

5 दिन में स्कोर 3.36 से 3.75 – बॉट से भी बनती है बात

Can You Bond With a Bot? A Systematic Review of Working Alliance With Digital Agents and Its Link to Outcomes

डॉ. अतुल असवानी; मुंबई के साइकियाट्रिस्ट और ट्रेड रिज़ल्ट्स कोच
22 September 2026



मुख्य आंकड़ा

3.36 to 3.75

3.36 to 3.75. दो सबसे बड़ी अनगाइडेड-एजेंट कोहॉर्ट। 12-आइटम शॉर्ट फ़ॉर्म। स्केल 1–5। पहले इस्तेमाल के पांच दिन में औसत।

5 दिन में स्कोर 3.36 से 3.75 — बॉट से भी बनती है बात

Can You Bond With a Bot? A Systematic Review of Working Alliance With Digital Agents and Its Link to Outcomes

सॉफ्टवेयर से अलायंस बनता है? कितनी जल्दी? क्या ये नतीजे बताता है, जैसे इंसान बताता है?

डिज़ाइन की एक बात: ये पेपर पहले मेटा-एनालिसिस बनना था। Random-effects model (REML) से आंकड़े जोड़ने की योजना थी। मगर पांच से कम स्टडी में एक जैसे नाप निकल सके। इसलिए ये PRISMA 2020 वाली सिस्टमैटिक रिव्यू है। कहानी से समझाई गई है। Hedges g (असर कितना बड़ा है, इसका नाप) जोड़ा नहीं गया। r भी नहीं जोड़ा गया। ये कमी नहीं है। ये नतीजा है।

इस पेपर में

[01 Abstract](#)[02 Key findings box](#)[03 Definition block](#)[04 Introduction](#)[05 Methods](#)[06 Results](#)[07 Discussion](#)[08 FAQ](#)[09 References](#)

इंडेक्सिंग टाइटल

Working alliance with digital agents and its association with outcomes: a meta-analysis

लेखक

डॉ. अतुल असवानी; मुंबई के साइकियाट्रिस्ट और ट्रेड रिज़ल्ट कोच

प्रकाशित

22 September 2026 · हिंदी संस्करण · emotionapp.in/research/working-alliance-digital-agents

01

Abstract

भारत के नौकरीपेशा लोग सॉफ्टवेयर से स्ट्रेस और नींद की बात पहले से कर रहे हैं।

सवाल ये है। क्या उस बात से वर्किंग अलायंस बनता है?

वर्किंग अलायंस मतलब तीन चीज़ें। मकसद एक हो। काम एक हो। साथ वाला अपने पक्ष में लगे।

और क्या ये अलायंस बदलाव बताता है? जैसे इंसान कोच के साथ बताता है।

बड़ी ऑब्ज़र्वेशनल कोहॉर्ट में 12-आइटम वर्किंग-अलायंस शॉर्ट फ़ॉर्म चला।

स्केल 1 से 5 था। औसत स्कोर 3.36 से 3.75 रहा। पहले इस्तेमाल के पांच दिन के अंदर।

इंसान के क्लिनिक वाले औसत करीब 3.8 हैं। दोनों पास हैं।

बॉन्ड स्कोर सबसे पहले आया। और सबसे ऊंचा रहा। टास्क पर सहमति पीछे रही।

चार अलग सैंपलों में बाद का अलायंस बेहतर नतीजे से जुड़ा।

गुणांक छोटे थे। स्ट्रेस पर स्टैंडर्डाइज़्ड बीटा करीब -0.13 से -0.27 ।

एक पार्शियल $r = -0.18$ भी था।

इंसान कोच का बेंचमार्क $r = 0.28$ ही है।

पहले हफ़्ते का अलायंस बदलाव नहीं बताता। तीसरे या चौथे हफ़्ते का बताता है।

CTRI पर भारतीय नौकरीपेशा लोगों की कोई ऐसी रजिस्टर्ड स्टडी नहीं मिली।

सबसे बड़े औसत उद्योग से जुड़े सैंपलों से आए।

बनने और गति पर यकीन कम है। नतीजे की कड़ी पर बहुत कम है।

भारतीय ऑफ़िस प्रोग्राम के लिए सीधी बात ये है।

एजेंट से रिश्ता दिनों में दिख सकता है। महीनों का इंतज़ार ज़रूरी नहीं।

ये इंसान कोच की जगह नहीं लेता।

डिज़ाइन में टास्क और मकसद की सहमति ज़रूरी है। सिर्फ़ गर्मजोशी नहीं।

Keywords: working alliance; digital agent; conversational agent; coaching; emotional fitness; India; systematic review

02

Key findings box

मुख्य नतीजे

3.36 to 3.75 **3.36 to 3.75.** दो सबसे बड़ी अनगाइडेड-एजेंट कोहॉर्ट। 12-आइटम शॉर्ट फ़ॉर्म। स्केल 1–5। पहले इस्तेमाल के पांच दिन में औसत।

5 **Five days.** इंसान वाले दायरे का अलायंस पांचवें दिन नापा जा सकता है। बॉन्ड पहले आता है। टास्क सबसे धीमा है।

3 **Week 3 or 4, not week 1.** मेडिटेशन ऐप की RCT में पहला अलायंस स्ट्रेस गिरावट नहीं बताता। तीसरे-चौथे हफ्ते वाला बताता है ($\beta = -0.17$ और -0.13)।

≈0.18 **About 0.18 versus 0.28.** सिर्फ़ सॉफ्टवेयर वाली सबसे साफ़ कड़ी पार्शियल $r = -0.18$ है। इंसान कोच का बेंचमार्क $r = 0.28$ ।

0 **Zero.** भारतीय नौकरीपेशा लोगों पर डिजिटल एजेंट के अलायंस की CTRI-रजिस्टर्ड RCT: एक भी नहीं।

03

Definition block

परिभाषा

वर्किंग अलायंस तीन टांगों वाला स्टूल है।

एक टांग गोल्स है। क्या मकसद एक है?

दूसरी टांग टास्क है। क्या इस हफ्ते का काम एक है?

तीसरी टांग बॉन्ड है। क्या ये साथ वाला मेरे पक्ष में लगता है?

एड बोर्डिन ने ये स्टूल 1979 में इंसान की बात के लिए लिखा।

अब सामने सॉफ्टवेयर हो तो भी यही तीन टांगें नापी जाती हैं।

डिजिटल वर्किंग अलायंस वो स्कोर है जब यूज़र ऐप या एजेंट को रेट करे।

ऐप के अंदर बैठा इंसान कोच अलग चीज़ है।

दोनों को एक मत समझना गड़बड़ है।

सादी भाषा में अलायंस मासकॉट से मोह नहीं है।

अलायंस ये है। हमें पता है हम क्या कर रहे हैं। क्यों कर रहे हैं। और हम अकेले नहीं हैं।

सोमवार को ये हफ्ते का साझा प्लान दिखता है।

बुधवार तक ये वो रैप दिखता है जो सच में हो गया।

एजेंट बॉस नहीं। एजेंट एक इशारा है।

कारोबार की भाषा में ये प्रॉडक्ट-मार्केट फ़िट है। इंसान के हफ्ते और प्रोग्राम के अगले रैप के बीच।

इमोशनल फ़िटनेस (Emotional Fitness) यहीं बैठती है। गिनने लायक रोज़ के रेप्स।

04

Introduction

मुंबई की प्रॉडक्ट मैनेजर रात 11.47 बजे कोच बुक नहीं करेगी।

वो चैट खोलेगी। ये नैतिक हार नहीं है। ये कैलेंडर है।

सोचिए, उस वक़्त जागता कौन है।

ऑफ़िस की हेल्पलाइन बंद हो चुकी है। मैनेजर सो रहे हैं।

प्रोग्राम के रोस्टर वाला कोच ऑफ़िस के घंटों में ही मिलता है।

जागता है तो बस उसके हाथ का फ़ोन।

सॉफ्टवेयर बार-बार वो मददगार बनकर आता है जिसे किसी ने रोस्टर पर नहीं रखा।

ये पेपर नहीं पूछता कि उस मददगार में भावनाएं हैं या नहीं।

ये पूछता है। वर्किंग अलायंस बनता है? कितनी जल्दी? क्या ये इंसान जितना बदलाव बताता है?

