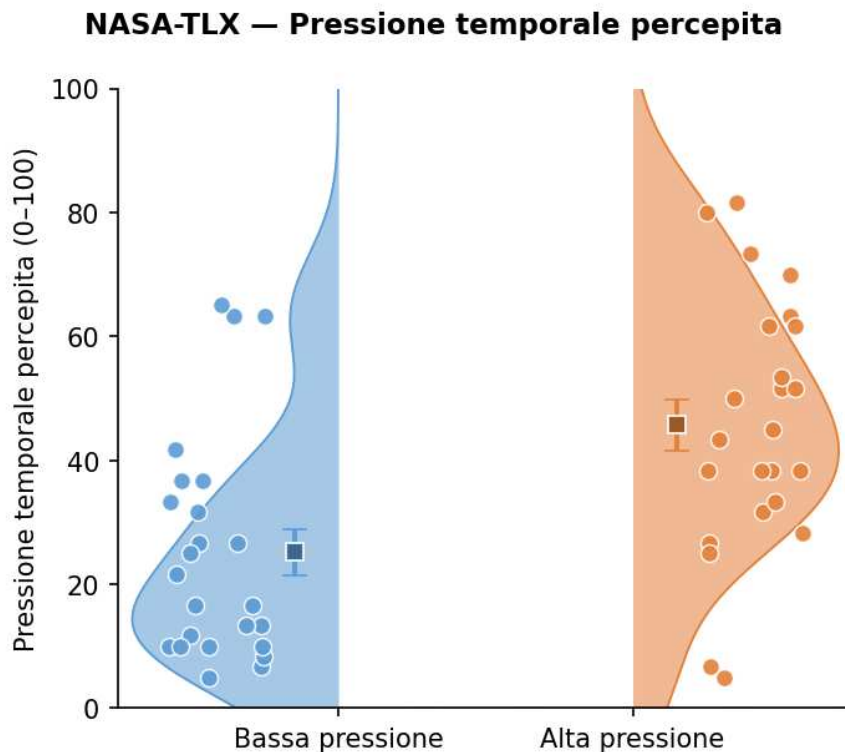
DSSQ — delta_distress (0 = nessun cambiamento dal baseline)



Pressione temporale

La pressione temporale percepita è stata misurata con la sottoscala *temporal demand* del NASA-TLX (Hart e Staveland, 1988) e analizzata con LMM (Tipo di supporto $\times$ Pressione temporale, intercetta casuale per partecipante; $N = 24$, 144 osservazioni).

La pressione temporale ha prodotto un effetto principale significativo, $\chi^2(1) = 15.66$, $p < .001$. I partecipanti riportavano una domanda temporale significativamente più elevata in condizione di alta pressione ($M = 45.69$, $DS = 20.47$) rispetto alla bassa pressione ($M = 25.14$, $DS = 18.32$; $t(23) = 7.29$, $p < .001$, $d = 1.06$).



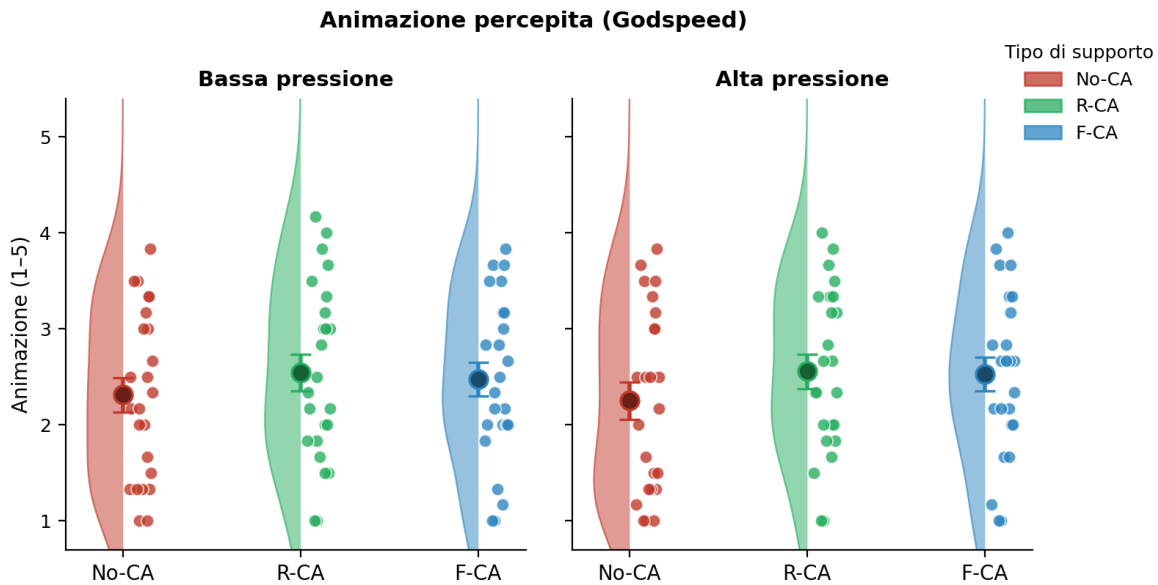
Complessivamente, la manipolazione della pressione temporale risulta confermata su entrambe le misure: il DSSQ-distress ha rilevato un incremento del distress autoriferito ($d = 0.34$), mentre il NASA *temporal demand* ha quantificato la percezione soggettiva del vincolo di tempo con un effetto nettamente più ampio ($d = 1.06$), coerente con la natura esplicita della manipolazione.

Percezione del sistema: animazione, antropomorfismo, simpatia

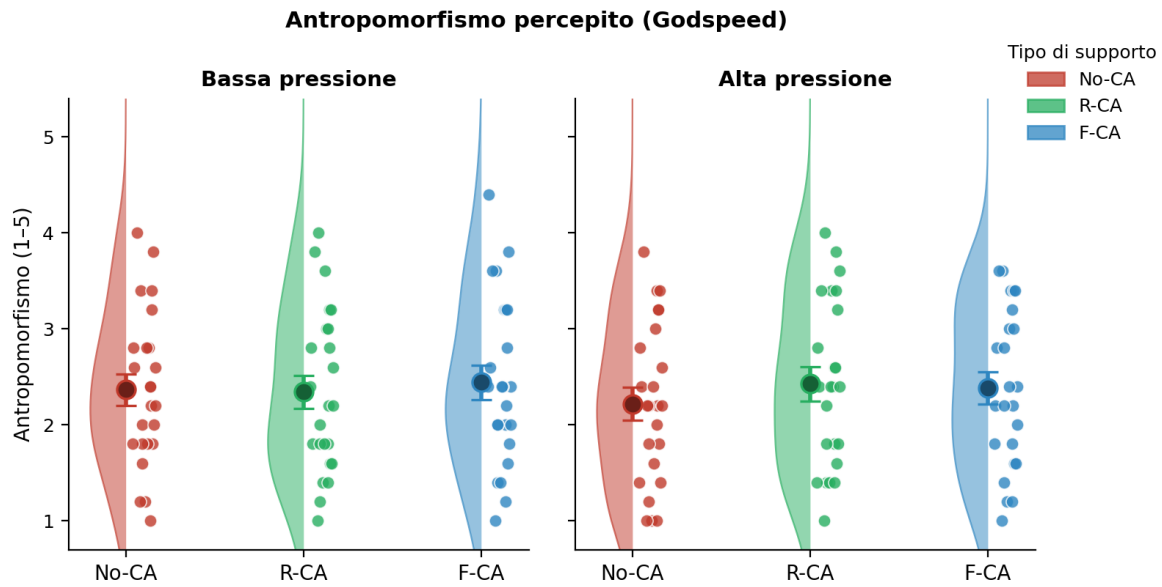
La percezione del robot è stata valutata con il Godspeed Questionnaire (Bartneck et al., 2009; adattamento italiano Operto, 2022). Le sottoscale animazione, antropomorfismo e simpatia sono

state analizzate con LMM (Tipo di supporto $\times$ Pressione temporale, intercetta casuale per partecipante; $N = 24$, 144 osservazioni).

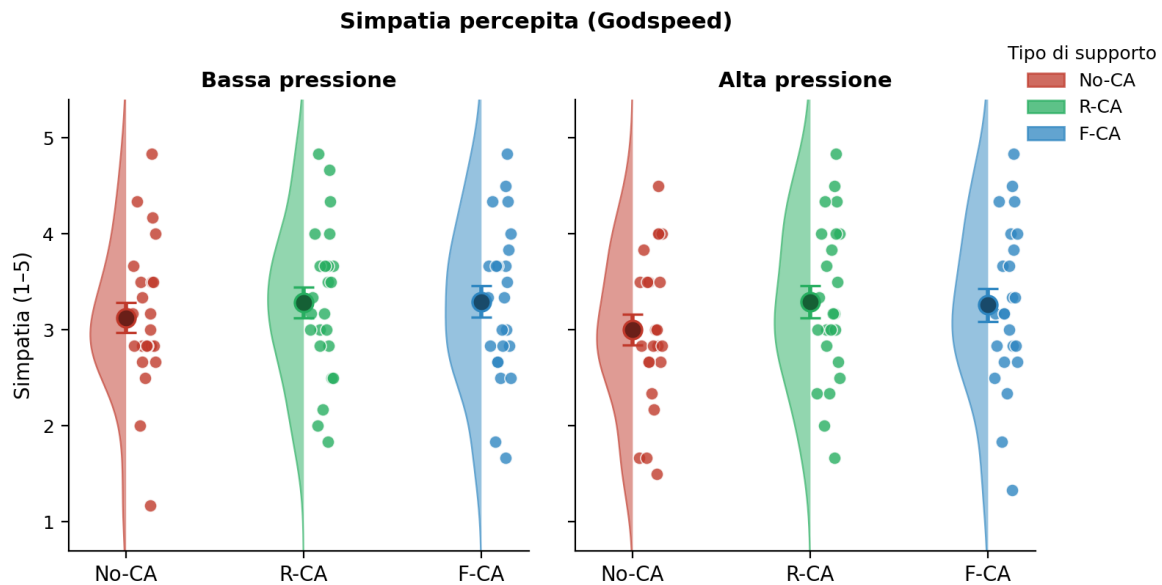
Animazione. Il tipo di supporto ha avuto un effetto significativo, $\chi^2(2) = 8.04$, $p = .018$. Entrambe le condizioni AC erano associate a punteggi più elevati rispetto a No-CA: R-CA ($M = 2.55$, $DS = 0.89$; $p = .033$, $g = 0.29$) e F-CA ($M = 2.50$, $DS = 0.85$; $p = .009$, $g = 0.24$) vs. No-CA ($M = 2.28$, $DS = 0.91$). R-CA e F-CA non differivano ($p = 1.000$, $g = 0.05$). La pressione temporale e l'interazione non erano significative ($p = .451$; $p = .593$).



