

मातृभाषा में कोचिंग: असर 0.28 बड़ा — भारत में ट्रायल शून्य

Mother Tongue, Bigger Effect? A Meta-Analysis of Language-of-Delivery in
Psychological Interventions

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



मुख्य आंकड़ा • g

0.28

बीच-स्टडी का सबसे अच्छा संकेत $g = 0.28$ है। मातृभाषा $d = 0.49$ । भाषा-मैच न बताने पर $d = 0.21$ ।

मातृभाषा में कोचिंग: असर 0.28 बड़ा — भारत में ट्रायल शून्य

Mother Tongue, Bigger Effect? A Meta-Analysis of Language-of-Delivery in Psychological Interventions

क्या मातृभाषा में दिया गया प्रोग्राम दूसरी भाषा से बेहतर चलता है?

PRISMA-2020 सिस्टमैटिक रिव्यू • नैरेटिव सिंथेसिस

इस पेपर में

01 सार

02 मुख्य नतीजे

03 परिभाषा

04 भूमिका

05 तरीका

06 नतीजे

07 चर्चा

08 FAQ

09 संदर्भ

इंडेक्सिंग टाइटल

Language of delivery and outcomes in psychological interventions: a systematic review and meta-analysis

लेखक

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

प्रकाशित

22 September 2026 • हिंदी संस्करण • emotionapp.in/research/language-of-delivery-meta-analysis

01

सार

क्या मातृभाषा में दिया गया वेलबीइंग प्रोग्राम बेहतर चलता है? बीच-स्टडी तुलना कहती है हाँ। फ़ायदा करीब $g = 0.28$ है।

Hedges g असर कितना बढ़ा है, इसका नाप है। पाँच से कम ट्रायल एक ही प्रोग्राम को पहली और दूसरी भाषा में जोड़कर देखे। इसलिए नया पूल नहीं बन सकता। भारत के हेड-टू-हेड ट्रायल की संख्या शून्य है।

यह PRISMA-2020 सिस्टमैटिक रिव्यू है। सवाल यह था। क्या बहुभाषी वयस्कों को मातृभाषा में दिया गया प्रोग्राम दूसरी भाषा से बेहतर है।

सर्च PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI, ClinicalTrials.gov और OSF पर चला। साल 2000 से 2026 तक। RCT ऐसी स्टडी है जिसमें लोगों को लॉटरी की तरह दो ग्रुप में बाँटा जाता है। ऐसा एक भी RCT नहीं मिला जिसने बहुभाषी वयस्कों में एक ही प्रोग्राम को L1 और L2 में बाँटा हो।

सबसे अच्छे नंबर पुरानी मेटा-एनालिसिस के अंदर की भाषा-मैचिंग से आते हैं। Griner और Smith (2006) ने 76 स्टडी और 25,225 लोगों को जोड़ा। मातृभाषा वाले प्रोग्राम अंग्रेज़ी से दोगुना असरदार दिखे। $d = 0.49$ बनाम $d = 0.21$ ।

Soto और साथियों (2018) ने 99 स्टडी अपडेट की। कुल $d = 0.50$ । पब्लिकेशन बायस काटने के बाद 0.35। पसंदीदा भाषा और अनुवादित टेस्ट दोनों असर बढ़ाते दिखे।

Hall और साथियों (2016) ने देखा। संस्कृति के हिसाब से बदला प्रोग्राम उसी के बिना-बदले रूप से बेहतर था। $g = 0.52$ । भाषा खुद स्थिर मॉडरेटर नहीं बनी।

दुभाषिया वाली तुलना मिली-जुली है। भारत के अच्छे प्रोग्राम स्थानीय भाषा में चले। Healthy Activity Programme, गोवा, उनमें है। किसी ने भाषा को लॉटरी से नहीं बाँटा।

यह कहना कमज़ोर है कि मातृभाषा ही नतीजा बेहतर बनाती है। भारतीय ऑफ़िस के लिए बात सरल है। अंग्रेज़ी मातृभाषा सिर्फ़ 0.02 फ़ीसदी भारतीयों की है। ऑफ़िस लगभग सबकी अंग्रेज़ी में है। जब साफ़ प्रयोग नहीं है, भाषा मिलाना कम-पछतावे वाला फ़ैसला है।

02

मुख्य नतीजे

मुख्य नतीजे

0.28

बीच-स्टडी का सबसे अच्छा संकेत $g = 0.28$ है। मातृभाषा $d = 0.49$ । भाषा-मैच न बताने पर $d = 0.21$ ।

0

भारत में शून्य ट्रायल एक ही वेलबीइंग प्रोग्राम को पहली बनाम दूसरी भाषा में जोड़ा।

0.66

जब नतीजे के टेस्ट खुद अनुवादित थे, असर दोगुने से ज़्यादा दिखा। $d = 0.66$ बनाम $d = 0.28$ ।

4

दुभाषिया वाली चार वयस्क तुलनाएँ हैं। वे एक राय नहीं रखतीं। एक बड़ी कोहोर्ट $n = 825$ में दुभाषिया के साथ फ़ायदा छोटा था। तीन छोटी स्टडी में नतीजे लगभग बराबर थे।

0.02

अंग्रेज़ी मातृभाषा 0.02 फ़ीसदी भारतीयों की है। इमोशनऐप 23 भाषाएँ देता है। भाषा को अलग करके मापने वाला ट्रायल अभी नहीं चला।

03

परिभाषा

परिभाषा

डिलीवरी की भाषा वह भाषा है जिसमें प्रोग्राम व्यक्ति से बात करता है। पहली भाषा, या L1, वह भाषा है जो पहले सीखी। घर और दिल उसी में बोलते हैं। दूसरी भाषा, या L2, बाद में आई। ऑफिस, परीक्षा, फॉर्म उसी में चलते हैं। मातृभाषा L1 का आम भारतीय नाम है।

भाषा मिलाना मतलब कोच, स्क्रिप्ट और टेस्ट व्यक्ति की पसंदीदा भाषा में हों। दुभाषिया वाली डिलीवरी में तीसरा व्यक्ति अर्थ ढोता है। मशीन-अनुवाद वाली डिलीवरी में मॉडल अंग्रेज़ी स्क्रिप्ट दूसरी भाषा में पलटता है। कमरे में द्विभाषी कोच नहीं होता।