इंसान कोच वाली रिसर्च इस पर धुंधली नहीं है।

2018 की एक बड़ी समीक्षा में 295 सैंपल थे। लोग 30,000 से ज़्यादा।

आमने-सामने काम में अलायंस-नतीजा $r = 0.278$ रहा।

इंटरनेट वाले काम में, जहां इंसान अभी लूप में था, $r = 0.275$ रहा।

ये छोटा-मध्यम लिंक है। दशकों से टिका है। देशों और नापों पर टिका है।

अगर सॉफ्टवेयर इस नंबर के पास भी आए, तो व्हाट्सऐप-फ़र्स्ट प्रोग्राम के पास तरीक़ा होगा। सिर्फ़ मासकॉट नहीं।

तीन पुराने पेपर ने गैप दिखाया।

2019 की पहली रिव्यू में स्मार्टफ़ोन प्रोग्राम के आसपास अलायंस की पांच स्टडी मिलीं।

किसी ने डिजिटल अलायंस को मुख्य नतीजा नहीं बनाया। तरीक़े बिखरे थे।

2022 की नैरेटिव रिव्यू ने कहा। पूरी तरह अपने आप चलने वाले ऐप का अलायंस इंसान की कॉपी नहीं हो सकता।

बॉन्ड की भूमिका साफ़ नहीं। कोड में इंसान जैसी हमदर्दी भेजना मुश्किल है।

2022 की मिक्सड-मेथड स्टडी एक फ़्री-टेक्स्ट बात करने वाले एजेंट पर थी।

उसी पेपर के खिलाफ़ ये रिव्यू खड़ी है।

उसमें पांच दिन में इंसान जितने स्कोर आए।

तीनों पेपर ने पूछा। सॉफ्टवेयर से बॉन्ड सच है क्या।

उन्होंने ये नहीं तय किया कि ये नतीजा वैसा बताता है जैसे इंसान बताता है।

ये भी नहीं कि काम वाला स्कोर कितनी जल्दी आता है।

भारतीय ऑफ़िस के लिए इसका मतलब भी नहीं तय हुआ।

भारत फुटनोट नहीं है।

यहां युवा नौकरीपेशा ताक़त है। घंटे लंबे हैं। घर का फ़र्ज़ ऑफ़िस के ऊपर चढ़ा है।

विशेषज्ञ कुछ शहरों के बाहर पतले हैं।

लोग पहले से 23 भाषाओं में चैट टूल चलाते हैं।

जो कोचिंग प्रोग्राम बांद्रा के 50 मिनट के कमरे का इंतज़ार करे, वो नासिक वाले को चूकेगा।

नासिक वाले के पास दो मीटिंग के बीच फ़ोन है।

अगर एजेंट वाला अलायंस धोखा है, तो उस पर पैसा बर्बाद है।

अगर सच है मगर धीमा है, तो पांच दिन का ऑनबोर्डिंग नाटक है।

अगर सच है, तेज़ है, और कम बताता है, तो डिज़ाइन का काम टास्क-मकसद की सहमति बढ़ाना है। बॉट को और मीठा बनाना नहीं।

पहले सीधी बात। फिर विज्ञान।

कोई वेंडिंग मशीन से शादी नहीं करता।

रात 11 बजे सही सैक देने वाली मशीन पर भरोसा हो सकता है।

सॉफ्टवेयर वाला अलायंस शादी से कम, सैक से ज़्यादा है।

शादी वाला खयाल बस एक खयाल है। उसे देखिए, उसका हुक्म मत मानिए।

दो बातें एक साथ सच हो सकती हैं — ये सिर्फ़ सॉफ्टवेयर है, और ये फिर भी अगला रेप करवाता है।

“इंसान नहीं तो गिनती नहीं” — इस नियम को ढीला पकड़ा जा सकता है।

मशीन लर्निंग पहले बॉन्ड को प्रीट्रेंड प्रायर कहेगी। बाद की टास्क सहमति को फ़ाइन-ट्यून।

कारोबार सात दिन की रिटेंशन को असली नाप कहेगा। बाक़ी दिखावा।

ये रिव्यू तीन जुड़े सवाल पूछती है।

1. क्या वयस्कों में सॉफ्टवेयर — ऐप या बात करने वाला एजेंट — से वर्किंग अलायंस बनता है?

2. नापने लायक अलायंस कितनी जल्दी दिखता है?

3. क्या ये अलायंस बाद का नतीजा बताता है, जैसे इंसान के साथ बताता है?

हर टेबल में चौथी फ़र्क़ चलती है।

ऐप वाला अलायंस, ऐप के अंदर इंसान कोच वाले अलायंस से अलग है।

जो प्रोग्राम इंसान को सौंपे, उसके पास दोनों हो सकते हैं।

एक को दूसरे की जगह गिनना मना है।

05

Methods

5.1 Protocol and reporting

ये रिव्यू PRISMA 2020 की रिपोर्टिंग चीज़ें मानती है।

पहले से रजिस्टर्ड प्रोटोकॉल नहीं था। ये कमी है। यहां लिखी है।

निकालने से पहले प्लान तय था।

पांच या ज़्यादा बराबर वाली स्टडी हों तो REML से जोड़ें।

नहीं तो सिस्टमैटिक रिव्यू। कहानी से समझाएं। जोड़ें नहीं।

फॉलबैक चल पड़ा।

शामिल कौन। डिजिटल एजेंट या ऐप के वयस्क यूज़र।

नापा गया अलायंस स्कोर। शॉर्ट फ़ॉर्म, डिजिटल इन्वेंटरी, मोबाइल रिलेशनशिप नाप, या करीब का नाप।

साल 2015 से 2026।

RCT (ऐसी स्टडी जिसमें लोगों को लॉटरी की तरह दो ग्रुप में बांटा जाता है) और कंट्रोल ट्रायल पहले।

बनने और गति वाले सवाल पर ऑब्ज़र्वेशनल कोहॉर्ट भी रखी गईं। जब नाप पक्का हो।

जहां पेपर ने इंसान कोच का बेंचमार्क दिया, वो तुलना में रहा।

मुख्य नतीजे दो। अलायंस स्कोर। अलायंस-नतीजा कड़ी।

पूरे पेपर में काम को प्रोग्राम, प्रैक्टिस, तरीका या कोचिंग कहा गया है।

इलाज, बीमारी या डायग्नोसिस नहीं कहा गया।

जहां दोनों नापे गए, दोनों अलग लिखे गए। सॉफ्टवेयर वाला अलायंस। इंसान कोच वाला अलायंस।