Antropomorfismo. Il modello non ha evidenziato effetti significativi né del tipo di supporto, $\chi^2(2) = 0.96$, $p = .618$, né della pressione temporale, $\chi^2(1) = 2.00$, $p = .157$, né della loro interazione, $\chi^2(2) = 2.46$, $p = .292$. R-CA ($M = 2.38$, $DS = 0.83$), F-CA ($M = 2.41$, $DS = 0.81$) e No-CA ($M = 2.29$, $DS = 0.80$) non differivano.



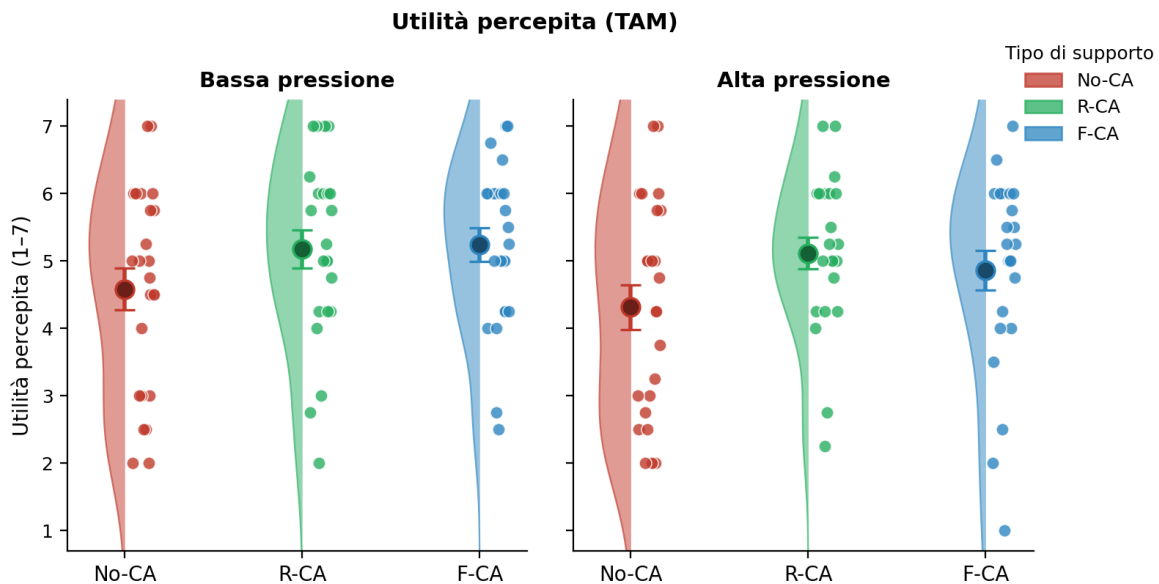
Simpatia. L'effetto del tipo di supporto non raggiungeva la significatività convenzionale, $\chi^2(2) = 4.54$, $p = .103$. La pressione temporale e l'interazione non erano significative ($p = .158$; $p = .559$).



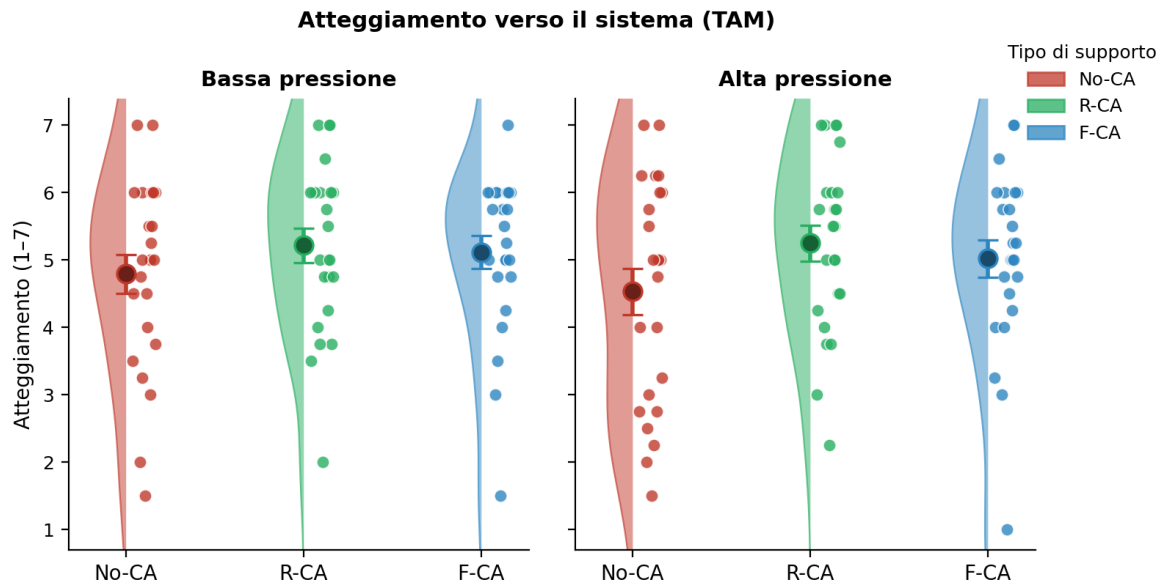
5.2 RQ2: Percezione soggettiva del supporto del AC

RQ2.1: TAM (Technology Acceptance Model)

Il tipo di supporto ha influito sull'**utilità percepita**, $\chi^2(2) = 9.72, p = .008$; la **pressione temporale** e l'**interazione** non risultano significativi ($\chi^2(1) = 1.36, p = .244$; $\chi^2(2) = 0.94, p = .626$). I confronti *post-hoc* mostrano che R-CA supera No-CA ($p = .029, g = 0.50$); i restanti confronti non sono significativi (F-CA vs. No-CA: $p = .076, g = 0.43$; R-CA vs. F-CA: $p = 1.000$). Le medie per condizione sono: R-CA $M = 5.15$ ($DS = 1.25$), F-CA $M = 5.05$ ($DS = 1.33$), No-CA $M = 4.45$ ($DS = 1.55$).



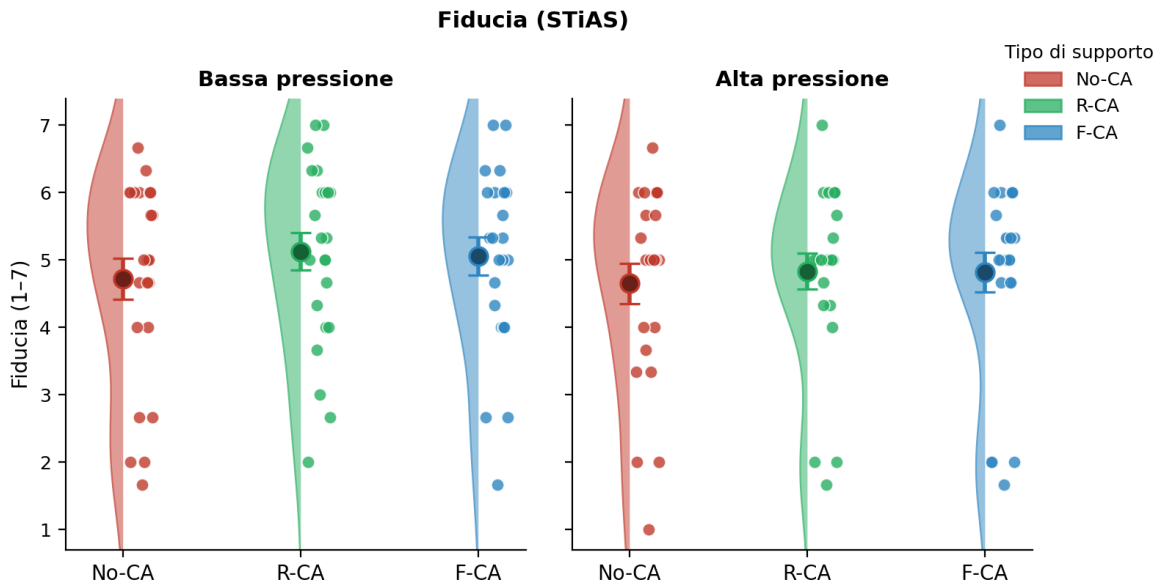
Nessun effetto risulta significativo sull'**atteggiamento verso il sistema** (supporto: $\chi^2(2) = 5.21, p = .074$; pressione temporale: $\chi^2(1) = 1.78, p = .182$; interazione: $\chi^2(2) = 1.12, p = .570$). Le medie sono: R-CA $M = 5.23$ ($DS = 1.27$), F-CA $M = 5.07$ ($DS = 1.26$), No-CA $M = 4.66$ ($DS = 1.54$).



Sulla **facilità d'uso** nessun effetto raggiunge la significatività (supporto: $\chi^2(2) = 0.64$, $p = .725$; pressione temporale: $\chi^2(1) = 3.73$, $p = .053$; interazione: $\chi^2(2) = 0.84$, $p = .656$). Le medie per condizione sono: R-CA $M = 5.63$ ($DS = 0.97$), F-CA $M = 5.58$ ($DS = 1.13$), No-CA $M = 5.40$ ($DS = 1.02$).

RQ2.1: Fiducia (S-TiAS)

La **fiducia** non varia in funzione del supporto, $\chi^2(2) = 4.45$, $p = .108$, della pressione temporale, $\chi^2(1) = 0.12$, $p = .734$, né dell'interazione, $\chi^2(2) = 0.64$, $p = .725$. Le medie sono: R-CA $M = 4.98$ ($DS = 1.33$), F-CA $M = 4.94$ ($DS = 1.39$), No-CA $M = 4.69$ ($DS = 1.47$).



RQ2.2: Controllo, autonomia e sorveglianza

Il supporto non produce effetti significativi su nessuna di esse: controllo $\chi^2(2) = 0.42$, $p = .812$; autonomia $\chi^2(2) = 0.03$, $p = .984$; sorveglianza $\chi^2(2) = 0.39$, $p = .821$. Anche le interazioni con la pressione temporale non sono significative ($p \geq .351$). Emergono unicamente due effetti tendenziali della pressione temporale sul controllo percepito, $\chi^2(1) = 3.15$, $p = .076$, e sulla sorveglianza percepita, $\chi^2(1) = 3.28$, $p = .070$, che non raggiungono comunque la significatività. Le medie per condizione di supporto sono prossime tra loro (controllo: R-CA $M = 4.67$, $DS = 1.35$; F-CA $M = 4.54$, $DS = 1.24$; No-CA $M = 4.28$, $DS = 1.39$; autonomia: 4.46/4.33/4.22; sorveglianza: 3.87/3.98/3.96). Dato il carattere non validato degli item, questi risultati sono da considerare come pattern indicativi.