यह रिव्यू भाषा को डिज़ाइन का फैसला मानता है। यह कोई बीमारी का नाम नहीं है, न संस्कृति की सजावट। L2 में लिखी चिंता अक्सर ठंडी और साफ़-सुथरी निकलती है, चुभन पीछे छूट जाती है। कड़वे विचार से दूरी बनाना तभी काम आता है जब शब्दों में वज़न बचा हो। प्रोडक्ट की भाषा में, यह अंदरूनी ज़िंदगी का स्क्रीन है।

04

भूमिका

पुणे की एक प्रोफेशनल अंग्रेज़ी में स्टैंडअप खत्म करती है। व्हाट्सऐप खोलती है। बहन को मराठी में लिखती है। थकान ऐसी है जिसका ऑफिस में नाम नहीं। उसका मैनेजर घर लौटते हुए यही करता है, गुजराती में। यह आदत नहीं है। यह भारतीय प्रोफेशनल ज़िंदगी का डिफॉल्ट है।

भारत की जनगणना 2011 ने 121 मातृभाषाएँ गिनीं। अनुसूचित भाषाएँ 22 हैं। हिंदी मातृभाषा 43.6 फ़ीसदी की है। अंग्रेज़ी मातृभाषा 0.02 फ़ीसदी की है। करीब 2.6 लाख लोग। लगभग दस में से एक अंग्रेज़ी दूसरी या तीसरी भाषा बताता है। बाकी घर पर क्षेत्रीय भाषा में ही सच बोलते हैं।

यह पेपर एक सीधा सवाल पूछता है। क्या मातृभाषा में दिया गया वेलबीइंग प्रोग्राम दूसरी भाषा से बेहतर है?

सवाल ज़रूरी है क्योंकि भारतीय ऑफ़िस ने बिना प्रयोग के जवाब चुन लिया। ईएपी लाइन अंग्रेज़ी में खुलती है। कोचिंग डेक अंग्रेज़ी में आता है। साँस वाले ऐप पर अंग्रेज़ी लेबल है। हिंदी टॉगल पर्यटक मेन्यू जैसा लगता है। जिस भाव को रेप चाहिए — शर्म, घबराहट, वह कसैलापन जो स्ट्रेस भी नहीं, फ़र्ज़ भी नहीं — वह दूसरी शब्द-सूची में रहता है।

स्टैंड-अप कॉमेडियन यह जानते हैं। मज़ाक अनुवाद में मर जाता है। इसलिए नहीं कि श्रोता मंद है। इसलिए कि टाइमिंग मूल मुँह में रहती है। उद्यमी यही बात दूसरी सूरत में जानते हैं। प्रोडक्ट-मार्केट फ़िट फ़ीचर लिस्ट नहीं है। ग्राहक पहले स्क्रीन में खुद को पहचानता है या नहीं। एआई मॉडल तीसरी सूरत जानता है। टोकनाइज़र अंग्रेज़ी पर ट्रेनिंग ले। फिर तमिल वाक्य पकड़े। मॉडल आधी ताकत लिपि जोड़ने में खर्च करता है।

कोई भी भारतीय ऑफ़िस चौथी सूरत जानता है। हैदराबाद का एक सेल्स मैनेजर पिता से मुश्किल कॉल पर है। अंग्रेज़ी में शुरू करता है। अटकता है। जो एक वाक्य ज़रूरी है, उसके लिए तेलुगु में चला जाता है। पिता अंग्रेज़ी समझते थे। जवाब तेलुगु को देते हैं। तथ्य नहीं बदले। शब्द बदले। भाषा अर्थ का रैपर नहीं है। अर्थ भाषा में ही रखा रहता है।

पुराना सबूत भाषा को संस्कृति-अनुकूलन के बड़े कटोरे का एक मसाला मानता था। Griner और Smith ने 2006 में 76 स्टडी जोड़ीं। अनुकूलित प्रोग्राम का मध्यम फ़ायदा $d = 0.45$ था। उसी पेपर में एक वाक्य था। APA Monitor ने उसे पोस्टर बना दिया। मातृभाषा वाले प्रोग्राम अंग्रेज़ी वाले से दोगुना असरदार दिखे। दस साल बाद Hall और साथियों ने देखा। अनुकूलित प्रोग्राम उसी के बिना-बदले रूप से बेहतर था। $g = 0.52$ । भाषा अलग से स्थिर मॉडरेटर नहीं बनी। भाषा के सवाल पर कोई अलग मेटा-एनालिसिस नहीं थी।

यह रिव्यू उसी छेद पर बैठता है। साफ़ कहें। यह बीच-स्टडी तुलना है। लगभग किसी ने एक ही प्रोग्राम को L1 और L2 बाँहों में नहीं बाँटा। फ़ील्ड ने उन प्रोग्रामों की तुलना की जो मातृभाषा में चले। उनकी तुलना उनसे की जो अंग्रेज़ी में चले। साथ में कहानी, रूपक, कोच का समुदाय और टेस्ट भी बदले। भाषा संस्कृति के साथ चलती है। दोनों को अलग करना यह पेपर पूरा नहीं कर सकता। भारत ने शुरू भी नहीं किया।

इमोशनऐप के लिए दांव व्यावहारिक है। प्लेटफ़ॉर्म भारतीय प्रोफ़ेशनल्स के लिए नॉन-क्लिनिकल इमोशनल फ़िटनेस (Emotional Fitness) कोचिंग है। पॉज़िटिव साइकोलॉजी के तरीके गिने-चुने रोज़ के रेप्स की तरह मिलते हैं। एआई कोच है। इंसान का हाथ भी है। पहले व्हाट्सऐप। 23 भारतीय भाषाएँ। डेटा भारत में रहता है। DPDP के हिसाब से। ISO 9001 और ISO 13485 के नीचे बना। अगर मातृभाषा सच में फ़ायदा है, 23 भाषाएँ विज्ञापन नहीं हैं। वे प्रोडक्ट हैं। अगर मातृभाषा सिर्फ़ संस्कृति-अनुकूलन की सवारी है, ईमानदार काम यह है कि कह दिया जाए। इंजीनियरिंग कहीं और लगे। यह रिव्यू फ़र्क़ बताने की कोशिश है।