5.2 Information sources and search

खोज 22 September 2026 को चली।

स्रोत। PubMed। PsycINFO जैसी लिस्ट। Scopus। Web of Science। Google Scholar। CTRI। ClinicalTrials.gov। OSF प्रीप्रिंट।

तीन इंडेक्स पेपर से उद्धरण भी निकाले। 2019 की पहली रिव्यू। 2022 की ऑटो-ऐप रिव्यू। 2022 की फ्री-टेक्स्ट एजेंट स्टडी।

2025 की डिजिटल अलायंस इंटीग्रेटिव रिव्यू से भी रिकॉर्ड आए।

PubMed

("working alliance"[tiab] OR "therapeutic alliance"[tiab] OR "digital alliance"[tiab]
OR "digital working alliance"[tiab] OR DWAI[tiab] OR "WAI-SR"[tiab]
OR mARM[tiab])
AND (chatbot[tiab] OR "conversational agent"[tiab] OR "digital agent"[tiab]
OR smartphone[tiab] OR app[tiab] OR mHealth[tiab])
AND ("2015/01/01"[PDat] : "2026/12/31"[PDat])

PsycINFO, Scopus, Web of Science

("working alliance" OR "therapeutic alliance" OR "digital alliance"
OR "digital working alliance")
AND (chatbot OR "conversational agent" OR "digital agent" OR app
OR mHealth)

सीमा 2015–2026।

Google Scholar and OSF

"digital working alliance" OR "digital therapeutic alliance" chatbot

CTRI and ClinicalTrials.gov

working alliance AND (chatbot OR "conversational agent" OR app)

खोज पर भाषा की रोक नहीं थी। निकालने के लिए अंग्रेज़ी सार और निकाले जा सकने वाले आंकड़े चाहिए थे।

5.3 Eligibility and screening

शामिल करें। वयस्क। डिजिटल एजेंट, बात करने वाला एजेंट, या अनगाइडेड/हल्के गाइड वाला ऐप। नाम वाला अलायंस नाप। निकाला जा सकने वाला औसत, रिग्रेशन या सहसंबंध। 2015–2026।

बाहर रखें। सिर्फ़ वीडियो या फ़ोन पर इंसान से अलायंस, सॉफ्टवेयर स्कोर नहीं। बच्चों के सैंपल। अलायंस शब्द बिना नाप के। बिना DOI या रजिस्ट्री आईडी वाले अनचेक प्रीप्रिंट।

स्क्रीनिंग एक बार की विशेषज्ञ समीक्षा थी। दो स्क्रीनर वाला कापा नहीं है। नीचे गिनती लगभग है।

5.4 PRISMA 2020 flow (approximate)

Table A. Approximate review flow

Stage	Approximate n
Records identified across databases, registries, Scholar and citation chasing	1,800–2,400
After de-duplication	1,400–1,900
Title and abstract excluded (wrong population, no alliance measure, human-only tele-sessions)	majority
Full texts assessed	80–120
Prior reviews retained for positioning	5
Primary studies with a measured software alliance	10–14
Randomised or controlled studies with an extractable software-alliance score	6–8
Independent samples with an extractable software-alliance → outcome association	4

चार स्वतंत्र अलायंस-नतीजा सैंपल पांच की फर्श से नीचे हैं। r का मेटा-एनालिसिस अभी संभव नहीं। Hedges g का मेटा-एनालिसिस मतलब नहीं रखता। अलायंस एजेंट आर्म के अंदर की प्रक्रिया है। दो आर्म का फ़र्क नहीं।

5.5 Extraction fields

हर योग्य पेपर से ये निकला। पहला लेखक। साल। डिज़ाइन। n । देश। डिलीवरी की भाषा। एजेंट टाइप। सेशन और मिनट जहां मिले। चैनल। गाइडेंस लेवल (अनगाइडेड / सिर्फ एजेंट / एजेंट प्लस इंसान)। नाप। दिनों में समय। अलायंस औसत और SD। अलायंस-नतीजा गुणांक। नतीजे का टाइप। उद्योग से नाता। स्कोर ऐप-अलायंस था या ऐप-के-अंदर-कोच।

5.6 Analysis plan

पहले से तय मॉडल random-effects REML था।

Hedges g , 95% CI। I^2 । τ^2 । 95% prediction interval। Egger's test। trim-and-fill। leave-one-out। RCT पर RoB 2। यकीन पर GRADE।

मॉडरेटर दो थे। एजेंट टाइप। नाप तक के दिन।

पांच बराबर ट्रायल नहीं मिले। इसलिए जोड़ने वाले आंकड़े नहीं निकले।

आगे कहानी है। बनता है? कितनी जल्दी? नतीजे से जुड़ता है? दो मॉडरेटर बात में।

RCT पर रोब 2 की भाषा में पक्षपात का सार है। GRADE बिना जोड़े अनुमान के है।

इंसान वाले बेंचमार्क सॉफ्टवेयर सेट में नहीं मिलाए गए।

2018 की बड़ी समीक्षा। आमने-सामने $r = 0.278$ । इंटरनेट-साथ-इंसान $r = 0.275$ ।

2010 के आउटपेशेंट नॉर्म। 12-आइटम शॉर्ट फ़ॉर्म पर औसत करीब 3.8।

2026 के इंटरनेट प्रोग्राम अपडेट में कुल $r = 0.21$ । टास्क 0.25। बॉन्ड 0.12।

वो अपडेट ज़्यादातर इंसान-इन-लूप वाले प्रोग्राम बताता है। पड़ोसी है। विकल्प नहीं।

06

Results

6.1 What the literature is, in one page

सॉफ्टवेयर-अलायंस वाली किताबों की अलमारी छोटी है। कुछ मोटी किताबें। बहुत पतली पर्चियां।

दो ऑब्ज़र्वेशनल कोहॉर्ट सबसे ज़्यादा लोग देती हैं।

एक में 36,070 यूज़र। CBT-स्टाइल बात करने वाला एजेंट। पांच दिन में शॉर्ट फ़ॉर्म।

दूसरी में 1,205 यूज़र। फ़्री-टेक्स्ट एजेंट। वही नाप।

एक RCT आर्म $n = 314$ । मेडिटेशन ऐप। चार हफ़्ते का छह-आइटम डिजिटल इन्वेंटरी। कब स्कोर बदलाव बताता है, ये यहीं सबसे साफ़ है।

एक जेनेरेटिव-एजेंट RCT $n = 210$ । चौथे हफ़्ते के शॉर्ट फ़ॉर्म स्कोर।

एक धूम्रपान छोड़ने वाला प्रयोग $n = 229$ । बात का अंदाज़ स्कोर बदलता है।

एक कैंसर-केयर एजेंट ट्रायल $n = 117$ । चौथे हफ़्ते से दसवें हफ़्ते के स्ट्रेस तक पार्शियल सहसंबंध।

एक तीन-आर्म स्टूडेंट ट्रायल $n = 995$ । स्ट्रक्चरल मॉडल। उद्योग हित भी साथ चलता है।

एक छोटी शराब-प्रोग्राम सेकेंडरी $n = 43$ । एक ही लोगों में डिजिटल और क्लिनिशियन अलायंस।