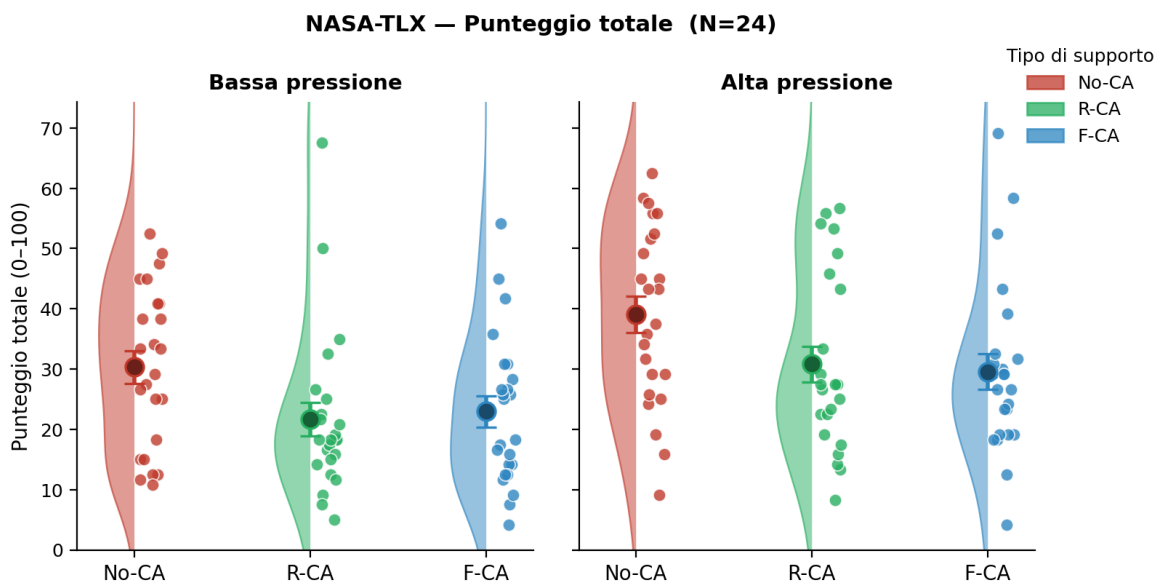
5.3 RQ3: Dinamiche cognitive di workload, attenzione e affaticamento

Carico di lavoro percepito

L'effetto del tipo di supporto sul carico di lavoro percepito è stato analizzato tramite LMM su ciascuna sottoscala del NASA-TLX (Tipo di supporto $\times$ Pressione temporale, intercetta casuale per partecipante; $N = 24$, 144 osservazioni). Tre sottoscale mostrano un effetto principale significativo del supporto: domanda mentale, $\chi^2(2) = 23.12$, $p < .001$; sforzo, $\chi^2(2) = 7.70$, $p = .021$; e punteggio totale composito, $\chi^2(2) = 9.48$, $p = .009$. La domanda temporale mostra un effetto tendenziale, $\chi^2(2)$

= 4.75, $p = .093$. Le sottoscale restanti non raggiungono la significatività: domanda fisica, $\chi^2(2) = 3.15, p = .207$; prestazione percepita, $\chi^2(2) = 0.64, p = .727$; frustrazione, $\chi^2(2) = 2.12, p = .346$.

I confronti *post-hoc* con correzione Bonferroni indicano, per la domanda mentale, valori significativamente inferiori sia in R-CA ($M = 29.7, DS = 21.7$) sia in F-CA ($M = 27.0, DS = 19.2$) rispetto a No-CA ($M = 42.7, DS = 22.6$): rispettivamente $T(23) = -3.85, p_{\text{corr}} = .002, g = -0.57$, e $T(23) = -5.06, p_{\text{corr}} < .001, g = -0.73$. R-CA e F-CA non differiscono ($p_{\text{corr}} = .281, g = 0.14$). La medesima struttura emerge per la sottoscala sforzo: R-CA ($M = 26.6, DS = 18.0$) e F-CA ($M = 27.8, DS = 17.0$) risultano entrambe inferiori a No-CA ($M = 37.5, DS = 19.0$), con $T(23) = -3.39, p_{\text{corr}} = .008, g = -0.60$, e $T(23) = -3.37, p_{\text{corr}} = .008, g = -0.57$; R-CA vs F-CA n.s. ($p_{\text{corr}} = 1.000$). Per il punteggio totale: R-CA ($M = 26.6, DS = 13.1$) e F-CA ($M = 26.7, DS = 12.8$) inferiori a No-CA ($M = 34.7, DS = 12.6$), $T(23) = -3.63, p_{\text{corr}} = .004, g = -0.66$, e $T(23) = -3.50, p_{\text{corr}} = .006, g = -0.66$. Le interazioni supporto $\times$ pressione temporale risultano tutte n.s. ($p \geq .215$).



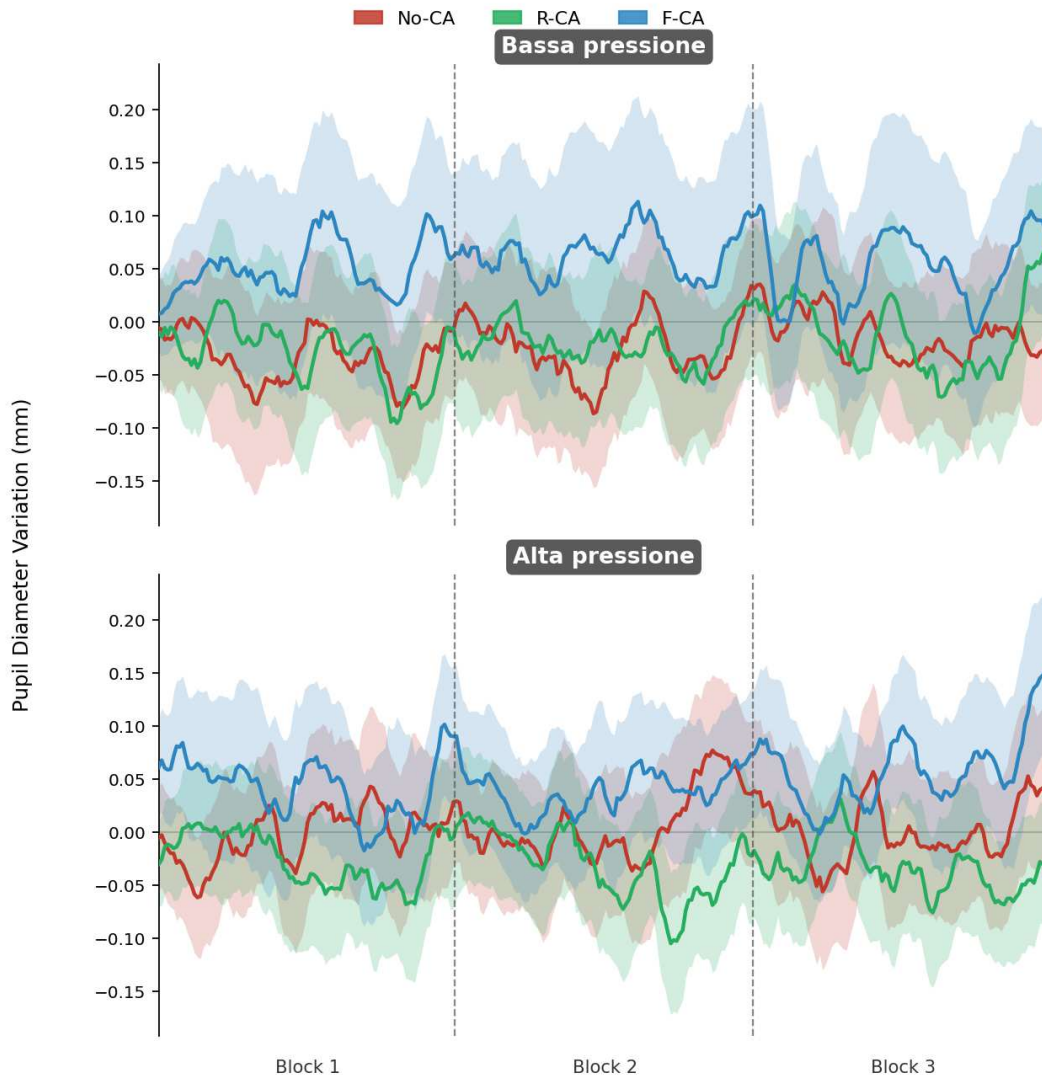
Attivazione fisiologica: pupillometria

Le medie del diametro pupillare medio (*baseline-corrected*) mostrano una struttura differenziata tra le condizioni: F-CA presenta i valori di dilatazione più elevati, No-CA si mantiene prossima allo zero, e R-CA registra i valori più contenuti, in entrambi i livelli di pressione temporale; per i valori di medie e DS si veda Tabella 5.1.