05

तरीका

प्रोटोकॉल

यह रिव्यू PRISMA 2020 के हिसाब से लिखा गया। पहले से रजिस्टर नहीं था। विश्लेषण पहले से तय था। रैंडम-इफ़ेक्ट्स REML। Hedges g । 95% CI। I^2 । τ^2 । 95 फ़ीसदी प्रेडिक्शन इंटरवल। Egger टेस्ट। trim-and-fill। leave-one-out। RoB 2। GRADE। 95% CI यानी नंबर के इर्द-गिर्द वह दायरा जिसमें असली असर के रहने की उम्मीद है। I^2 बताता है स्टडी आपस में कितनी जुदा हैं।

वही प्रोटोकॉल एक बैकअप भी तय करता था। अगर पाँच से कम पात्र ट्रायल मिलें, तो नैरेटिव सिंथेसिस चले। कह दिया जाए कि मेटा-एनालिसिस अभी संभव नहीं। पूल न हो। वही बैकअप चला।

पात्रता

आबादी: बहुभाषी वयस्क। उम्र 18 साल या ऊपर। प्रोग्राम: कोई भी मनोवैज्ञानिक या वेलबीइंग प्रोग्राम जिसने डिलीवरी की भाषा बताई। तुलना: वही प्रोग्राम दूसरी भाषा में। या L1 बनाम L2 की बीच-स्टडी तुलना। मुख्य नतीजे: कोई भी वेलबीइंग या लक्षण वाला नतीजा। साल: 2000 से 2026। डिज़ाइन: पूल के लिए RCT या कंट्रोलड ट्रायल ज़रूरी थे। दुभाषिया वाली कहानी के लिए अवलोकन स्टडी रखी गई।

बच्चों के भाषा प्रोग्राम बाहर रहे। अफ़ेजिया बाहर रहा। क्लास की भाषा-ऑफ़-इंस्ट्रक्शन बाहर रही। जो स्टडी सिर्फ़ एक ही भाषा के प्रैक्टिशनर से मिलाती थी, वह भी भाषा-कंट्रास्ट से बाहर रही। Griner और Smith का पुराना नियम यही था।

समीक्षा के प्रोग्रामों के लिए इलाज, बीमारी या डायग्नोसिस शब्द नहीं लगाए गए। स्रोत पेपर के टाइटल में वे शब्द हों तो टाइटल ज्यों का त्यों है।

जानकारी के स्रोत और सर्च

स्रोत: PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI (भारत), ClinicalTrials.gov, और OSF प्रीप्रिंट। सर्च 1 January 2000 से 22 September 2026 तक चली। Griner और Smith (2006), Hall और साथी (2016), Soto और साथी (2018), Chowdhary और साथी (2014), Benish और साथी (2011) और Arundell और साथी (2021) की संदर्भ सूची हाथ से देखी गई। भारतीय रजिस्ट्री अंग्रेज़ी और भाषा-नामों से खोजी गई। हिंदी, तमिल, तेलुगु, मराठी, बंगाली, गुजराती, कन्नड़, मलयालम, पंजाबी, उर्दू, कोंकणी।

PubMed	("language"[Title/Abstract] OR "mother tongue"[Title/Abstract] OR bilingual[Title/Abstract] OR "first language"[Title/Abstract] OR "native language"[Title/Abstract] OR "language matching"[Title/Abstract] OR "preferred language"[Title/Abstract]) AND (intervention[Title/Abstract] OR coaching[Title/Abstract] OR "wellbeing"[Title/Abstract] OR "well-being"[Title/Abstract] OR counsel*[Title/Abstract]) AND (outcome[Title/Abstract] OR efficacy[Title/Abstract] OR effectiveness[Title/Abstract]) AND ("2000/01/01"[Date - Publication] : "2026/12/31"[Date - Publication]) बिना उम्र और डिज़ाइन की सीमा के चौड़ी स्ट्रिंग ने 2000–2026 के लिए 20,267 PubMed रिकॉर्ड दिए। ज़्यादातर बच्चों के भाषा प्रोग्राम थे। बाकी अफ़ेजिया या शिक्षा ट्रायल। इसलिए शीर्षक और सार नियमों से पढ़े गए।
PsycINFO (Ovid)	वही बूलियन। सीमा: 2000–2026, इंसान, वयस्क 18+।
Scopus	TITLE-ABS-KEY((language OR "mother tongue" OR bilingual OR "first language") AND (intervention OR coaching OR wellbeing OR counsel*) AND (outcome OR efficacy OR effectiveness)) AND PUBYEAR > 1999 AND PUBYEAR < 2027
Web of Science	TS=(language OR "mother tongue" OR bilingual OR "first language") AND TS=(intervention OR coaching OR wellbeing) AND TS=(outcome OR efficacy) Timespan=2000-2026
Google Scholar	"mother tongue" OR "first language" OR "native language" intervention (wellbeing OR depression OR anxiety) (RCT OR randomised OR randomized) after:1999. पहले 200 रिकॉर्ड देखे।
CTRI	counselling OR coaching OR wellbeing AND language OR Hindi OR Tamil OR bilingual
ClinicalTrials.gov	(language matching OR "mother tongue" OR bilingual) Interventional Adult 01/01/2000–22/09/2026
OSF preprints	language matching psychological intervention
चुनाव और निकाला गया डेटा	दो समीक्षक शीर्षक और सार नियमों से देखे। पूरा पाठ तब पढ़ा जब भाषा लिखी थी या अनुमान लग सकता था। स्थानीय भाषा का प्रैक्टिशनर। अनुवादित मैनुअल। दुभाषिया। द्विभाषी बाँह। मतभेद बात से सुलझे। निकाले गए फ़ील्ड: n, डिज़ाइन, देश, डिलीवरी की भाषा, डोज़ सेशन और मिनट में, चैनल, गाइडेंस, नतीजे का नाप, समय-बिंदु, और L1 बनाम L2 कंट्रास्ट भीतर-स्टडी था या बीच-स्टडी। जो DOI या स्थिर URL से न पुष्ट हो, वह UNVERIFIED रहा। पूल में नहीं

गया। एक भारतीय “क्षेत्रीय भाषा हेल्पलाइन RCT” बहु-विषय जर्नल में था। ट्रायल रजिस्ट्रेशन नहीं। एथिक्स रिकॉर्ड नहीं। UNVERIFIED। हटाया।

विश्लेषण योजना, जैसी तय थी और जैसी चली

तय मॉडल: रैंडम-इफेक्ट्स REML। Hedges g और 95% CI। I^2 , τ^2 , 95 फ़ीसदी प्रेडिक्शन इंटरवल। Egger और trim-and-fill। leave-one-out। RoB 2। GRADE। तय मॉडरेटर: पहली बनाम दूसरी भाषा। दुभाषिया वाली। मशीन-अनुवाद वाली।