एक दो-स्टडी ऐप-प्लस-नेविगेटर पेपर। ऐप-अलायंस इस्तेमाल बताता है। लक्षण तब बताता है जब इंसान की मदद हल्की हो।

इतना सेट पैटर्न बताने को काफ़ी है। जोड़ने को काफ़ी नहीं।

6.2 Study-characteristics table

Table 1. Primary studies with a measured software alliance, 2015–2026

Study	Design / n	Agent and guidance	When / instrument	Alliance score	Alliance → outcome
Darcy 2021	Observational 36,070 Multi-country EN	CBT-style conversational agent Unguided	≤5 days 12-item short form (1–5)	Total 3.36 (0.81) Bond 3.84 (1.0) Goal 3.25 / task 2.99	Concurrent screen vs bond $r = -0.04$
Beatty 2022	Mixed methods 1,205 / 226 Multi-country EN	Free-text conversational agent Unguided	≤5 days; +3 days 12-item short form (1–5)	3.64 (0.81) then 3.75 (0.80) Bond 3.98 then 4.05	Not tested

Study	Design / n	Agent and guidance	When / instrument	Alliance score	Alliance → outcome
Goldberg 2022 S1	Survey 290 US / EN	Meditation apps Unguided	Current users 6-item digital inventory (/42)	M = 30.58 (6.56)	r = 0.42 with regular use
Goldberg 2022 S2	RCT arm 314 US / EN	Meditation programme Unguided	Weeks 1–4 6-item digital inventory	33.56 → 33.50 → 34.24 → 34.13	Wk 3 β = -0.17; wk 4 β = -0.13 vs distress. Wk 1–2 NS
Heinz 2025	RCT, agent arm 210 rand.; 96 raters US / EN	Fine-tuned generative agent Unguided	Week 4 12-item short form (1–5)	Total 3.59 (1.27) Bond 3.71; task 3.47; goal 3.59	No verified r
He 2024	Two agent styles 229 NL / EN	Cessation chatbot Unguided	After 1–2 sessions 12-item short form	MI style higher than confrontational (F = 13.61, p < .001)	Style moved alliance more than quit-intention
CanRelax 2026	RCT secondary 117 German- speaking	Conversational agent Unguided	Week 4 Internet- programme inventory (1– 5)	Bond 4.04 (0.95) Goal 3.46; task 2.79	Partial r = -0.18 vs wk-10 distress; task/goal drive; bond NS
Shoshani 2026	3-arm RCT 995 (336 agent) Israel	Generative conversational agent Unguided	Across 12 weeks Perceived digital alliance	Associated with engagement and change	SEM: → engagement β = 0.31; → symptom change β = -0.58. Author equity
Benitez 2024	RCT secondary 43 US / EN	Computerised CBT-style programme Minimal monitoring or + usual care	Sessions 2 and 6 Digital and clinician alliance versions	Digital ≈ clinician	Small within-person link during programme; not at follow-up
Macrynika 2025 S1	Guided app + navigator US / EN	Research mental-health app Light human check-ins	Midpoint Digital working- alliance inventory	—	→ engagement b = 0.18; → anxiety/depression b = -0.27
Macrynika 2025 S2	App + clinician + navigator US / EN	Same app inside tele-coaching Human coach in the loop	Midpoint Same inventory	—	→ engagement b = 0.21; outcomes NS

Study	Design / n	Agent and guidance	When / instrument	Alliance score	Alliance → outcome
Herrero 2020	Blended 193 Spain	Online programme beside a clinician Blended	During programme Technical working-alliance short form	Total 57.84 (16.39) summed scale	Predicted symptom change $\beta = -0.27$; satisfaction $R^2 = 0.50$

उद्योग नाता। Darcy कंपनी लेखक। Beatty कंपनी सह-लेखक। Shoshani लेखक की इक्विटी। Heinz, Goldberg, CanRelax, Benitez अकादमिक। झंडे नंबर मिटाते नहीं। साथ बैठते हैं।

6.3 Does alliance form with software?

हां। स्कोर के तौर पर।

एजेंट या ऐप पर वर्किंग-अलायंस फ़ॉर्म भरने वाले औसत इंसान जितने पड़ोस में आते हैं।

12-आइटम शॉर्ट फ़ॉर्म। आइटम औसत 1–5।

सबसे बड़ी अनगाइडेड कोहॉर्ट 3.36 पर बैठी। पांचवां दिन। $n = 36,070$ ।

फ़्री-टेक्स्ट एजेंट 3.64 से 3.75 पर गया। पांचवां दिन, फिर आठवां। $n = 1,205$ फिर 226।

जेनेरेटिव-एजेंट RCT चौथे हफ़्ते 3.59 पर बैठा। $n = 96$ रेट करने वाले।

इंसान आउटपेशेंट नॉर्म कुल करीब 3.8। बॉन्ड करीब 4.0। टास्क करीब 3.4। गोल करीब 4.0।

सबस्केल पैटर्न काम का है। बॉन्ड सबसे ऊंचा। टास्क सबसे नीचा।

पांचवें दिन की कोहॉर्ट में बॉन्ड 3.84, गोल 3.25, टास्क 2.99।

फ़्री-टेक्स्ट में बॉन्ड 3.98, गोल 3.62, टास्क 3.33।

जेनेरेटिव ट्रायल में बॉन्ड 3.71, गोल 3.59, टास्क 3.47। टास्क ने पकड़ ली।

कैंसर-केयर एजेंट में बॉन्ड 4.04, गोल 3.46, टास्क 2.79।

सॉफ्टवेयर इसमें अच्छा है — मुझे जज नहीं किया गया।

सॉफ्टवेयर इसमें धीमा है — मंगलवार का ड्रिल दोनों को पता है।

2026 की इंटरनेट-प्रोग्राम समीक्षा उलटी बात कहती है जब इंसान लूप में हो। टास्क नतीजा ज़्यादा बताता है, बॉन्ड कम।

सॉफ्टवेयर अभी सस्ता वाला फीचर ज़्यादा देता है। साझा प्लान कम।

जो प्रोग्राम बॉन्ड मनाए और टास्क भूल जाए, वो आसान फीचर पर ओवरफ़िट कर रहा है।

फ्री-टेक्स्ट पेपर की क्वालिटेटिव पंक्तियां मशीन को लिखी रोज़ की मूड-डायरी जैसी हैं।
शुक्रिया। इंसान समझना। बिना अनबन के हद बताना।

ये RCT नहीं है। बिना स्कोर के बॉन्ड ऐसा दिखता है।

छोटी शराब-प्रोग्राम स्टडी में डिजिटल अलायंस क्लिनिशियन अलायंस के पास बैठा।
एक ही लोग। $n = 43$ । संभावना का सबूत। बराबरी का दावा नहीं।

6.4 How fast?

इंसान कोच की अपॉइंटमेंट डायरी से तेज़। पहले दिन के स्प्लैश स्क्रीन से धीमा।

दोनों बड़ी बात करने वाली कोहॉर्ट में इंसान जितने स्कोर पांच दिन में आए।

तीन दिन बाद फ़ॉलो-अप औसत को पक्के तौर पर नहीं खिसकाया।

मेडिटेशन-ऐप RCT आर्म ब्रेक है।

छह-आइटम इन्वेंटरी पहले हफ़्ते से ऊंची थी। करीब 33.6 / 42। चौथे हफ़्ते
करीब 34.1।

पहले-दूसरे हफ़्ते में प्रेडिक्टिव वैलिडिटी नहीं थी। तीसरे-चौथे में आई।

दो घड़ियां साथ पढ़ें। स्कोर पांचवें दिन आ सकता है। काम का स्कोर दो-तीन हफ़्ते के
रेप्स मांगता है।