Condizione	Bassa pressione		Alta pressione	
	$M (mm)$	DS	$M (mm)$	DS
No-CA	-0.024	0.273	0.003	0.244
R-CA	-0.021	0.289	-0.030	0.207
F-CA	0.056	0.385	0.046	0.243

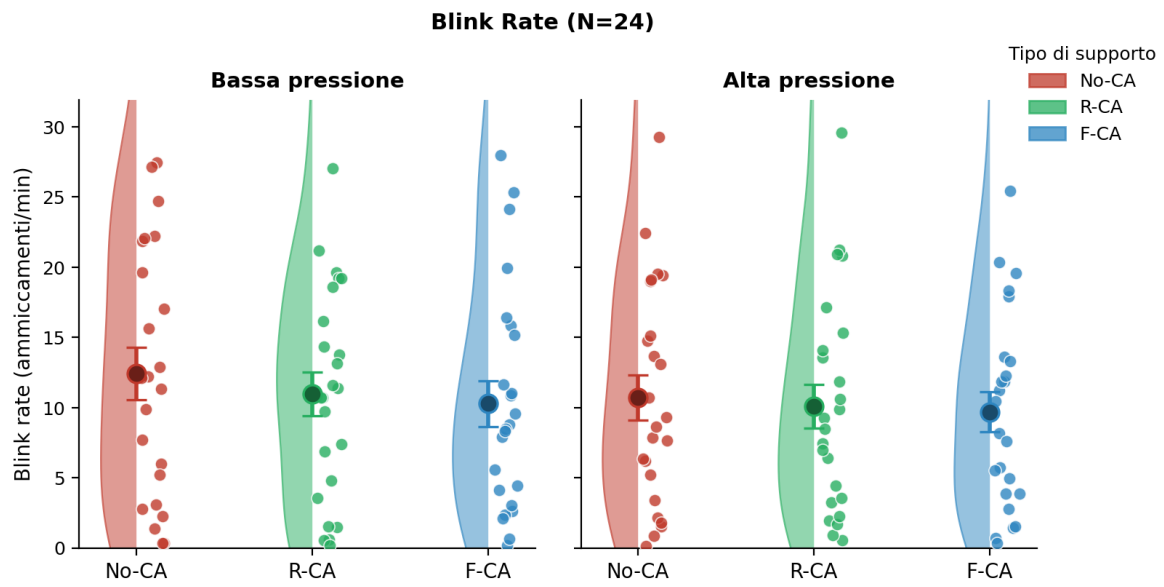
Tabella 5.1 *Diametro pupillare medio (baseline-corrected) per condizione e livello di pressione temporale ($N = 24$)*

Il *time course* del diametro pupillare mostra, in entrambe le condizioni di pressione, una dilatazione sistematicamente maggiore in F-CA rispetto alla propria baseline di condizione, mentre R-CA tende verso valori negativi — in particolare sotto alta pressione nel Block 2 — e No-CA rimane prossima allo zero. Il pattern suggerisce traiettorie di attivazione distinte tra i tre stili di supporto, ma le ampie deviazioni standard e la natura esplorativa dell'analisi ne limitano l'interpretazione.



Frequenza dei blink

Il *blink rate* (frequenza di ammiccamento spontaneo) è stato analizzato con un LMM con *random intercept* e *random slopes* per i fattori *within-subjects*, su dati a livello di singolo blocco ($N = 24$, 428 blocchi totali). L'effetto principale del tipo di supporto è significativo, $\chi^2(5) = 16.11$, $p = .007$. Anche l'effetto della pressione temporale risulta significativo, $\chi^2(2) = 8.63$, $p = .013$, mentre l'interazione supporto $\times$ pressione non raggiunge la significatività, $\chi^2(4) = 6.86$, $p = .143$. I test *post-hoc* Bonferroni, condotti sui dati aggregati per partecipante, mostrano una differenza tendenziale tra No-CA ($M = 11.5$ blinks/min, $DS = 8.24$) e F-CA ($M = 10.0$ blinks/min, $DS = 7.40$): $t(23) = 2.36$, $p_{\text{corr}} = .081$, $g = 0.19$. Il confronto R-CA ($M = 10.5$ blinks/min, $DS = 7.37$) vs. No-CA non raggiunge la significatività ($p_{\text{corr}} = .347$, $g = 0.12$); i due stili di supporto non differiscono tra loro ($p_{\text{corr}} = 1.000$).



6. DISCUSSIONE

Nel complesso, i risultati mostrano un numero contenuto di effetti nella direzione attesa. La manipolazione della pressione temporale si è dimostrata efficace, confermando la validità del paradigma sperimentale. Gli effetti del tipo di supporto AC riguardano principalmente la riduzione del *workload* percepito e, in misura tendenziale, la maggiore utilità del sistema rispetto alla condizione senza supporto; le differenze tra R-CA e F-CA, dove presenti, emergono nelle misure fisiologiche esplorative piuttosto che nelle scale soggettive.

6.1 Domanda di ricerca 1

Manipulation check: pressione temporale. La pressione temporale ha prodotto effetti significativi su entrambe le misure previste. Il DSSQ (Dundee Stress State Questionnaire; Matthews et al., 2002) ha registrato un incremento del *distress* soggettivo in condizione di alta rispetto alla bassa pressione ($d = 0.34$); la sottoscala *temporal demand* del NASA-TLX (Hart e Staveland, 1988) ha riflesso la percezione del vincolo temporale con un effetto di ampiezza nettamente superiore ($d = 1.06$). La convergenza dei due indici su scale diverse conferma la validità della manipolazione.

Manipulation check: percezione del sistema. La Godspeed Scale (Bartneck et al., 2009; Operto, 2022) ha rilevato un effetto del tipo di supporto sull'animazione percepita: entrambe le condizioni R-CA e F-CA sono state valutate come più animate rispetto a No-CA ($g \approx 0.24-0.29$), senza differenze tra i due stili. Antropomorfismo e simpatia non hanno mostrato effetti significativi.

L'effetto sull'animazione riflette la struttura concettuale della sottoscala. L'animazione misura il grado in cui il sistema appare vivo e responsivo: gli item valutano i continuum apatico–responsivo, inerte–interattivo, morto–vivo (Bartneck et al., 2009). Un AC che eroga messaggi — indipendentemente dal registro comunicativo — trasforma il cobot da sistema silenzioso a sistema che agisce e risponde, aumentando la percezione di attività comportamentale. L'equivalenza tra R-CA e F-CA sull'animazione è coerente con il disegno: frequenza e timing dei messaggi erano identici nelle due condizioni, e ciò che le distingueva era il contenuto comunicativo, non il livello di attività del sistema.

L'antropomorfismo misura invece la somiglianza percepita con l'umano attraverso item morfologici (macchina–umano, meccanico–organico, artificiale–naturale) e si riferisce quindi a proprietà fisiche del sistema, non al suo comportamento comunicativo. Messaggi testuali, anche caldi e personali come quelli di R-CA, non modificano la forma del cobot né introducono segnali vocali incarnati che tipicamente sostengono l'attribuzione di tratti umani. La valutazione del sistema è rimasta ancorata alla sua morfologia robotica. Un orientamento simile è documentato da Nielsen et al. (2026) in futuri lavoratori industriali, i quali resistevano attivamente all'attribuzione di caratteristiche umane al cobot, considerandola inappropriata in un contesto produttivo.

Il mancato effetto dello stile comunicativo sulla simpatia è il risultato meno atteso, considerando che R-CA includeva messaggi di incoraggiamento ed espressioni di empatia. Una meta-analisi sui

AC *text-based* ha rilevato un effetto medio delle *social cues human-like* sulle risposte sociali degli utenti ($g = 0.36$), con un'attenuazione sistematica in contesti ad alta complessità di compito (Klein, 2025). L'assemblaggio industriale con cobot rientra in questa categoria: il criterio di valutazione del sistema appare prevalentemente strumentale, e il calore comunicativo di R-CA non aggiunge valore percepibile oltre la disponibilità funzionale già garantita da F-CA. Indicazioni simili emergono da Chattaraman et al. (2019), dove l'effetto positivo dello stile relazionale risultava moderato dalla competenza digitale dell'utente e si manifestava in un contesto di *e-commerce* distante dall'assemblaggio industriale. La coerenza con i risultati di RQ2 rafforza questa interpretazione: poiché utilità percepita, atteggiamento verso il sistema e fiducia non differivano tra R-CA e F-CA, è atteso che non differisse neppure la simpatia globale rilevata a fine sessione.

Nel complesso, il pattern della Godspeed indica che la manipolazione ha prodotto differenziazione comportamentale tra condizioni con e senza AC, registrata come aumento dell'animazione, senza introdurre il *confound* dell'umanizzazione differenziale del sistema.

6.2 Domanda di ricerca 2

Utilità percepita, atteggiamento e fiducia. I due sistemi conversazionali risultano sostanzialmente equivalenti per fiducia verso il sistema automatizzato (STiAS; Körber, 2019), atteggiamento verso il sistema e facilità d'uso. Sull'utilità percepita emerge invece una distinzione tra i due stili: R-CA differisce significativamente da No-CA ($g = 0.50$), mentre F-CA non raggiunge la significatività statistica; i due stili non differiscono tra loro.

Il risultato è in parte contro-intuitivo. Entrambe le condizioni AC alleggerivano il compito nella stessa misura (il sistema gestiva il conteggio di uno dei due colori target indipendentemente dal registro comunicativo), per cui era atteso che F-CA e R-CA producessero effetti analoghi sull'utilità. Che solo R-CA raggiunga la significatività suggerisce che la *perceived usefulness* nel TAM (Davis, 1989) non cattura soltanto l'efficienza operativa del supporto, ma anche la salienza con cui quel supporto viene percepito come orientato all'utente. Il registro relazionale, con le sue segnalazioni esplicite di incoraggiamento e riconoscimento, potrebbe amplificare la percezione di essere attivamente assistiti, aumentando la valutazione soggettiva di utilità indipendentemente dal contenuto informativo. Questo meccanismo è coerente con quanto osservato da Klein (2025) in una meta-analisi sugli AC *text-based*: le *cues* verbali di tipo relazionale producono effetti sulle percezioni dell'utente anche quando il contenuto informativo è equivalente, grazie all'attivazione di schemi di risposta sociale mediati dallo stile (Lombard e Xu, 2021). Gli effect size di R-CA ($g = 0.50$) e F-CA ($g = 0.43$) sono analoghi: la differenza di significatività riflette plausibilmente la potenza statistica limitata del campione ($N = 24$) piuttosto che una differenza sostanziale tra i due stili; con campioni più ampi, entrambe le condizioni AC potrebbero differire significativamente da No-CA.

L'atteggiamento verso il sistema e la fiducia mostrano un pattern uniforme: le medie si collocano nella parte alta delle rispettive scale in tutte e tre le condizioni, senza differenze significative tra R-CA, F-CA e No-CA. L'introduzione dell'AC non ha eroso l'atteggiamento né ridotto la fiducia

rispetto alla condizione priva di supporto. Il dato si inserisce in un contesto in cui l'integrazione di sistemi IA in ambienti produttivi è documentata come potenziale fonte di preoccupazione per i lavoratori, in relazione al controllo della propria attività e al rischio percepito di sostituzione tecnologica (Kopp et al., 2022; Nielsen et al., 2026). Il pattern osservato suggerisce che, nel campione esaminato, la presenza dell'AC non abbia compromesso la valutazione complessiva del sistema. La generalizzabilità di questo esito a campioni più ampi e con profili professionali diversificati resta da verificare.

I risultati di RQ2 indicano un sistema percepito come utile e accettato indipendentemente dal registro comunicativo.