चला विश्लेषण: नैरेटिव सिंथेसिस। प्राथमिक L1-बनाम-L2 ट्रायल से नई पूल की g नहीं बनी। पुरानी मेटा-एनालिसिस की भाषा-मैचिंग सबसे अच्छा मात्रा संकेत है। वैसा ही लिखा है। दुभाषिया स्टडी अलग कही गई। भाषा-मैचिंग से नहीं जोड़ी गई। मशीन-अनुवाद का कोई पात्र कंट्रोल ट्रायल नहीं मिला।

पक्षपात का खतरा और पक्कापन

RoB 2 उस RCT पर लागू होता जिसका पूल में जाना संभव था। भाषा-कंट्रास्ट के लिए लगभग हर उम्मीदवार या तो बिना लॉटरी का था, या लॉटरी किसी और बात पर थी। GRADE दो दावों पर लगा। एक, मातृभाषा दूसरी भाषा से बेहतर है। दो, संस्कृति-अनुकूल प्रोग्राम बिना-अनुकूल से बेहतर हैं।

PRISMA 2020 फ़्लो (इस सर्च से फिर बनाया)

यह फ़्लो इस रिव्यू की सर्च से फिर बनाया गया। रजिस्टर्ड सॉफ़्टवेयर पाइपलाइन का आउटपुट नहीं है।

- रिकॉर्ड मिले: करीब 4,530। PubMed ~1,100। PsycINFO ~900। Scopus ~1,400। Web of Science ~800। Google Scholar पहले 200। CTRI ~40। ClinicalTrials.gov ~60। OSF ~30।
- डुप्लीकेट हटाने के बाद: करीब 2,800।
- शीर्षक-सार से बाहर: करीब 2,650। मुख्य वजह: बच्चों के भाषा प्रोग्राम। क्लास की भाषा। अफ़ेजिया। वेलबीइंग नतीजा नहीं। वयस्क नहीं।
- पूरा पाठ देखा: करीब 150।
- पूरे पाठ पर बाहर: भाषा अलग न निकली। बच्चा सैंपल। नतीजा नहीं। वेलबीइंग प्रोग्राम नहीं। 2000 से पहले। कंट्रोल डिज़ाइन नहीं। UNVERIFIED स्रोत।
- नैरेटिव में शामिल: 6 पुरानी मेटा-एनालिसिस। 5 दुभाषिया या भाषा-फ़ॉर्मेट स्टडी। 4 भारतीय या दक्षिण-एशियाई प्रोग्राम स्थानीय भाषा में, भाषा पर लॉटरी नहीं।
- एक ही प्रोग्राम की नई पूल L1-बनाम-L2 मेटा-एनालिसिस के योग्य: 0।

06

नतीजे

स्टडी की तस्वीर

कोई पात्र RCT बहुभाषी वयस्कों को एक ही वेलबीइंग प्रोग्राम की पहली भाषा बनाम दूसरी भाषा में नहीं बाँटता।

फ़्रील्ड में तीन परतें हैं।

परत 1 – पुरानी मेटा-एनालिसिस जिनमें भाषा मॉडरेटर है। यह बीच-स्टडी तुलना है। स्पेनिश में चली स्टडी की तुलना अंग्रेज़ी वाली दूसरी स्टडी से होती है। भाषा समुदाय, आदत, टेस्ट के अनुवाद और कंटेंट के साथ चलती है।

परत 2 – दुभाषिया बनाम एक ही भाषा का कोच। यह भाषा-कंट्रास्ट के करीब है। पर तीन लोगों की बात की तुलना दो लोगों की बात से है। ज़्यादातर लॉटरी नहीं।

परत 3 – भारत और बहुभाषी जगहों पर स्थानीय भाषा में चले प्रोग्राम, L2 बाँह नहीं। ये बताते हैं मातृभाषा चल सकती है। आम देखभाल से बेहतर भी दिख सकती है। भाषा की परीक्षा नहीं करते।

Table 1. Prior meta-analyses that speak to language of delivery

Source	k (N)	Language contrast	Main finding	Language-specific number
Griner & Smith 2006	76 (25,225)	Between-study moderator	Adapted programmes d = 0.45 (0.36–0.53)	Native-language (if not English) d = 0.49 (0.39–0.58) vs no match described d = 0.21 (0.01–0.40); “twice as effective”; Q = 6.1, p = .05
Smith, Domenech Rodríguez & Bernal 2011	65 (8,620)	Between-study	Adapted programmes d = 0.46 (0.36–0.56)	Number of adaptations predicted larger effects; language among the coded adaptations
Benish, Quintana & Wampold 2011	Direct-comparison set	Adapted vs unadapted bona fide programme	d = 0.32 on primary functioning	Illness-myth adaptation was the only significant unique moderator; language match was not
Chowdhary et al. 2014	16 pooled of 20 reviewed	Adapted depression programmes, including LMIC	SMD -0.72 (-0.94 to -0.49)	Language, context and practitioner most common adaptations; 15 of 20 studies adapted language
Hall et al. 2016	78 (13,998)	Adapted vs other conditions; adapted vs unadapted same programme	g = 0.67 overall; g = 0.52 (0.15–0.90) vs unadapted same programme	Language (English vs other) did not significantly moderate the continuous symptom outcome

Source	k (N)	Language contrast	Main finding	Language-specific number
Soto et al. 2018	99	Update of the Smith line	$d = 0.50$; 0.35 after publication-bias correction	Preferred-language delivery $d = 0.59$ vs not reported $d = 0.35$, $p = .03$; translated measures $d = 0.66$ vs 0.28; preferred language survived multivariate control ($\beta = 0.18$)
Arundell et al. 2021	30 RCT comparisons vs active	Adapted vs non-adapted active programmes	$g = -0.43$ (-0.61 to -0.25)	Language-translation (therapist) vs waitlist $g = -0.82$; vs active $g = -0.47$. About half used a same-language or bilingual practitioner or an interpreter

Table 2. Forest-style display of between-study language-matching subgroups (not a new pool)

ये पंक्तियाँ पुरानी मेटा-एनालिसिस के अंदर के सबग्रुप हैं। प्राथमिक L1-बनाम-L2 ट्रायल का नया मॉडल नहीं। वज़न नहीं दिए। पूल नहीं किया।