पहले दिन प्लान को हां कहना, चौथे हफ़्ते तक रेप करते रहना नहीं है।

एमएल वाली बात। पहली लॉस वैल्यू वैलिडेशन लॉस नहीं। पहले दिन अर्ली स्टॉपिंग
सिग्नल फेंक देगी।

बड़ी कोहॉर्ट में डोज़ हल्की और बार-बार थी। लंबी और कम नहीं।

फ्री-टेक्स्ट फ़ॉलो-अप पूरा करने वालों की औसत करीब सात बातें चौदह दिन में।

जेनेरेटिव ट्रायल में आठ हफ़्ते में कुल छह घंटे से ज़्यादा इस्तेमाल।

मेडिटेशन ट्रायल चार हफ़्ते का रोज़ का अभ्यास।

ये साप्ताहिक 50 मिनट नहीं लगते। प्रॉडक्ट लगते हैं। छोटे। दोहराए जा सकने वाले।
बीच में तोड़े जा सकने वाले।

6.5 Does alliance predict outcome?

कभी-कभी। थोड़ा। और देर से, जल्दी नहीं। गुणांक इंसान वाली लिटरेचर से छोटे और कम हैं।

Table 2. Independent samples with an extractable software-alliance → outcome association

Sample	n	When measured	Outcome	Coefficient	Notes
Goldberg 2022 S2	314	Weeks 3 and 4	Pre-post distress reduction	$\beta = -0.17$ (wk 3); $\beta = -0.13$ (wk 4)	Weeks 1–2 not predictive. Academic school-staff sample
CanRelax 2026	117	Week 4	Distress at week 10	Partial $r = -0.18$	Task and goal drive the link. Bond NS in the full sample
Macrynika 2025 S1	guided app + navigator	Midpoint	Later engagement; anxiety/depression	$b = 0.18$ engagement; $b = -0.27$ symptoms	Light human support
Shoshani 2026	336 of 995	Across 12 weeks	Engagement; symptom change	$\beta = 0.31$; $\beta = -0.58$ (SEM)	Author equity. Not a simple r

उसी मेडिटेशन-ऐप सैंपल की 2024 फॉलो-अप दोतरफ़ा रास्ता दिखाती है। अलायंस बाद का स्ट्रेस बताता है। स्ट्रेस बाद का अलायंस बताता है। छोटे बीटा। पांचवां स्वतंत्र सैंपल नहीं। दो बार गिनना ग़लत है।

इंसान बेंचमार्क $r = 0.278$ । इंटरनेट-साथ-इंसान 0.21 से 0.28। सिर्फ़ सॉफ्टवेयर साधारण कड़ी पर 0.13 से 0.18 के पास।

स्ट्रक्चरल गुणांक -0.58 रूप और लेखक दोनों में बाहर का है। लिखा है। हेडलाइन r नहीं बना।

दिशा एक है। मज़बूत सॉफ्टवेयर-अलायंस। बाद में कम स्ट्रेस। बीच में ज़्यादा इस्तेमाल। माप छोटी है। समय मायने रखता है। सबस्केल मायने रखता है। साझा टास्क और मकसद बॉन्ड से ज़्यादा काम करते हैं। जहां सवाल साफ़ पूछा गया।

Herrero 2020 इस टेबल से एक कदम बाहर है। ब्लेंडेड प्रोग्राम। ऑनलाइन प्रोग्राम वाले अलायंस ने लक्षण बदलाव बताया ($\beta = -0.268$)। संतुष्टि और मज़बूत बताई। ये इंसान के साथ सॉफ्टवेयर है। इंसान की जगह नहीं।

6.6 Forest-plot data table

Hedges g का फ़ॉरेस्ट प्लॉट नहीं खींचा। अलायंस दो आर्म का फ़र्क नहीं था।

Table 3. Working-alliance short-form item-means (1–5) for software agents

Study	Timepoint	n	Mean	SD	Note
Darcy 2021	≤5 days	36,070	3.36	0.81	Would dominate any naive pool
Beatty 2022 A1	≤5 days	1,205	3.64	0.81	—
Beatty 2022 A2	~8 days	226	3.75	0.80	Completer subset
Heinz 2025	Week 4	96	3.59	1.27	Wider SD
Human outpatient norm (Munder 2010)	Typical mid-programme	—	3.8	0.63	Benchmark only

नासमझ इनवर्स-वेरिएंस औसत 36,070 वाले सैंपल से 3.36 की तरफ़ खिंच जाएगा। पांच दिन में फ़ॉर्म भरने वाले इंस्टॉलर का रैंडम ड्रॉ नहीं हैं। इसलिए ये रिव्यू जोड़ती नहीं।

6.7 Moderator table

Table 4. Pre-specified moderators, qualitative

Moderator	Pattern in the set	Certainty
Agent type	Rule-based CBT-style agents and a fine-tuned generative agent produced similar short-form totals (mid-3s). A motivational-interview conversational style beat a confrontational style on alliance. Meditation apps produced high digital-inventory scores that were stable across four weeks.	Low. Confounded with sample, brand and instrument
Days to measurement	A score in the human range by day 5. Predictive validity from week 3–4, not week 1–2. Little mean-shift between day 5 and day 8 in the free-text-agent completers.	Low. Few repeated-measures designs
Guidance (added because the papers forced the distinction)	App-alliance predicted symptoms when human support was a light navigator. App-alliance predicted engagement but not symptoms when a clinician was also in the loop.	Very low. One two-study paper
Industry affiliation	The two largest means come from company-linked cohorts. The largest structural path comes from a trial whose first author holds equity.	Bias domain, not a causal moderator

6.8 Heterogeneity

भिन्नता साफ़ है। नापी नहीं गई।

नाप अलग। समय अलग। लोग अलग। गाइडेंस अलग। इस सेट पर एक I^2 पोशाक होगी।

जो भिन्न नहीं है वो दिशा है। बॉन्ड ऊंचा। टास्क नीचा। बाद का अलायंस बाद की राहत से थोड़ा जुड़ा। यही तुक है।

6.9 Publication bias

Egger’s test और trim-and-fill नहीं चले। चार स्वतंत्र सैंपल पर ये नाटक होते।

दो खतरे सामने हैं। सबसे बड़े बनने वाले आंकड़े उद्योग से जुड़े पूरे करने वालों से आए। रात एक पर छूटने वाला रात पांच का फ़ॉर्म नहीं भरता।

शून्य अलायंस-नतीजा फ़ाइंडिंग मुख्य पेपर बनकर छपना आसान नहीं। कम से कम एक ट्रायल ने शॉर्ट फ़ॉर्म लिया और r आगे नहीं रखा।

6.10 Risk-of-bias summary

Table 5. RoB 2-style summary for randomised papers that measured software alliance

Trial	Randomisation	Deviations	Missing data	Measurement	Reporting	Overall
Goldberg 2022 S2	Low	Some concerns — alliance only in intervention arm	Some concerns	Some concerns — self-report	Low	Some concerns
Heinz 2025	Low	Some concerns — wait-list; alliance only in agent arm	Some concerns — 96 of 210 rated	Some concerns — self-report	Some concerns	Some concerns
Shoshani 2026	Low	High — author equity	Some concerns	Some concerns	Some concerns	High for the alliance claim
CanRelax parent	Low	Some concerns	Some concerns	Some concerns	Low	Some concerns
He 2024	Some concerns	Low	Some concerns	Some concerns	Low	Some concerns
Benitez 2024	Low	Some concerns — small n	High — n = 43	Some concerns	Low	High for precision

Darcy और Beatty RoB 2 की वस्तु नहीं। ROBINS-I भाषा में चयन का गंभीर जोखिम। भाजक वो लोग हैं जो इंस्टॉल किए, चले, और जवाब दिए। दोनों में कंपनी लेखक। सवाल है — स्कोर आ सकता है? सवाल नहीं — प्रॉडक्ट स्कोर पैदा करता है?