Controllo, autonomia e sorveglianza. Nessuna delle tre dimensioni (controllo percepito, autonomia, sorveglianza percepita) differisce tra le condizioni. Questo risultato segnala l'assenza di effetti destabilizzanti sulla percezione di agenzia dell'operatore: l'introduzione di un agente conversazionale (anche nella variante socio-emotiva R-CA) non compromette la percezione di controllo né genera sensazioni di sorveglianza o perdita di autonomia. In un paradigma di Industria 5.0 che pone il lavoratore al centro del processo produttivo (Lu et al., 2022), la neutralità dell'AC rispetto a queste dimensioni costituisce una premessa favorevole per la sua adozione in contesti reali. Va considerato, tuttavia, che il carattere non validato degli item e la cornice sperimentale (dove la percezione di essere osservati è strutturalmente presente) possono aver ridotto la sensibilità delle misure alle differenze tra condizioni.

6.3 Domanda di ricerca 3

Workload percepito (NASA-TLX). Entrambe le condizioni AC riducono significativamente la domanda mentale e lo sforzo percepiti rispetto a No-CA, con effetti di ampiezza medio-grande ($g=0.57-0.73$). R-CA e F-CA risultano indistinguibili su tutte le sottoscale del NASA-TLX ($|g| < 0.16$ nei confronti diretti). La riduzione del *workload* appare quindi legata alla presenza del supporto operativo piuttosto che al suo stile comunicativo.

Questo *pattern* è coerente con la letteratura sull'HRC: un sistema di supporto allevia selettivamente la componente cognitiva del carico, lasciando inalterate le dimensioni fisiche e temporali (Prati et al., 2023). Va rilevato che nelle condizioni AC il partecipante eseguiva un solo conteggio, delegando al sistema il secondo colore, mentre in No-CA era richiesto il conteggio completo. Questa asimmetria riflette la funzione operativa reale di un AC integrato in HRC: il *task-sharing* è parte integrante del supporto, non una variabile separabile da esso. La riduzione del *workload* osservata cattura quindi l'effetto del sistema nella sua interezza (presenza dell'AC più stile comunicativo), in linea con la validità ecologica del paradigma; isolare il contributo specifico dello stile richiederebbe un design in cui l'AC fornisca messaggi senza svolgere alcuna funzione operativa, condizione che si discosterebbe però da un sistema di supporto realistico.

Attivazione fisiologica: pupillometria. L'analisi esplorativa del diametro pupillare medio (*baseline-corrected*) mostra un gradiente differenziato tra le condizioni: F-CA si associa alla

dilatazione più ampia, No-CA a valori prossimi allo zero e R-CA ai valori più contenuti, in entrambi i livelli di pressione temporale. Il *pattern* non segue la semplice logica del *task-sharing* (presenza dell'AC = minore carico): F-CA, pur riducendo il carico di conteggio del partecipante, mostra la dilatazione pupillare più elevata, superiore anche a No-CA. Questo risultato, di natura esplorativa, suggerisce che lo stile comunicativo possa modulare il profilo di attivazione fisiologica in modo non direttamente riducibile al carico operativo oggettivo.

Il gradiente $R-CA < F-CA$ suggerisce che, a livello implicito, i due stili non siano del tutto equivalenti, e che differenze sottili nel profilo di attivazione cognitiva emergono più chiaramente su misure fisiologiche continue che sulle scale soggettive. Kontogiorgos e Gustafson (2021) hanno esaminato, tramite *eye-tracker*, le risposte pupillari di coppie umano-umano impegnate in un compito collaborativo di assemblaggio (sgabello IKEA), confrontando istruzioni verbali frammentate (distribuite in unità incrementali successive) con istruzioni complete in enunciazione singola. I risultati mostrano che le istruzioni frammentate riducono il carico cognitivo di chi le riceve, consentendo l'elaborazione passo-passo dell'informazione, e che queste differenze si riflettono nel diametro pupillare in modi non sempre catturabili dalle misure soggettive del carico cognitivo.

Nel presente studio, il registro collaborativo e incrementale dei messaggi R-CA (uso della prima persona plurale, riconoscimento progressivo del lavoro svolto) potrebbe aver prodotto un effetto analogo, riducendo la domanda attentiva rispetto allo stile più diretto di F-CA. Pur provenendo da un contesto umano-umano anziché da un sistema AC integrato a un cobot, lo studio di Kontogiorgos e Gustafson (2021) documenta la sensibilità della pupilla a differenze nel *design* delle istruzioni che le misure soggettive non rilevano, fornendo una cornice interpretativa coerente con il gradiente osservato. L'analisi rimane esplorativa, priva di test inferenziale formale, e richiede verifica con campioni più ampi.

Frequenza dei blink. H3 prevedeva una riduzione della frequenza di ammiccamento spontaneo in condizioni di maggiore carico cognitivo. Il modello LMM evidenzia un effetto principale della pressione temporale ($\chi^2(2) = 8.63$, $p = .013$): il blink rate risulta inferiore in condizione di alta pressione ($M = 10.15$ blinks/min) rispetto alla condizione di bassa pressione ($M = 11.22$ blinks/min). I confronti post-hoc tra condizioni AC non raggiungono la significatività dopo correzione di Bonferroni (tutti $p_{\text{corr}} \geq .347$), indicando che H3 trova supporto parziale limitato all'effetto della pressione temporale, mentre nessun effetto attribuibile al tipo di AC emerge dai dati.

La frequenza di ammiccamento spontaneo è sensibile a molteplici fattori (orientamento attentivo, *arousal*, fatica, salienza dello stimolo); la sua interpretazione in assenza di un disegno mirato alla loro disambiguazione richiede cautela (Maffei e Angrilli, 2018). Una riduzione è generalmente associata a un maggiore orientamento attentivo verso lo stimolo: il riflesso viene soppresso in modo adattivo per minimizzare la perdita di informazione visiva durante l'elaborazione di stimoli salienti (Maffei e Angrilli, 2018; Ranti et al., 2020). Il *pattern* osservato è coerente con questa lettura: l'alta pressione temporale richiede una maggiore allocazione di risorse cognitive al compito (Lang, 2000), inducendo una soppressione più marcata del riflesso. La tendenza $F-CA < No-CA$ ($p_{\text{corr}} = .081$, g

= 0.19), pur non raggiungendo la significatività statistica, potrebbe essere compatibile con un'attenzione più ancorata al compito nelle condizioni con supporto strutturato: i messaggi periodici di F-CA forniscono segnali esterni che scandiscono il ritmo operativo, riducendo lo spazio per il *mind-wandering* rispetto a No-CA, dove l'operatore gestisce autonomamente il proprio orientamento attentivo. Si tratta di una lettura plausibile tra più alternative, non discriminabile in base ai soli dati disponibili.

L'entità dell'effetto è piccola e il confronto critico non supera la correzione per confronti multipli. Nel complesso, il quadro di RQ3 indica che la presenza dell'AC produce effetti rilevabili sul carico cognitivo percepito e, in misura più tenue, su misure fisiologiche implicite di orientamento attentivo; le differenze tra R-CA e F-CA rimangono contenute nelle misure soggettive, mentre segnali preliminari emergono sul piano fisiologico.

6.4 Limiti e direzioni future

Il limite metodologico più rilevante riguarda la distinzione tra le condizioni sperimentali. R-CA e F-CA condividevano lo stesso contenuto operativo e differivano per lo strato socio-emotivo sovrapposto: registro amichevole, uso della prima persona plurale e messaggi di incoraggiamento in R-CA; tono anonimo e puramente informativo in F-CA. Questa differenza, pur reale, era contenuta, e su costrutti percettivi relativamente stabili come antropomorfismo e fiducia (Bartneck et al., 2009), per i quali la letteratura documenta effetti della *human-likeness* che si attenuano al crescere della complessità del contesto operativo (Klein, 2025), una manipolazione dello stile comunicativo di questa entità potrebbe non essere sufficiente a produrre effetti rilevabili.

Il *design within-subjects*, pur consentendo un controllo statistico efficiente della variabilità individuale, espone al rischio di *carry-over* e *demand characteristics*: avendo sperimentato tutte e tre le condizioni nella medesima sessione, i partecipanti potrebbero aver confrontato attivamente gli stili di supporto, orientando le valutazioni successive sulla base dell'esposizione precedente. Il quadrato latino bilancia l'ordine ma non elimina questi effetti di contrasto.

La dimensione campionaria ($N = 24$) riduce la potenza statistica: gli *effect size* analoghi tra R-CA e F-CA su diverse misure di percezione lasciano aperta la possibilità che, con campioni più ampi, entrambe le condizioni differiscano significativamente da No-CA. La selezione per convenienza produce un campione eterogeneo per istruzione e occupazione (il 79.2% dei partecipanti risulta occupato o studente-lavoratore) ma con scarsa esposizione al contesto manifatturiero: solo il 29.2% aveva sentito parlare di cobot e l'8.3% ne aveva esperienza diretta. Al contrario, l'83.3% del campione aveva già utilizzato sistemi IA, una familiarità che potrebbe aver generato aspettative precostituite sullo stile comunicativo dell'AC. Questa combinazione (bassa familiarità con i robot, alta familiarità con l'IA) riduce la generalizzabilità dei risultati a operatori industriali con esposizione consolidata alla tecnologia collaborativa.

Le misure *self-report* erano raccolte al termine di ogni condizione, all'interno della stessa sessione sperimentale. La ripetizione dei questionari su sei blocchi consecutivi espone a effetti di fatica da

risposta e a derive sistematiche nelle valutazioni, che potrebbero aver attenuato la sensibilità degli strumenti nelle condizioni più tarde della sessione.

Il contesto di laboratorio non riproduce le condizioni operative reali: ritmi produttivi sostenuti, familiarità con i sistemi e dinamiche di squadra possono modulare le risposte in modi non prevedibili dal *setting* controllato.

Direzioni future. Un'estensione immediata delle analisi riguarda il materiale qualitativo raccolto nelle interviste post-sperimentali, che può offrire elementi interpretativi sull'esperienza dei partecipanti non catturabili dai questionari strutturati. Sul piano fisiologico, la pupillometria è rimasta sul piano descrittivo ed esplorativo: studi futuri dovrebbero applicare modelli inferenziali che integrino diametro pupillare e *blink rate*, consentendo un confronto diretto con le misure di attivazione cardiovascolare registrate ma non analizzate nel presente studio.

Sul piano campionario, studi futuri dovrebbero coinvolgere campioni più ampi e operatori con esperienza diretta in contesti di produzione avanzata, possibilmente attraverso *field study* in ambienti industriali reali, per verificare la generalizzabilità dei risultati e testare se le tendenze emerse sulle misure fisiologiche si consolidano in condizioni operative più ecologiche. Sarebbe inoltre utile esaminare il ruolo di variabili disposizionali, come l'attitudine verso i robot (Carradore et al., 2024) e l'ansia da IA (Wang e Wang, 2022), attraverso analisi di moderazione, per identificare i profili individuali più sensibili agli effetti dello stile comunicativo dell'AC.