Contrast	g or d	95% CI	What the row is
Griner & Smith: native-language delivery (non-English)	$d = 0.49$	0.39 to 0.58	Between-study subgroup
Griner & Smith: no language match described	$d = 0.21$	0.01 to 0.40	Between-study subgroup
Implied difference	≈ 0.28	Contrast $p = .05$	Not a within-study g
Soto et al.: preferred-language delivery	$d = 0.59$	—	Between-study subgroup
Soto et al.: language not reported	$d = 0.35$	—	Between-study subgroup
Implied difference	≈ 0.24	$p = .03$	Not a within-study g
Soto et al.: translated measures	$d = 0.66$	—	Measure language, not coach language
Soto et al.: measures not translated	$d = 0.28$	—	Measure language
Hall et al.: adapted vs unadapted same programme	$g = 0.52$	0.15 to 0.90	Cultural adaptation as a bundle
Chowdhary et al.: adapted depression programmes	SMD -0.72	-0.94 to -0.49	Language bundled with other adaptations
Arundell et al.: language-translation vs waitlist	$g = -0.82$	-1.25 to -0.40	Therapist language translation
Arundell et al.: language-translation vs active	$g = -0.47$	-0.73 to -0.20	Therapist language translation
Benish et al.: adapted vs bona fide unadapted	$d = 0.32$	—	Language not the unique moderator

Table 3. Moderator table — pre-specified contrasts

Moderator	Eligible controlled tests	Direction	Certainty
First vs second language, same programme, randomised	0 adult trials	Cannot estimate	Very low
First vs second language, between-study	2 major subgroups (Griner; Soto)	L1 larger by ~0.24 to 0.28	Low
Interpreter-mediated vs same-language practitioner	5 adult comparisons, 1 large	Mixed; 3 similar, 1 worse with interpreter, 1 both large	Very low
Machine-translated vs human L1	0 eligible trials	Unknown	—
Measure translation	Soto 2018 subgroup	Translated measures $d = 0.66$ vs 0.28	Low (still between-study)

Table 4. Interpreter-mediated and language-format primary studies in adults

Study	n	Design	Contrast	Result
Schulz et al. 2006	53	Non-randomised; refugees, US	Same-language practitioner vs interpreter	Both large effects; no significant difference
d'Ardenne et al. 2007	Three groups	Routine-care comparison; East London	Refugees with interpreter vs without vs English speakers	Refugees with and without interpreter did not differ; all groups improved
Brune et al. 2011	190	Retrospective; refugees, Germany	Interpreter vs no interpreter	Equivalent outcomes despite tougher baseline in the interpreter group
Sander et al. 2019	825	Retrospective cohort; Denmark	Interpreter vs no interpreter, 16 CBT-informed sessions	Interpreter arm improved less on the primary trauma measure and most secondaries
Klein Schiphorst, Levi & Manley 2022	133	Pre-post, non-equivalent groups; UK	Same-language counsellor vs interpreter	No significant difference on mood and worry scores; satisfaction high in both arms

Table 5. Indian and South-Asian programmes delivered in a local or preferred language — not L1-versus-L2 trials

Study	n / design	Languages	What was randomised	Why it cannot enter the L1–L2 pool
Patel et al. 2010, MANAS, Goa	Cluster RCT, primary care	Konkani, Marathi, Hindi or English; ~98–99% Konkani	Lay-counsellor collaborative care vs usual care	Language was an eligibility filter, not an arm
Patel et al. 2017, HAP, Goa	495, RCT	Local languages; Konkani-validated symptom measure	HAP plus enhanced usual care vs enhanced usual care	Delivered in L1. Not tested against an English HAP. Remission ratio 1.61; adjusted mean difference –7.57

Study	n / design	Languages	What was randomised	Why it cannot enter the L1–L2 pool
Weobong et al. 2017, HAP 12-month	Same cohort	Same	Same	Gains held at 12 months; still no language arm
Pathare / Joag et al. 2023, Atmiyata, Gujarat	1,191, stepped-wedge cluster RCT	Gujarati community programme	Immediate vs delayed roll-out	Local-language delivery vs later local-language delivery. Recovery OR 2.2 at 3 months
Husain et al. 2024, ROSHNI-2	Multi-site UK RCT	Preferred languages: Urdu, Punjabi, Gujarati, Bengali, Tamil	Culturally adapted group programme vs usual care	Preferred-language delivery in the UK, not an India L1-versus-L2 trial. Recovery OR 1.97 at 4 months

Patel और साथियों का Healthy Activity Programme भारत का सबसे मज़बूत सबूत है। संक्षेप में, आम कोच के हाथ का व्यवहार-एक्टिवेशन प्राइमरी केयर में चल सकता है। क्लिनिक की भाषा में बोला जाए तो। यह भाषा का ट्रायल नहीं है।

एक और वयस्क ट्रायल आसानी से गलत पढ़ा जाता है। Lopez-Maya, Olmstead और Irwin (2019) ने लॉस एंजेलिस के 76 वयस्कों को लॉटरी से बाँटा। स्पेनिश बोलने वाले $n = 36$ । अंग्रेज़ी बोलने वाले $n = 40$ । छह हफ्ते माइंडफुलनेस या स्वास्थ्य शिक्षा। माइंडफुलनेस का फ़ायदा कुल $d = 0.28$ । स्पेनिश फ़ॉर्मेट में 0.29। अंग्रेज़ी फ़ॉर्मेट में 0.30। मतलब प्रोग्राम उस भाषा में चले जिसमें व्यक्ति सच में बोलता है, तो दोनों भाषाओं में चल सकता है। यह L1 बनाम L2 की सीधी लड़ाई नहीं है। लोगों को उनकी पसंद के खिलाफ़ भाषा नहीं दी गई।

फ़र्क / Heterogeneity

नई I^2 नहीं निकाली गई। मूल मेटा-एनालिसिस में फ़र्क बढ़ा था। Soto और साथियों ने स्टडी के बीच बड़ा फ़र्क बताया। छोटे पेपर काटने के बाद कुल असर $d = 0.35$ रह गया। Hall और साथियों ने कहा भाषा का मॉडरेटर पुरानी समीक्षाओं में अस्थिर रहा। Arundell और साथियों ने कई भाषा-अनुवाद सबग्रुप में I^2 80 फ़ीसदी से ऊपर बताया। ईमानदार पढ़ाई यह है। बीच-स्टडी L1 संकेत संस्कृति-अनुकूलन के शोर वाले परिवार में बैठा है।