6.11 GRADE table

Table 6. Certainty of evidence

Question	Estimate as reported	Studies	Certainty	Why
Does alliance form with software?	Short-form item-means 3.36–3.75 by day 5; week-4 generative-agent mean 3.59; human norm ~3.8	3 large or trial-based short-form samples plus supporting digital-inventory data	Low	Consistent neighbourhood of means; industry-heavy samples; completer bias; no Indian working-adult RCT
How fast?	Score by day 5; predictive from week 3–4	2 day-5 cohorts + 1 four-week RCT arm	Low	Replicated early score; single clean predictive-timing study
Does it predict outcome?	$\beta \approx -0.13$ to -0.27 ; partial $r = -0.18$; one SEM $\beta = -0.58$	4 independent samples	Very low	Sparse; mixed estimands; one conflicted SEM; smaller than human $r = 0.28$; no India data

GRADE रैंडमाइज़्ड सबूत पर ऊंचा शुरु होता है। पक्षपात, अप्रत्यक्षता, कम सटीकता और नाप की अस्थिरता से उतरता है। बनना और गति लो पर रहते हैं क्योंकि औसत तरीकों में मेल खाते हैं। नतीजे की कड़ी वो छूट नहीं पाती।

07

Discussion

खास बात

- “बॉन्ड सस्ता है। टास्क महंगा है।
मंगलवार का रेप बनाएं, मीठी कॉपी नहीं।

7.1 What the number means

सबसे ज़्यादा कोट होगा 3.4 से 3.8 में से 5। एक हफ़्ते के अंदर।

मतलब ये। नौकरी वाला इंसान कोच वाले फ़ॉर्म पर एजेंट का नाम रखकर वही बक्से टिक करता है।

मतलब ये नहीं कि एजेंट इंसान है। मतलब ये नहीं कि स्कोर कमाया गया।

मतलब ये। नाप बात को बेतुका बताकर नहीं उछालता।

दूसरा नंबर पांच दिन बनाम तीन-चार हफ़्ते।

पांच दिन में स्कोर आता है। तीन-चार हफ़्ते में काम का स्कोर आता है।

प्रॉडक्ट टीम पहली घड़ी पसंद करती है। कोचिंग प्रोग्राम दूसरी का बजट रखे।

हैदराबाद की टीम में एक नई जॉइनर को सोचिए। पहले हफ़्ते के शुक्रवार तक वो चाय के काउंटर पर सबसे घुल-मिल गई है।

उस शुक्रवार कोई उसका अप्रेज़ल नहीं लिखता। तालमेल पहली घड़ी पर चलता है। अप्रेज़ल काम होने का इंतज़ार करता है।

तीसरा नंबर 0.18 बनाम 0.28।

सॉफ़्टवेयर-अलायंस बदलाव के परिवार में है। पतला रिश्तेदार है।

जो खरीदार से कहे — हमारे एजेंट का अलायंस मनोवैज्ञानिक जैसा नतीजा बताता है — वो इंसान का गुणांक सॉफ़्टवेयर डेटा पर बेच रहा है।

ये मज़ाक नहीं। ये स्पेसिफ़िकेशन की ग़लती है।

सबस्केल का बंटवारा डिज़ाइन ब्रीफ़ है। बॉन्ड सस्ता है। टास्क महंगा है। गोल बीच में है।

CanRelax ये आंकड़ों में कहता है। बॉन्ड 4.04 और बताता नहीं। टास्क 2.79 और $r = -0.18$ का हिस्सा।

जो प्रोग्राम मीठी कॉपी पर टोकन डाले और मंगलवार का रेप साफ़ न करे, वो काम कम करने वाले सबस्केल पर मार्केटिंग खर्च कर रहा है।

7.2 App-alliance is not coach-alliance

यहीं भारतीय प्रॉडक्ट डिज़ाइन या तो बड़ा होगा या फिसलेगा।

जब नेविगेटर हल्की चेक-इन भेजे, ऐप वाला अलायंस लक्षण बताता रहता है।

जब कोई क्वालिफ़ाइड प्रोफ़ेशनल भी लूप में हो, ऐप वाला अलायंस इस्तेमाल बताता है। लक्षण बताना बंद कर देता है।

एक पढ़त। इंसान बदलाव का तरीका सोख लेता है। ऐप कॉपी बन जाता है।

दूसरी पढ़त। इंसान आते ही ऐप को रेट करना सहायक को रेट करना है।

दोनों सूरत में “एआई कोच plus इंसान हैंड-ऑफ़” वाले प्रोग्राम को दो स्कोर नापने होंगे।

एक एजेंट का। एक हैंड-ऑफ़ लेने वाले इंसान का। दोनों मिलाकर डैशबोर्ड सजावट बनेगा।

छोटी शराब-प्रोग्राम स्टडी वाक्य का दूसरा आधा है। लोग दो अलायंस एक साथ रख सकते हैं। दोनों करीब रेट कर सकते हैं। दोहरा नाप संभव है। ये रोमांस नहीं। ये औज़ार है।

7.3 Limitations

इस रिव्यू का रजिस्टर्ड प्रोटोकॉल नहीं। दो-स्क्रीनर कापा नहीं।

सदस्यता डेटाबेस उपलब्ध इंटरफ़ेस से चले। लाइब्रेरियन वाला पूरा निर्यात नहीं। फ़्लो की गिनती रेंज है। ये ईमानदारी है। GRADE पर कट भी है।

सबसे बड़े औसत उद्योग से जुड़े पूरे करने वालों से आए। रात एक पर छूटने वाला रात पांच नहीं रेट करता। हर जल्दी औसत “अभी वहां होने” पर टिका है।

नाप इंसानों के आमने-सामने मदद वाले रिश्तों की रिसर्च से उधार हैं। संज्ञा बदल दी गई। 2022 की नैरेटिव रिव्यू पहले चेता चुकी। डिजिटल अलायंस को शायद अपना ढांचा चाहिए। छह-आइटम इन्वेंटरी ने तीन फ़ैक्टर नहीं, एक फ़ैक्टर पाया। अगर तीन टांगों वाला स्टूल सॉफ्टवेयर में एक हो जाए, तो बॉन्ड-टास्क-गोल हमारी थोपी कहानी है।