Manipolazioni più marcate degli attributi del AC, come personalità dell'agente, grado di *embodiment* visivo (Kiuchi et al., 2024, Nunnari et al., 2025), prosodia, consentirebbero di verificare se una caratterizzazione più differenziata produca effetti robusti su costrutti percettivi stabili e di verificare i possibili effetti della dimensione relazionale rispetto all'utilità operativa percepita. Infine, l'analisi degli errori di conteggio e dei tempi di completamento del blocco, registrati ma non analizzati nel presente studio, consentirebbe di valutare se la riduzione del *workload* percepito si traduca in un miglioramento degli esiti produttivi, dato rilevante per la sostenibilità dell'HRC con AC in contesti industriali reali.

7. CONCLUSIONI

Il presente studio ha esaminato se e come lo stile comunicativo di un AC integrato a un cobot industriale, nelle varianti R-CA e F-CA, modifichi la risposta psicologica e fisiologica degli operatori in un compito di disassemblaggio collaborativo sotto pressione temporale variabile.

Il risultato più robusto è anche quello meno scontato: integrare un AC al cobot rende il sistema percepito più animato e responsivo, un agente attivo (Bartneck et al., 2009), senza produrre nessuno dei costi che la letteratura indica come più probabili nell'introduzione di IA che monitorano e valutano l'operato in ambienti di lavoro (Tausch e Kluge, 2020; Berretta et al., 2023; Stock-Homburg e Nguyen, 2023). Controllo percepito, autonomia e sorveglianza restano invariati tra le tre condizioni: l'agente conversazionale, anche nella variante socio-emotiva R-CA, non genera vissuti di monitoraggio negativo né erode la percezione di agenzia dell'operatore. Questo *pattern* è coerente con il framework CASA (Nass et al., 1994), ma lo estende a un contesto produttivo reale in cui tali effetti avrebbero potuto essere neutralizzati o rovesciati da preoccupazioni legate alla sorveglianza tecnologica (Kopp et al., 2022, 2023; Nielsen et al., 2026).

Sul piano dell'utilità e del carico cognitivo, i risultati convergono in una direzione chiara: R-CA supera No-CA sull'utilità percepita ($g = 0.50$), e entrambe le condizioni AC riducono il carico mentale e lo sforzo percepiti rispetto all'assenza di supporto ($g = 0.52-0.70$), indipendentemente dal livello di pressione temporale. L'AC non agisce come fonte di carico aggiuntivo ma come *offload* cognitivo stabile (Hmedan et al., 2021), liberando risorse attentive a prescindere dalle condizioni operative (Lang, 2000; Kahneman, 1973). È precisamente la direzione richiesta da una progettazione HRC centrata sul lavoratore nel paradigma Industria 5.0 (Lu et al., 2022): un sistema che amplifica le capacità dell'operatore senza sottrarne l'autonomia.

Sul piano metodologico, lo studio presenta alcune caratteristiche che ne rafforzano l'interpretabilità. Il compito di disassemblaggio con un cobot UR10e in ambiente di laboratorio controllato garantisce validità ecologica superiore agli studi su scenari ipotetici o stimoli virtuali. La triangolazione tra misure soggettive (TAM, NASA-TLX, Godspeed, S-TIAS), fisiologiche (*blink rate*, pupillometria) e prestazionali consente di valutare l'impatto dell'AC su più livelli di risposta dell'operatore. Il disegno *within-subjects* controbilanciato tramite quadrato latino permette il confronto intra-individuale tra tutti e tre gli stili di supporto. Non secondario è il contributo sul piano della letteratura: la mappatura sistematica dei due poli comunicativi (R-CA e F-CA) a partire da una terminologia frammentata (Shum et al., 2018; Van Pinxteren et al., 2020) rappresenta di per sé un contributo alla chiarificazione concettuale del campo. Che questi effetti emergano in un campione con familiarità con i cobot molto limitata (8.3% con esperienza diretta) li rende ancora più incoraggianti per l'adozione in contesti produttivi reali: se l'AC produce *offload* cognitivo e accettazione positiva anche in chi non ha mai interagito con un robot collaborativo, le prospettive di scalabilità sono favorevoli.

La diffusione dell'IA nei contesti produttivi (Brynjolfsson et al., 2024) sposta il problema progettuale: non si tratta più di decidere se dotare il cobot di una componente conversazionale, ma di scegliere con quale stile farlo. I risultati di questo studio indicano che orientarsi verso il polo relazionale è una scelta praticabile senza i costi che la letteratura associa all'introduzione di IA valutativa in ambito lavorativo: l'AC arricchisce la percezione del sistema e alleggerisce il carico cognitivo dell'operatore, senza eroderne il senso di controllo né l'autonomia.

Workload, accettazione, fiducia, controllo e autonomia percepita definiscono lo spazio entro cui questa scelta va progettata e verificata. Calibrare l'AC su questi indicatori non è un dettaglio tecnico: è ciò che determina se l'operatore rimane agente del processo produttivo o ne diventa semplicemente destinatario.

BIBLIOGRAFIA

- Adamopoulou, E., & Moussiades, L. (2020). Chatbots: History, technology, and applications. *Machine Learning with Applications*, 2, Article 100006. <https://doi.org/10.1016/j.mlwa.2020.100006>
- Allouch, M., Azaria, A., & Azoulay, R. (2021). Conversational agents: Goals, technologies, vision and challenges. *Sensors*, 21(24), Article 8448. <https://doi.org/10.3390/s21248448>
- Arai, T., Kato, R., & Fujita, M. (2010). Assessment of operator stress induced by robot collaboration in assembly. *CIRP Annals*, 59(1), 5–8.
- Babamiri, M., Heidarimoghadam, R., Ghasemi, F., Tapak, L., & Morteza pour, A. (2024). Going beyond general stress scales: Developing a new questionnaire to measure stress in human-robot interaction. *International Journal of Social Robotics*, 16, 2243–2259. <https://doi.org/10.1007/s12369-024-01183-5>
- Baltrusch, S. J., Krause, F., de Vries, A. W., van Dijk, W., & de Looze, M. P. (2022). What about the human in human robot collaboration? A literature review on HRC's effects on aspects of job quality. *Ergonomics*, 65(5), 719–740. <https://doi.org/10.1080/00140139.2021.1984585>
- Bartneck, C., Kulić, D., Croft, E., & Zoghbi, S. (2009). Measurement instruments for the anthropomorphism, animacy, likeability, perceived intelligence, and perceived safety of robots. *International Journal of Social Robotics*, 1(1), 71–81.
- Bassi, G., Orso, V., Salcuni, S., & Gamberini, L. (2025). Understanding workers' well-being and cognitive load in human-cobot collaboration: Systematic review. *Journal of Medical Internet Research*, 27, e75658. <https://doi.org/10.2196/75658>
- Berretta, S., Tausch, A., Peifer, C., & Kluge, A. (2023). The Job Perception Inventory: Considering human factors and needs in the design of human–AI work. *Frontiers in Psychology*, 14, Article 1128945. <https://doi.org/10.3389/fpsyg.2023.1128945>
- Bousdekis, A., Foosherian, M., Fikardos, M., Wellsandt, S., Lepenioti, K., Bosani, E., Mentzas, G., & Thoben, K.-D. (2025). Augmented intelligence with voice assistance and automated machine learning in Industry 5.0. *Frontiers in Artificial Intelligence*. <https://doi.org/10.3389/frai.2025.1538840>
- Brynjolfsson, E., Li, D., & Raymond, L. (2024). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889–942. <https://doi.org/10.1093/qje/qjae044>
- Cai, N., Gao, S., & Yan, J. (2024). How the communication style of chatbots influences consumers' satisfaction, trust, and engagement in the context of service failure. *Humanities and Social Sciences Communications*, 11, 687. <https://doi.org/10.1057/s41599-024-03212-0>

- Carissoli, C., Negri, L., Bassi, M., Storm, F. A., & Delle Fave, A. (2024). Mental workload and human-robot interaction in collaborative tasks: A scoping review. *International Journal of Human-Computer Interaction*, 40(20), 6458–6477. <https://doi.org/10.1080/10447318.2023.2254639>
- Carradore, M., Artioli, G., & Sarli, A. (2024). The General Attitudes Towards Robots Scale (GAToRS): A preliminary validation of the Italian version. *International Journal of Social Robotics*, 16, 2001–2018. <https://doi.org/10.1007/s12369-024-01170-w>
- Chattaraman, V., Kwon, W.-S., Gilbert, J. E., & Ross, K. (2019). Should AI-based, conversational digital assistants employ social- or task-oriented interaction style? A task-competency and reciprocity perspective for older adults. *Computers in Human Behavior*, 90, 315–330. <https://doi.org/10.1016/j.chb.2018.08.048>
- Chen, J., Zhang, Y., & Wang, L. (2025). The impact of service robot communication style on consumers' continued willingness to use. *Collabra: Psychology*, 11(1). <https://doi.org/10.1525/collabra.128020>
- Chen, Q., Chen, M., Lin, L., & Bai, X. (2024). The challenge–hindrance–threat appraisal framework and the differential effects on employees' work well-being and behaviors. *Behavioral Sciences*, 14, 734. <https://doi.org/10.3390/bs14090734>
- Dahlbäck, N., Jönsson, A., & Ahrenberg, L. (1993). Wizard of Oz studies: Why and how. *Knowledge-Based Systems*, 6(4), 258–266. [https://doi.org/10.1016/0950-7051\(93\)90017-N](https://doi.org/10.1016/0950-7051(93)90017-N)
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
- Davis, F. D., Bagozzi, R. P., & Warshaw, P. R. (1989). User acceptance of computer technology: A comparison of two theoretical models. *Management Science*, 35(8), 982–1003. <https://doi.org/10.1287/mnsc.35.8.982>
- Deci, E. L., & Ryan, R. M. (1985). *Intrinsic motivation and self-determination in human behavior*. Plenum Press.
- Deci, E. L., & Ryan, R. M. (2000). The “what” and “why” of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227–268. https://doi.org/10.1207/S15327965PLI1104_01
- Di Pasquale, V., De Simone, V., Giubileo, V., & Miranda, S. (2023). A taxonomy of factors influencing worker's performance in human–robot collaboration. *IET Collaborative Intelligent Manufacturing*. <https://doi.org/10.1049/cim2.12069>
- Eese, A. K., Kostolani, D., Kang, T., Schlund, S., Medvegy, T., Abonyi, J., & Ruppert, T. (2024). May I have your attention?! Exploring multitasking in human-robot collaboration. *IFAC PapersOnLine*, 58(19).
- Faccio, M., Granata, I., Menini, A., Milanese, M., Rossato, C., Bottin, M., Minto, R., Pluchino, P., Gamberini, L., Boschetti, G., & Rosati, G. (2023). Human factors in cobot era: A

review of modern production systems features. *Journal of Intelligent Manufacturing*, 34, 85–106. <https://doi.org/10.1007/s10845-022-01953-w>