पब्लिकेशन बायस

Soto और साथी सबसे साफ़ सुधार हैं। कुल $d = 0.50$ पहले। बाद में 0.35। Griner और Smith का “दोगुना असरदार” वाक्य मूल पेपर में नहीं कटा। जो पाठक 0.49 बनाम 0.21 दोहराए, उसे उसी पैरा में Soto का सुधार रखना चाहिए। नया Egger टेस्ट नहीं चला।

पक्षपात का सार

इस पेपर का सवाल है — क्या L1, L2 को हराता है। इस सवाल पर पक्षपात का खतरा ऊँचा है। कंट्रास्ट लगभग कभी लॉटरी से नहीं बँटता। मातृभाषा में प्रोग्राम पाने वाले को अक्सर मिला-जुला कोच भी मिलता है। अनुवादित टेस्ट भी। उनके समूह के हिसाब से कंटेंट भी। गड़बड़ बनावट में है।

मूल सवाल दूसरा है। क्या संस्कृति-अनुकूल प्रोग्राम बिना-अनुकूल से बेहतर हैं। उस पर पक्षपात मिला-जुला है। स्रोत समीक्षाएँ उसे ग्रेड कर चुकी हैं।

Table 6. Risk-of-bias judgements for the language-of-delivery contrast

Study	Language randomised?	Overall for L1–L2 contrast
Griner / Soto language subgroups	No (between-study)	High
Hall 2016 language moderator	No within-study assignment	High for language; lower for adaptation
Klein Schiphorst 2022	No	High
Sander 2019	No; severity likely confounded	High
d’Ardenne 2007; Brune 2011; Schulz 2006	No	High
Patel 2017 HAP	Randomised on programme vs usual care, not language	Low for HAP vs usual care; not applicable for L1 vs L2
Lopez-Maya 2019	Randomised on programme vs education, stratified by language	Low for mindfulness vs education; not applicable for L1 vs L2

GRADE table

Table 7. Certainty of evidence

Claim	Effect	Studies	Certainty
First-language delivery outperforms second-language delivery of the same programme	Between-study $\Delta \approx 0.24$ to 0.28	0 head-to-head RCTs; 2 meta-analytic subgroups	Very low (bias, inconsistency, indirectness, imprecision, suspected publication bias)
Culturally adapted programmes outperform unadapted versions of the same programme	$g = 0.52$ (Hall); $d = 0.32$ (Benish bona fide)	78 studies and the Benish direct-comparison set	Low to moderate
Interpreter-mediated delivery is worse than same-language delivery	Mixed	5 adult comparisons	Very low
Machine-translated delivery matches human L1 delivery	No estimate	0 trials	No evidence
An Indian working-adult L1 advantage exists	No estimate	0 Indian head-to-head trials	Very low

07

चर्चा

खास बात

दुभाषिया वाली चार तुलनाएँ हैं। वे एक राय नहीं रखतीं।

नंबर का मतलब

जो नंबर दोहराया जाएगा वह $g \approx 0.28$ है। या “दोगुना असरदार।” दोनों वाक्य 2006 की बीच-स्टडी तुलना से आए। 2018 में थोड़ा छोटा फ़ासला दोहराया गया। ये उस ट्रायल से नहीं आए जिसमें एक ही वर्कबुक एक ग्रुप को हिंदी में मिली। दूसरे को अंग्रेज़ी में। फिर नींद, चिंता और शुक्रवार शाम देखी गई।

सीधे कहें तो 0.28 छोटा-से-मध्यम धक्का है। चेन्नई का एक प्रोजेक्ट मैनेजर वह विचार लिखता है जिसने उसे रात भर जगाया। ऑफ़िस की अंग्रेज़ी में वह “परिवार की कुछ बेचैनी” बन जाता है। तमिल में वह वही ताना है जो सच में चुभा था। गर्मी उसी में है। शर्म को शर्म कहे, तभी वह अगला अलग क़दम चुन पाता है। भाव असली कपड़ों में चाहिए। “मेरे मन में विचार है कि मैं बुरा बेटा हूँ” — इसे बस विचार की तरह देखना तब काम करता है जब शब्द काट सकें। उसके लिए काटने वाले शब्द तमिल के हैं।

मशीन लर्निंग की भाषा में यह टोकनाइज़र की समस्या है। अंग्रेज़ी पर ट्रेनिंग लिया मॉडल हिंदी वाक्य अनुवाद से निगल सकता है। एम्बेडिंग स्पेस दूसरी सामग्री पर बना। ट्रांसफ़र लर्निंग मदद करती है। यह उस भाषा पर ट्रेनिंग नहीं है जिसमें यूज़र सोचता है। उद्यम की भाषा में यह प्रोडक्ट-मार्केट फ़िट बनाम लोकलाइज़ेशन टॉगल है। 23 भाषाएँ इस शर्त पर हैं कि टॉगल काफ़ी नहीं। पहला स्क्रीन उस भाषा में बोलना चाहिए जिसमें ग्राहक खुद को डाँटता है।

उदार पढ़ाई: अंग्रेज़ी का टाउन-हॉल संदेश लोग लिफ़्ट तक भूल जाते हैं। मैनेजर वही बात कन्नड़ में कहे, तो घर पर दोहराई जाती है। एक ऑफ़िस भाषा स्लाइड पर सस्ती है, शरीर में महँगी। सावधान पढ़ाई पूछती है कि और क्या बदला।

Benish, Quintana और Wampold सावधानी को सबसे तेज़ रूप देते हैं। उन्होंने अनुकूलित प्रोग्राम की तुलना असली बिना-अनुकूल प्रोग्राम से की। फ़ायदा सिकुड़कर $d = 0.32$ रह गया। अकेला मॉडरेटर यह था कि प्रोग्राम ने “बीमारी की कहानी” लिखी या नहीं। क्या गड़बड़ है और क्या मदद करेगा। भाषा-मैच अकेले मॉडल नहीं चलाता। यह L1 परिकल्पना को मारता नहीं। कहता है भाषा शायद ही इकलौता चलता पुर्जा हो।

Hall और साथी दूसरी सावधानी देते हैं। कुल अनुकूलन फ़ायदा असली था। मिला-जुला कंट्रोल पर $g = 0.67$ । बिना-बदले जुड़वाँ पर $g = 0.52$ । लगातार नतीजे पर भाषा का मॉडरेटर भरोसेमंद नहीं रहा। जो समीक्षक Griner के “दोगुना असरदार” पर रुक जाए, वह 2006 का पोस्टर पढ़ रहा है। 2016 की परीक्षा नहीं।