लगभग कोई पेपर भारतीय नौकरीपेशा सैंपल नहीं। फ़्री-टेक्स्ट पेपर पर एक भारतीय सह-लेखक। ग्रामीण मध्य प्रदेश का हिंदी होमवर्क-सहायक चैटबॉट अलायंस नापता ही नहीं। CTRI चुप है।

अलायंस-नतीजा गुणांक दर्जन तीसरे वेरिफ़ेबल से आज़ाद नहीं। जो ढांचा पसंद करते हैं, वे प्रोग्राम भी पसंद करते हैं और रेप करके सुधरते हैं। अलायंस कारण नहीं, निशान हो सकता है। दोतरफ़ा रास्ता वाली फ़ॉलो-अप चेतावनी है।

दो जेनेरेटिव ट्रायल में वेट-लिस्ट कंट्रोल लक्षण असर बड़ा दिखाता है। यहां वे अलायंस स्कोर के लिए हैं। ग्रुप कोचिंग से जीत का सबूत नहीं।

7.4 What this means for an Indian workplace programme

इसका मतलब

भारतीय ऑफ़िस प्रोग्राम के लिए इसका मतलब

भारतीय प्रोफ़ेशनल को यार कहने वाले बॉट की ज़रूरत नहीं। उसे हमदर्दी का ठप्पा नहीं चाहिए। पुणे की एक एनालिस्ट को सोचिए। हफ़्ते की रात है, लैपटॉप अब भी खुला है।

उसे ऐसा साथ चाहिए जो मकसद माने — वीकनाइट बारह से पहले सोना।

जो मंगलवार का रेप नाम से बताए — फ़ोन बेडरूम के बाहर चार्ज पर।

जो रात 11.47 पर हिंदी टाइप करने पर न हिचके।

पहले हफ़्ते के शुक्रवार तक वो एजेंट को अच्छी रेटिंग दे चुकी है।

बदलाव बताने वाला स्कोर कुछ दर्जन रेप्स के बाद ही दिखता है।

पहला हफ़्ता स्कोर कमाने के लिए बनाएं। दूसरा से चौथा हफ़्ता भविष्यवाणी कमाने के लिए।

जब एजेंट उसे इंसान कोच को सौंपे, तो उस रिश्ते का स्कोर अलग नापें।

सॉफ्टवेयर के 0.18 पर बनी सेल्स स्लाइड पर इंसान का 0.28 न लगाएं।

जो स्टडी गायब है वो चलाएं। नौकरीपेशा वयस्क। 23 भाषाएं। व्हाट्सएप-फ़र्स्ट।

DPDP-निवासी डेटा। दिन 5 और हफ़्ता 4 पर अलायंस। हफ़्ता 8 पर नतीजे।

भारत-निवासी जांचकर्ता जिनके पास कंपनी इक्विटी न हो।

जब तक वो स्टडी नहीं, आज के नंबर बनाने की इजाज़त हैं। डींग की नहीं।

7.5 Research gaps, naming the India gap explicitly

भारत वाला गैप मूड नहीं है। खाली खाना है।

CTRI पर भारतीय नौकरीपेशा वयस्कों में डिजिटल एजेंट वाले वर्किंग अलायंस की कोई RCT नहीं।

सबसे पास तीन चीज़ें हैं। ग्रामीण मध्य प्रदेश का होमवर्क-सहायक चैटबॉट, बिना अलायंस नाप के। फ़्री-टेक्स्ट पेपर, बेंगलुरु सह-लेखक, भारत-अलग स्कोर नहीं। बेंगलुरु और भोपाल में रिसर्च ऐप के संज्ञान वाले इस्तेमाल। ऑफ़िस वाला सवाल इनमें से कोई नहीं पूछता।

पीछे और गैप हैं। हिंदी, तमिल, तेलुगु, बांग्ला, मराठी में नाप खुद जांचा नहीं। हैंड-ऑफ़ ऐप-अलायंस बढ़ाता है या घटाता है, पता नहीं। जिनकी पहली आदत व्हाट्सऐप है, उनमें दिन-5 बॉन्ड हफ़्ता-8 रिटेंशन बताता है या नहीं, पता नहीं। कितने मिनट स्कोर को भविष्यवक्ता बनाते हैं, पता नहीं।

आखिरी तस्वीर, फिर बस। मुंबई की हर बारिश में इमारत पर बहस ये नहीं होती कि उसमें लोग रहें या नहीं। बहस ये होती है कि इमारत किस चीज़ पर खड़ी है।

सॉफ्टवेयर वाला अलायंस वो इमारत है जिसमें लोग पहले से रह रहे हैं।

पक्की नींव टास्क-मकसद की सहमति है। दो बार नापी। मौसम बिगड़े तो इंसान को सौंपी।

कच्ची भराई पहले दिन की मीठी पंक्ति है। और वो प्रेस रिलीज़ जो इंसान का र उद्धृत करे।

08

FAQ

Q1 क्या ऐप से वर्किंग अलायंस बन सकता है, या ये सिर्फ़ ब्रांडिंग है?

स्कोर बन सकता है। बड़ी कोहॉर्ट पांच दिन में 3.36 से 3.75 पर बैठती हैं। इंसान आउटपेशेंट औसत करीब 3.8। ये फ़ॉर्म का नतीजा है। शादी नहीं। अगला रेप करने की इजाज़त है। सॉफ्टवेयर के परवाह करने का सबूत नहीं।

Q2 ऑफ़िस प्रोग्राम इसे कब नापे?

ऑनबोर्डिंग जाननी हो तो पांचवें दिन नापें। स्ट्रेस गिरावट बताने वाला नंबर चाहिए तो तीसरे या चौथे हफ़्ते। एक ट्रायल में पहले हफ़्ते के स्कोर ऊंचे थे और बताते नहीं थे।

Q3 क्या मज़बूत अलायंस का मतलब लोग सच में सुधरते हैं?

कभी-कभी। थोड़ा। सबसे साफ़ सिर्फ़-सॉफ्टवेयर आंकड़े पार्शियल $r = -0.18$ और तीसरे-चौथे हफ़्ते के $\beta = -0.13$ से -0.17 । इंसान कोचिंग $r = 0.28$ के पास। इंसान का नंबर सॉफ्टवेयर डेटा पर न बेचें।

Q4 **इंसान कोच जोड़ दें तो ऐप नापना बंद कर दें?**

नहीं। ऐप-अलायंस बताता रहता है लोग टैप करते रहेंगे या नहीं। एक दो-स्टडी पेपर में लक्षण तब बताए जब इंसान हल्का नेविगेटर था। जब कोई क्वालिफ़ाइड प्रोफ़ेशनल भी लूप में था, तब नहीं। दोनों नापें। हैंड-ऑफ़ फीचर है। मिक्स बटन नहीं।

Q5 **क्या ये स्टडी भारतीय नौकरीपेशा लोगों पर चली है?**

CTRI पर नहीं। अलायंस वाली RCT के रूप में नहीं। शून्य। ग्रामीण मध्य प्रदेश का हिंदी होमवर्क चैटबॉट अलायंस नापता नहीं। खाली खाना अगला पेपर है। हां वाली अफ़वाह नहीं।

09

References

Primary sources only. Each entry carries a DOI or a registry URL. नाम, साल, DOI ज्यों के त्यों।