Freire, I. T., Guerrero-Rosado, O., Amil, A. F., & Verschure, P. F. M. J. (2024). Socially adaptive cognitive architecture for human-robot collaboration in industrial settings. *Frontiers in Robotics and AI*, 11, Article 1248646. <https://doi.org/10.3389/frobt.2024.1248646>

Gervasi, R., Capponi, M., Mastrogiacomo, L., & Franceschini, F. (2023). Analyzing psychophysical state and cognitive performance in human-robot collaboration for repetitive assembly processes. *Production Engineering*, 18, 19–33. <https://doi.org/10.1007/s11740-023-01230-6>

Gilles, M., Bach, C., Quentel, M., Pallamin, N., Jégou, G., Querrec, R., & Bevacqua, E. (2025). Pilot study on virtual assistant design for new generation aircraft cockpit. *Behaviour & Information Technology*. <https://doi.org/10.1080/0144929X.2025.2507691>

Hart, S. G., & Staveland, L. E. (1988). Development of NASA-TLX (Task Load Index): Results of empirical and theoretical research. In P. A. Hancock & N. Meshkati (Eds.), *Human mental workload* (pp. 139–183). North-Holland.

Hermanns, L., & Teubner, T. (2025). Under pressure: How time constraints, task complexity, and AI reliability shape human-AI interaction. *Behaviour & Information Technology*. <https://doi.org/10.1080/0144929X.2025.2587732>

Hmedan, B., Kilgus, D., Fiorino, H., Landry, A., & Pellier, D. (2021). Adapting cobot behavior to human task ordering variability for assembly tasks. In *Proceedings of the 2021 IEEE International Conference on Human-Machine Systems (ICHMS)*.

Hobfoll, S. E. (1989). Conservation of resources: A new attempt at conceptualizing stress. *American Psychologist*, 44(3), 513–524. <https://doi.org/10.1037/0003-066X.44.3.513>

Hoff, K. A., & Bashir, M. (2015). Trust in automation: Integrating empirical evidence on factors that influence trust. *Human Factors*, 57(3), 407–434. <https://doi.org/10.1177/0018720814547570>

Jacob, F., Grosse, E. H., Morana, S., & König, C. J. (2023). Picking with a robot colleague: A systematic literature review and evaluation of technology acceptance in human–robot collaborative warehouses. *Computers & Industrial Engineering*, 180, 109262. <https://doi.org/10.1016/j.cie.2023.109262>

Kahneman, D. (1973). *Attention and effort*. Prentice-Hall.

Keppel, G., & Wickens, T. D. (2004). *Design and analysis: A researcher's handbook* (4th ed.). Pearson.

Kiuchi, K., Otsu, K., & Hayashi, Y. (2024). Psychological insights into the research and practice of embodied conversational agents, chatbots and social assistive robots: A systematic meta-review. *Behaviour & Information Technology*, 43(16), 3696–3736. <https://doi.org/10.1080/0144929X.2023.2286528>

- Klein, S. H. (2025). The effects of human-like social cues on social responses towards text-based conversational agents: A meta-analysis. *Humanities and Social Sciences Communications*, 12. <https://doi.org/10.1057/s41599-025-05618-w>
- Kohn, S. C., de Visser, E. J., Wiese, E., Lee, Y.-C., & Shaw, T. H. (2021). Measurement of trust in automation: A narrative review and reference guide. *Frontiers in Psychology*, 12, Article 604977. <https://doi.org/10.3389/fpsyg.2021.604977>
- Kolomaznik, M., Petrik, V., Slama, M., & Jurik, V. (2024). The role of socio-emotional attributes in enhancing human-AI collaboration. *Frontiers in Psychology*, 15, Article 1369957. <https://doi.org/10.3389/fpsyg.2024.1369957>
- Kontogiorgos, D., & Gustafson, J. (2021). Measuring collaboration load with pupillary responses: Implications for the design of instructions in task-oriented HRI. *Frontiers in Psychology*, 12, Article 623657. <https://doi.org/10.3389/fpsyg.2021.623657>
- Kopp, T. (2024). Facets of trust and distrust in collaborative robots at the workplace: Towards a multidimensional and relational conceptualisation. *International Journal of Social Robotics*, 16, 1445–1462. <https://doi.org/10.1007/s12369-023-01082-1>
- Kopp, T., Baumgartner, M., & Kinkel, S. (2022). How linguistic framing affects factory workers' initial trust in collaborative robots: The interplay between anthropomorphism and technological replacement. *International Journal of Human-Computer Studies*, 158, 102730. <https://doi.org/10.1016/j.ijhcs.2021.102730>
- Kopp, T., Baumgartner, M., & Kinkel, S. (2023). “It’s not Paul, it’s a robot”: The impact of linguistic framing and the evolution of trust and distrust in a collaborative robot during a human-robot interaction. *International Journal of Human-Computer Studies*, 178, 103095. <https://doi.org/10.1016/j.ijhcs.2023.103095>
- Körber, M. (2019). Theoretical considerations and development of a questionnaire to measure trust in automation. In *Proceedings of the 20th Congress of the International Ergonomics Association (IEA 2018)*. Springer, Cham.
- Kret, M. E., & Sjak-Shie, E. E. (2019). Preprocessing pupil size data: Guidelines and code. *Behavior Research Methods*, 51(3), 1336–1342. <https://doi.org/10.3758/s13428-018-1075-y>
- Lampi, A., Salo, M., Venermo, K., & Pirkkalainen, H. (2023). Emergence of technostress among employees working with physical robots. In *Proceedings of the 31st European Conference on Information Systems (ECIS 2023)*. Kristiansand, Norway.
- Lang, A. (2000). The limited capacity model of mediated message processing. *Journal of Communication*, 50(1), 46–70. <https://doi.org/10.1111/j.1460-2466.2000.tb02833.x>
- Lazarus, R. S., & Folkman, S. (1984). *Stress, appraisal, and coping*. Springer.
- Li, J., Zhang, L., & Mao, M. (2023). How does human-AI interaction affect employees' workplace procrastination? *Technological Forecasting and Social Change*, 212, 123951. <https://doi.org/10.1016/j.techfore.2024.123951>

- Liao, S., Lin, L., & Chen, Q. (2023). Research on the acceptance of collaborative robots for the Industry 5.0 era: The mediating effect of perceived competence and the moderating effect of robot use self-efficacy. *International Journal of Industrial Ergonomics*, 95, Article 103492. <https://doi.org/10.1016/j.ergon.2023.103492>
- Lombard, M., & Xu, K. (2021). Social responses to media technologies in the 21st century: The media are social actors paradigm. *Human-Machine Communication*, 2, 29–55. <https://doi.org/10.30658/hmc.2.2>
- Lu, L., Xie, Z., Wang, H., Li, L., & Xu, X. (2022a). Mental stress and safety awareness during human-robot collaboration. *Applied Ergonomics*, 105, 103832. <https://doi.org/10.1016/j.apergo.2022.103832>
- Lu, Y., Zheng, H., Chand, S., Xia, W., Liu, Z., Xu, X., Wang, L., Qin, Z., & Bao, J. (2022b). Outlook on human-centric manufacturing towards Industry 5.0. *Journal of Manufacturing Systems*, 62, 612–627. <https://doi.org/10.1016/j.jmsy.2022.02.001>
- Maffei, A., & Angrilli, A. (2018). Spontaneous eye blink rate: An index of dopaminergic component of sustained attention and fatigue. *International Journal of Psychophysiology*, 123, 58–63. <https://doi.org/10.1016/j.ijpsycho.2017.11.009>
- Maffei, A., & Angrilli, A. (2019). Spontaneous blink rate as an index of attention and emotion during film clips viewing. *Physiology & Behavior*, 204, 256–263. <https://doi.org/10.1016/j.physbeh.2019.02.037>
- Mariscal, M. A., Ortiz Barcina, S., García Herrero, S., & López Perea, E. M. (2024). Working with collaborative robots and its influence on levels of working stress. *International Journal of Computer Integrated Manufacturing*, 37(7), 900–919. <https://doi.org/10.1080/0951192X.2023.2263428>
- Mathôt, S., Fabius, J., Van Heusden, E., & Van der Stigchel, S. (2018). Safe and sensible preprocessing and baseline correction of pupil-size data. *Behavior Research Methods*, 50(1), 94–106. <https://doi.org/10.3758/s13428-017-1007-2>
- Matthews, G., Campbell, S. E., Falconer, S., Joyner, L. A., Huggins, J., Gilliland, K., Grier, R. A., & Warm, J. S. (2002). Fundamental dimensions of subjective state in performance settings: Task engagement, distress, and worry. *Emotion*, 2(4), 315–340. <https://doi.org/10.1037/1528-3542.2.4.315>
- Molan, J., Saad, L., Roesler, E., McCurry, J. M., Gyory, N., & Trafton, J. G. (2025). Perceived danger scale: Development and validation. *arXiv*. <https://doi.org/10.48550/arXiv.2501.14099>
- Moon, Y., & Nass, C. (1996). How “real” are computer personalities? Psychological responses to personality types in human-computer interaction. *Communication Research*, 23(6), 651–674. <https://doi.org/10.1177/009365096023006002>
- Nass, C., Steuer, J., & Tauber, E. R. (1994). Computers are social actors. In *Proceedings of CHI* (pp. 72–78). ACM.