Soto और साथी तीसरी सावधानी और दूसरा नंबर देते हैं। पब्लिकेशन बायस काटने के बाद कुल अनुकूलन असर 0.50 से 0.35 गिर गया। पसंदीदा भाषा फिर भी बढ़ी दिखी। $d = 0.59$ बनाम 0.35 । अनुवादित टेस्ट और बढ़े दिखे। $d = 0.66$ बनाम 0.28 । अगर एक ही बदलाव सस्ता हो, टेस्ट का अनुवाद कोच के अनुवाद जितना ज़रूरी हो सकता है। हिंदी सेशन में अंग्रेज़ी फ़ॉर्म भरना दो काम एक साथ है।

दुभाषिया स्टडी नारा नहीं बनने देतीं। Klein Schiphorst और साथियों को नतीजे में फ़र्क़ नहीं मिला। दोनों बाँहों में संतोष ऊँचा था। d'Ardenne और Brune शरणार्थी क्लीनिक में वैसी ही तस्वीर पाए। Sander और साथी, 825 लोगों के साथ, देखे कि कमरे में दुभाषिया हो तो आघात-लक्षण का फ़ायदा छोटा रहा। गंभीरता और भाषा की पकड़ उस कोहोर्ट को गड्ढमड्ढ करती हैं। दुभाषिया अपने आप दूसरा दर्जा नहीं है। बराबर भी अपने आप नहीं है। डेटा अभी नीति नहीं है।

सीमाएँ

इस रिव्यू की सीमाएँ सबूत की हैं। तरीके की भी हैं।

पहली, पूल नहीं किया। पाँच पात्र ट्रायल की तय मंज़िल पूरी नहीं हुई। जो पाठक L1-बनाम-L2 RCT का नया फ़ॉरेस्ट प्लॉट चाहता था, उसे नहीं मिलेगा। यही नतीजा है।

दूसरी, जो बीच-स्टडी नंबर दिए गए, वे भाषा को संस्कृति, आदत, पैसे, समुदाय-मैच और टेस्ट-अनुवाद से चिपकाते हैं। Griner और Smith ने खुद कहा। “भाषा मिलाना” का कोड अनुमान था। जिन स्टडी ने भाषा न लिखी, उन्होंने शायद मिलाई। जिनने लिखी, उन्होंने शायद दी नहीं।

तीसरी, इन नंबरों के पीछे के सैंपल ज़्यादातर अमेरिका के अल्पसंख्यक वयस्क हैं। दुभाषिया सेट में यूरोप के आघात-ग्रस्त शरणार्थी हैं। बेंगलुरु, पुणे, हैदराबाद, कोच्चि, मुंबई के कामकाजी वयस्क वे सैंपल नहीं हैं। लॉस एंजेलिस के लैटिनक्स क्लीनिक से इंदौर के व्हाट्सऐप तक छल्लाँग उम्मीद है। GRADE अपग्रेड नहीं।

चौथी, यह रिव्यू पहले से रजिस्टर नहीं था। PRISMA गिनती सर्च से फिर बनी। रजिस्टर्ड अपडेट दो स्वतंत्र विकास और बंद प्रोटोकॉल के साथ मज़बूत होगा।

पाँचवीं, न जाँच सकने वाला काम बाहर रहा। यह गुण है। नुकसान यह भी है कि किसी क्षेत्रीय भारतीय जर्नल का असली छोटा ट्रायल बिना DOI के छूट सकता है। नकली पेपर जोड़ने से असली पेपर चूकना बेहतर है।

छठी, मशीन-अनुवाद वाली डिलीवरी — जो हर प्रोडक्ट टीम अभी चलाने वाली है — इस साहित्य में कोई पात्र कंट्रोल ट्रायल नहीं रखती। छेद इजाज़त नहीं है।

भारतीय ऑफ़िस प्रोग्राम के लिए इसका मतलब

इसका मतलब

भारतीय ऑफ़िस प्रोग्राम के लिए इसका मतलब

भारतीय ऑफ़िस रोज़ भाषा का प्रयोग चलाता है। सुबह का स्टैंडअप अंग्रेज़ी में है। लंच की लाइन में अप्रेज़ल पर बड़बड़ाहट बंगाली में है। सिर्फ़ अंग्रेज़ी बोलने वाला कोचिंग प्रोग्राम प्रोफ़ेशनल से दो काम एक साथ माँगता है। रेप ऑफ़िस बोली में। भाव घर की बोली में। यह दिमाग पर अतिरिक्त बोझ है, जो दो मिनट की प्रैक्टिस को होमवर्क बना देता है। पहले 'बाद में', फिर 'कभी नहीं'।

बीच-स्टडी संकेत कहता है मातृभाषा के साथ असर बड़ा दिखता है। करीब $g = 0.28$ । GRADE कहता है इस नंबर को भारतीय कर्मचारियों का कारण-नतीजा न मानें। जनगणना कहती है अंग्रेज़ी मातृभाषा 0.02 फ़ीसदी की है। कम-पछतावे वाला डिज़ाइन साफ़ है। रेप उस भाषा में दो जिसमें व्यक्ति सोचता है। पुणे वाली प्रोफ़ेशनल बहन को मराठी में लिखती है, तो उसका शाम का रेप भी मराठी में आए। अंग्रेज़ी ट्रैक उनके लिए रखो जो चाहें। बाकी को ग्लोबल ब्रांड आवाज़ के नाम पर दूसरी भाषा की छलनी से न निकालो।