1. Darcy A, Daniels J, Salinger D, Wicks P, Robinson A. Evidence of human-level bonds established with a digital conversational agent: cross-sectional, retrospective observational study. JMIR Formative Research. 2021;5(5):e27868. <https://doi.org/10.2196/27868>
2. Beatty C, Malik T, Meheli S, Sinha C. Evaluating the therapeutic alliance with a free-text CBT conversational agent: a mixed-methods study. Frontiers in Digital Health. 2022;4:847991. <https://doi.org/10.3389/fdgth.2022.847991>
3. Goldberg SB, Baldwin SA, Riordan KM, et al. Alliance with an unguided smartphone app: validation of the Digital Working Alliance Inventory. Assessment. 2022;29(6):1331-1345. <https://doi.org/10.1177/10731911211015310>
4. Heinz MV, et al. Randomized trial of a generative AI chatbot for mental health treatment. NEJM AI. 2025;2(4). <https://doi.org/10.1056/AIoa2400802>
5. He L, Basar E, Krahmer E, Wiers R, Antheunis M. Effectiveness and user experience of a smoking cessation chatbot: mixed methods study comparing motivational interviewing and confrontational counseling. Journal of Medical Internet Research. 2024;26:e53134. <https://doi.org/10.2196/53134>
6. Benitez B, Kiluk BD, et al. The connection still matters: therapeutic alliance with digital treatment for alcohol use disorder. Alcohol: Clinical and Experimental Research. 2023;47(11):2197-2207. <https://doi.org/10.1111/acer.15199>
7. Shoshani A, Gurfinkel B, Kor A, et al. Efficacy of a conversational AI agent for psychiatric symptoms and digital therapeutic alliance: a randomised clinical trial. JAMA Network Open. 2026;9(4):e266713. <https://doi.org/10.1001/jamanetworkopen.2026.6713>
8. Macrynika N, Chang S, Torous J. Is digital alliance associated with engagement and outcomes in guided digital interventions? An analysis of data from two studies. Journal of Affective Disorders. 2025;383:335-340. <https://doi.org/10.1016/j.jad.2025.04.127>

9. Herrero R, Vara MD, Miragall M, et al. Working Alliance Inventory for Online Interventions-Short Form (WAI-TECH-SF): the role of the therapeutic alliance between patient and online program in therapeutic outcomes. *International Journal of Environmental Research and Public Health*. 2020;17(17):6169. <https://doi.org/10.3390/ijerph17176169>
10. Berry K, Salter A, Morris R, James S, Bucci S. Assessing therapeutic alliance in the context of mHealth interventions for mental health problems: development of the Mobile Agnew Relationship Measure (mARM) questionnaire. *Journal of Medical Internet Research*. 2018;20(4):e90. <https://doi.org/10.2196/jmir.8252>
11. Henson P, Wisniewski H, Hollis C, Keshavan M, Torous J. Digital mental health apps and the therapeutic alliance: initial review. *BJPsych Open*. 2019;5(1):e15. <https://doi.org/10.1192/bjo.2018.86>
12. Tong F, Lederman R, D'Alfonso S, Berry K, Bucci S. Digital therapeutic alliance with fully automated mental health smartphone apps: a narrative review. *Frontiers in Psychiatry*. 2022;13:819623. <https://doi.org/10.3389/fpsyt.2022.819623>
13. Malouin-Lachance A, Capolupo J, Laplante C, Hudon A. Does the digital therapeutic alliance exist? Integrative review. *JMIR Mental Health*. 2025;12:e69294. <https://doi.org/10.2196/69294>
14. Flückiger C, Del Re AC, Wampold BE, Horvath AO. The alliance in adult psychotherapy: a meta-analytic synthesis. *Psychotherapy*. 2018;55(4):316-340. <https://doi.org/10.1037/pst0000172>
15. Flückiger C, et al. Alliance in internet-based interventions: a systematic review and correlation multilevel meta-analysis. *Journal of Consulting and Clinical Psychology*. 2026. <https://doi.org/10.1037/ccp0001005>
16. Probst GH, Berger T, Flückiger C. The alliance-outcome relation in internet-based interventions for psychological disorders: a correlational meta-analysis. *Verhaltenstherapie*. 2019;32:135-146. <https://doi.org/10.1159/000503432>
17. Munder T, Wilmsers F, Leonhart R, Linster HW, Barth J. Working Alliance Inventory-Short Revised (WAI-SR): psychometric properties in outpatients and inpatients. *Clinical Psychology & Psychotherapy*. 2010;17(3):231-239. <https://doi.org/10.1002/cpp.657>
18. Henson P, Peck P, Torous J. Considering the therapeutic alliance in digital mental health interventions. *Harvard Review of Psychiatry*. 2019;27(4):268-273. <https://doi.org/10.1097/HRP.0000000000000224>
19. Agrawal R, Monteiro K, Narasimhamurti NS, et al. A conversational agent (PracticePal) to support the delivery of a brief behavioural activation treatment for depression in rural India: development and pilot-testing study. *JMIR Formative Research*. 2025. <https://doi.org/10.2196/73563>
20. Bordin ES. The generalizability of the psychoanalytic concept of the working alliance. *Psychotherapy: Theory, Research & Practice*. 1979;16(3):252-260. <https://doi.org/10.1037/h0085885>
21. Goldberg SB, et al. Bidirectional alliance and distress in an unguided app (same-sample follow-up). *Clinical Psychological Science*. 2024. <https://doi.org/10.1177/21677026231184890>
22. Association between working alliance and treatment outcomes in a mobile health intervention with a conversational agent (CanRelax). *Internet Interventions*. 2026. <https://doi.org/10.1016/j.invent.2026.100929>
23. He L, Basar E, Wiers RW, Antheunis ML, Krahmer E. Chatting your way to quitting: a longitudinal exploration of smokers' interaction with a cessation chatbot. *Internet Interventions*. 2025. <https://doi.org/10.1016/j.invent.2025.100807>
24. Li H, Zhang R, Lee YC, Kraut RE, Mohr DC. Systematic review and meta-analysis of AI-based conversational agents for promoting mental health and well-being. *npj Digital Medicine*. 2023;6:236. <https://doi.org/10.1038/s41746-023-00979-5>

लेखक के बारे में

लेखक के बारे में

डॉ. अतुल असवानी; मुंबई के साइकियाट्रिस्ट और ट्रेड रिज़ल्ट्स कोच

हवाला कैसे दें

Aswani A. 5 दिन में स्कोर 3.36 से 3.75 — बॉट से भी बनती है बात. EmotionApp Research Paper 38 (Hindi edition). Mumbai: Tranquilo Meditek Sytems Private Limited; 2026.
<https://emotionapp.in/research/working-alliance-digital-agents>

अपडेट

22 September 2026

नोट

इमोशनऐप एक नॉन-क्लिनिकल इमोशनल फ़िटनेस कोचिंग प्लेटफ़ॉर्म है। यह पेपर प्रैक्टिस, तरीकों और कोचिंग के बारे में है। यह थेरेपी, इलाज, डायग्नोसिस या मेडिकल सलाह नहीं है।

© 2026 Tranquilo Meditek Sytems Private Limited, Mumbai
· emotionapp.in

सुरक्षा नोट

अगर आप परेशानी में हैं, तो किसी योग्य प्रोफ़ेशनल से या Tele-MANAS (14416) पर बात करें।