- Nass, C., Moon, Y., Fogg, B. J., Reeves, B., & Dryer, D. C. (1995). Can computer personalities be human personalities? *International Journal of Human-Computer Studies*, 43(2), 223–239. <https://doi.org/10.1006/ijhc.1995.1042>
- Nenna, F., Orso, V., Zanardi, D., & Gamberini, L. (2022). The virtualization of human–robot interactions: A user-centric workload assessment. *Virtual Reality*. <https://doi.org/10.1007/s10055-022-00667-x>
- Nenna, F., Zanardi, D., & Gamberini, L. (2023). Enhanced interactivity in VR-based telerobotics: An eye-tracking investigation of human performance and workload. *International Journal of Human-Computer Studies*, 177, 103079. <https://doi.org/10.1016/j.ijhcs.2023.103079>
- Nielsen, S., Jochum, E. A., Li, C., Chrysostomou, D., & Ordoñez, R. (2026). “Robot emotions are not real!”: Future factory workers’ perceptions, attitudes, and experience of collaborative robots, conversational AIs, and AI-empowered, voice-enabled collaborative robots. *ACM Transactions on Human-Robot Interaction*, 15(2). <https://doi.org/10.1145/3779295>
- Nunnari, F., Tsovaltzi, D., Lavit Nicora, M., Beyrodt, S., Prajod, P., Chehayeb, L., Brdar, I., Delle Fave, A., Negri, L., André, E., Gebhard, P., & Malosio, M. (2025). Socially interactive industrial robots: A PAD model of flow for emotional co-regulation. *Frontiers in Robotics and AI*, 11, Article 1418677. <https://doi.org/10.3389/frobt.2024.1418677>
- Operto, S. (2021). HRI: L’interazione tra esseri umani e macchine. Dall’interazione sociale all’interazione sociotecnica [Tesi di dottorato]. Università di Genova. https://doi.org/10.15167/operto-stefania_phd2021-11-22
- Orlando, E. M., Nenna, F., Zanardi, D., Buodo, G., Mingardi, M., Sarlo, M., & Gamberini, L. (2025). Understanding workers' psychological states and physiological responses during human–robot collaboration. *International Journal of Human-Computer Studies*, 200, 103516. <https://doi.org/10.1016/j.ijhcs.2025.103516>
- Ososky, S., Sanders, T., Jentsch, F., Hancock, P. A., & Chen, J. Y. C. (2014). Determinants of system transparency and its influence on trust in and reliance on unmanned robotic systems. *Proceedings of SPIE*, 9084, Article 90840E. <https://doi.org/10.1117/12.2050622>
- Panchetti, T., Pietrantoni, L., Puzzo, G., Gualtieri, L., & Fraboni, F. (2023). Assessing the relationship between cognitive workload, workstation design, user acceptance and trust in collaborative robots. *Applied Sciences*, 13(3), 1720. <https://doi.org/10.3390/app13031720>
- Pluchino, P., Pernice, G. F. A., Nenna, F., Mingardi, M., Bettelli, A., Bacchin, D., Spagnolli, A., Jacucci, G., Ragazzon, A., Miglioranzi, L., Pettenon, C., & Gamberini, L. (2023). Advanced workstations and collaborative robots: Exploiting eye-tracking and cardiac activity indices to unveil senior workers' mental workload in assembly tasks. *Frontiers in Robotics and AI*, 10, 1275572. <https://doi.org/10.3389/frobt.2023.1275572>

- Prati, E., Borsci, S., Peruzzini, M., & Pellicciari, M. (2022). A systematic literature review of user experience evaluation scales for human-robot collaboration. In *Transdisciplinarity and the Future of Engineering*. IOS Press. <https://doi.org/10.3233/ATDE220627>
- Pupil Labs. (2023). Neon: Technical documentation. Pupil Labs GmbH.
- Puspanathan, T., Bitsch, G., & Louw, L. (2024). Systematic literature review on the acceptance of human-robot collaboration among assembly workers. *Procedia CIRP*, 126, 188–193. <https://doi.org/10.1016/j.procir.2024.08.322>
- Ranti, C., Jones, W., Klin, A., & Shultz, S. (2020). Blink rate patterns provide a reliable measure of individual engagement with scene content. *Scientific Reports*, 10, 8267. <https://doi.org/10.1038/s41598-020-64999-x>
- Rieger, T., & Manzey, D. (2022). Human performance consequences of automated decision aids: The impact of time pressure. *Human Factors*, 64(4), 617–634. <https://doi.org/10.1177/0018720820965019>
- Rieger, T., Kugler, L., Manzey, D., & Roesler, E. (2023). The (im)perfect automation schema: Who is trusted more, automated or human decision support? *Human Factors*, 66(8), 1995–2007. <https://doi.org/10.1177/00187208231197347>
- Saleem, N., Rihet, M., Sarthou, G., Clodic, A., & Roy, R. N. (2026). Engagement fluctuations during human-robot collaboration: Preliminary assessment using subjective, behavioral, and physiological metrics. In *Extended Abstracts of the 2026 CHI Conference on Human Factors in Computing Systems (CHI EA '26)*. ACM. <https://doi.org/10.1145/3772363.3799347>
- Sandi, C., & Pinelo-Nava, M. T. (2007). Stress and memory: Behavioral effects and neurobiological mechanisms. *Neural Plasticity*, 2007, 78970. <https://doi.org/10.1155/2007/78970>
- Schiffman, H. R. (2001). *Sensation and perception: An integrated approach* (5th ed.). Wiley.
- Schmidhuber, J., Schlogl, S., & Ploder, C. (2021). Cognitive load and productivity implications in human-chatbot interaction. In *Proceedings of the 2021 IEEE 2nd International Conference on Human-Machine Systems (ICHMS)*. IEEE. <https://doi.org/10.1109/ICHMS53169.2021.9582445>
- Seabold, S., & Perktold, J. (2010). Statsmodels: Econometric and statistical modeling with Python. *Proceedings of the 9th Python in Science Conference*, 57–61.
- Shirakura, N., Takase, R., Yamanobe, N., Domae, Y., & Ogata, T. (2022). Time pressure based human workload and productivity compatible system for human-robot collaboration. In *Proceedings of the 2022 IEEE 18th International Conference on Automation Science and Engineering (CASE)* (p. 659). IEEE. <https://doi.org/10.1109/CASE49997.2022.9926685>
- Shum, H.-Y., He, X.-D., & Li, D. (2018). From Eliza to XiaoIce: Challenges and opportunities with social chatbots. *Frontiers of Information Technology & Electronic Engineering*, 19(1), 10–26. <https://doi.org/10.1631/FITEE.1700826>

- Srivastava, S. C., Chandra, S., & Shirish, A. (2015). Technostress creators and job outcomes: Theorising the moderating influence of personality traits. *Information Systems Journal*, 25(4), 355–401. <https://doi.org/10.1111/isj.12067>
- Stock-Homburg, R., & Nguyen, M. A. (2023). See no evil, hear no evil: How users blindly overrely on robots with automation bias. In *Proceedings of the 44th International Conference on Information Systems (ICIS 2023)*. AIS Electronic Library. https://aisel.aisnet.org/icis2023/adv_info_systems_res/adv_info_systems_res/3/
- Sussman, R. F., & Sekuler, R. (2022). Feeling rushed? Perceived time pressure impacts executive function and stress. *Acta Psychologica*, 229, 103702. <https://doi.org/10.1016/j.actpsy.2022.103702>
- Tausch, A., & Kluge, A. (2020). The best task allocation process is to decide on one's own: Effects of the allocation agent in human–robot interaction on perceived work characteristics and satisfaction. *Cognition, Technology & Work*, 24(1), 39–55. <https://doi.org/10.1007/s10111-020-00656-7>
- Urasaki, M., & Niitsuma, M. (2021). Mental fatigue evaluation using blink frequency and pupil diameter on human-robot collaboration with different motion timing. In *2021 IEEE 30th International Symposium on Industrial Electronics (ISIE)*. IEEE. <https://doi.org/10.1109/ISIE45552.2021.9576273>
- Vallat, R. (2018). Pingouin: Statistics in Python. *Journal of Open Source Software*, 3(31), 1026. <https://doi.org/10.21105/joss.01026>
- van Dijk, W., Baltrusch, S. J., Dessers, E., & de Looze, M. P. (2023). The effect of human autonomy and robot work pace on perceived workload in human-robot collaborative assembly work. *Frontiers in Robotics and AI*, 10, 1244656. <https://doi.org/10.3389/frobt.2023.1244656>
- Van Pinxteren, M. M. E., Pluymaekers, M., & Lemmink, J. G. A. M. (2020). Human-like communication in conversational agents: A literature review and research agenda. *Journal of Service Management*, 31(2), 203–225. <https://doi.org/10.1108/JOSM-06-2019-0175>
- Verhagen, R. S., Neerincx, M. A., & Tielman, M. L. (2022). The influence of interdependence and a transparent or explainable communication style on human-robot teamwork. *Frontiers in Robotics and AI*, 9, Article 993997. <https://doi.org/10.3389/frobt.2022.993997>
- Vänni, K. J., Salin, S. E., Cabibihan, J.-J., & Kanda, T. (2019). Robostress, a new approach to understanding robot usage, technology, and stress. In M. A. Salichs et al. (Eds.), *Social Robotics. ICSR 2019. Lecture Notes in Computer Science*, vol. 11876 (pp. 515–524). Springer. https://doi.org/10.1007/978-3-030-35888-4_48
- Wang, S., Fatima, N., Shahbaz, M., & Asif, M. (2026). Building user trust in AI chatbots for customer service through human-like cues and perceived reliability. *Scientific Reports*, 16, Article 7860. <https://doi.org/10.1038/s41598-026-38179-2>
- Wang, S., Yan, Q., & Wang, L. (2023). Task-oriented vs. social-oriented: Chatbot communication styles in electronic commerce service recovery. *Electronic Commerce Research*, 25, 1793–1825. <https://doi.org/10.1007/s10660-023-09741-1>

Wang, Y.-Y., & Wang, Y.-S. (2022). Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. *Interactive Learning Environments*, 30(4), 619–634.

Xu, Y., Zhang, J., & Deng, G. (2022). Enhancing customer satisfaction with chatbots: The influence of communication styles and consumer attachment anxiety. *Frontiers in Psychology*, 13, 902782. <https://doi.org/10.3389/fpsyg.2022.902782>

Zielonka, J. T. (2022). The impact of trust in technology on the appraisal of technostress creators in a work-related context. In *Proceedings of the 55th Hawaii International Conference on System Sciences (HICSS)* (p. 5881). <https://hdl.handle.net/10125/80055>

Zielonka, J. T., & Rothlauf, F. (2021). Techno-eustress: The impact of perceived usefulness and perceived ease of use on the perception of work-related stressors. In *Proceedings of the 54th Hawaii International Conference on System Sciences (HICSS)*. <https://doi.org/10.24251/HICSS.2021.785>

Złotowski, J., Yogeewaran, K., & Bartneck, C. (2017). Can we control it? Autonomous robots threaten human identity, uniqueness, safety, and resources. *International Journal of Human-Computer Studies*, 100, 48–54. <https://doi.org/10.1016/j.ijhcs.2016.12.008>