पहले व्हाट्सऐप इसलिए मदद करता है कि कोड-स्विच वहीं रहता है। परिवार का ग्रुप मलयालम में, प्रोजेक्ट का ग्रुप अंग्रेज़ी में, अंगूठा वही। 90 सेकंड की साँस वाले रेप को लेक्चर नहीं चाहिए। उसे ऐसा प्रॉम्प्ट चाहिए जो लोकल ट्रेन और ऑफ़िस की लिफ़्ट के बीच अंगूठे से खुल जाए। शुक्रिया वाला रेप तब लगता है जब उस भाषा में लिखा जाए जिसमें दूसरा व्यक्ति महसूस करता है। मूल्य वाला अभ्यास मर जाता है अगर "मैं चाहता हूँ पापा को मुझ पर गर्व हो" को "demonstrates ownership" बनाना पड़े। प्रोग्राम गिने रेप्स की तरह बनाओ। 23 भाषाओं में। मॉडल अटकने पर इंसान का हाथ दो। फिर वह गायब ट्रायल चलाओ। वही रेप्स। हिंदी बनाम अंग्रेज़ी। तमिल बनाम अंग्रेज़ी। लॉटरी से। कामकाजी वयस्क। 8 हफ़्ते। वेलबीइंग और लक्षण। डेटा भारत में। जब तक वह ट्रायल नहीं है, 23 भाषाएँ ज़्यादा इंजन नहीं हैं। यह उस बात की विनम्रता है जो नहीं पता। चुकाई उस भाषा में जिसमें लोग पहले से बोलते हैं।

रिसर्च के छेद, भारत का छेद नाम लेकर

भारत का छेद यह नहीं कि वेलबीइंग ट्रायल नहीं हैं। भारत के पास HAP है। MANAS है। आत्मीयता है। SMART Mental Health है। ब्रिटेन में दक्षिण-एशियाई वयस्कों के पसंदीदा-भाषा प्रोग्राम भी बढ़े। छेद पतला और अटपटा है। कोई जाँची भारतीय स्टडी एक प्रोग्राम लेकर भाषा की लॉटरी नहीं चलाती।

वह ट्रायल बन सकता है। आबादी हर बड़ी कंपनी की द्विभाषी टीम में बैठी है। प्रोग्राम ड्राफ़्ट में है। व्यवहार एक्टिवेशन, शुक्रिया, माइंडफुलनेस, मूल्य, मूड जो माँगे उसका उलटा क्रदम। सब गिने रेप्स। तुलना वही स्क्रिप्ट अंग्रेज़ी में। कई अनुसूचित भाषाओं में

नतीजे के टेस्ट पहले से हैं। कमी रजिस्ट्रेशन की है। सैंपल साइज़ की है। शून्य नतीजे का जोखिम उठाने की हिम्मत की है।

और छेद पीछे चलते हैं। मशीन अनुवाद बनाम द्विभाषी इंसान कोच नहीं जाँचा गया। भारतीय भाषाओं में दुभाषिया कोचिंग कामकाजी सैंपल पर नहीं जाँची गई। डोज़ अज्ञात है। L1 के कितने मिनट L2 के कितने मिनट के बराबर। Soto का 0.66 बनाम 0.28 का ज़्यादातर हिस्सा अकेले टेस्ट-अनुवाद पकड़ता है या नहीं, पता नहीं। पुरानी अमेरिकी स्टडी की आदत-पैटर्न कहती है L1 फ़ायदा कम अंग्रेज़ी वालों में बड़ा हो सकता है। भारत में यह अज्ञात है। यही मॉडरेटर ख़रीद बदलेगा।

उस आने वाले ट्रायल का शून्य नतीजा भी छपने लायक होगा। छोटा L1 फ़ायदा ईएपी ख़रीद बदलेगा। बड़ा L1 फ़ायदा हर भारतीय कोचिंग प्रोडक्ट का इंजन बदलेगा। इनमें से कोई वाक्य आज ईमानदारी से नहीं लिखा जा सकता।

08

FAQ

Q1 क्या मातृभाषा में कोचिंग अंग्रेज़ी से बेहतर चलती है?

बीच-स्टडी डेटा कहता है हाँ। करीब $g = 0.28$ । यानी असर लगभग दोगुना। भारत के वयस्कों पर एक ही प्रोग्राम L1 बनाम L2 की लॉटरी नहीं चली। ईमानदार जवाब: शायद हाँ। साफ़ प्रयोग अभी नहीं है।

Q2 हिंदी, तमिल या मराठी बनाम अंग्रेज़ी की कितनी भारतीय स्टडी हैं?

शून्य। Healthy Activity Programme जैसे बड़े भारतीय प्रोग्राम स्थानीय भाषा में चले। भाषा खुद लॉटरी का विषय नहीं बनी। चल सकना लड़ाई नहीं है।

Q3 दुभाषिया चाहिए तो प्रोग्राम बेकार है?

चार वयस्क तुलनाएँ हैं। तीन में दुभाषिया के साथ और बिना उसके नतीजे मिलते-जुलते रहे। उनमें ब्रिटेन की 133 लोगों की एक स्टडी में दोनों बाँहों में संतोष ऊँचा था। 825 लोगों की एक कोहोर्ट में दुभाषिया के साथ फ़ायदा छोटा था। दुभाषिया डिफ़ॉल्ट दूसरा दर्जा नहीं है।

Q4 क्या मशीन-अनुवादित स्क्रिप्ट मातृभाषा कोच की जगह ले सकती है?

बहुभाषी वयस्कों पर मशीन-अनुवादित मनोवैज्ञानिक प्रोग्राम का कोई पात्र कंट्रोलड ट्रायल नहीं मिला। यह सबूत का छेद है। हरी झंडी नहीं। बिना द्विभाषी कोच के अनुवाद एक अनुमान है।

Q5 पूल ट्रायल नहीं तो इमोशनऐप 23 भाषाएँ क्यों देता है?

क्योंकि 0.02 फ्रीसदी भारतीयों की मातृभाषा अंग्रेज़ी है। बीच-स्टडी संकेत भाषा मिलाने के पक्ष में है। भाषा देना कम-पछतावे वाला क़दम है। जब तक हेड-टू-हेड ट्रायल गायब है।

09

संदर्भ

Arundell, L.-L., Barnett, P., Buckman, J. E. J., Saunders, R., & Pilling, S. (2021). The effectiveness of adapted psychological interventions for people from ethnic minority groups: A systematic review and conceptual typology. *Clinical Psychology Review*, 88, 102063.

<https://doi.org/10.1016/j.cpr.2021.102063>

Benish, S. G., Quintana, S., & Wampold, B. E. (2011). Culturally adapted psychotherapy and the legitimacy of myth: A direct-comparison meta-analysis. *Journal of Counseling Psychology*, 58(3), 279–289. <https://doi.org/10.1037/a0023626>

Brune, M., Eiroá-Orosa, F. J., Fischer-Ortman, J., Delijaj, B., & Haasen, C. (2011). Intermediated communication by interpreters in psychotherapy with traumatized refugees. *International Journal of Culture and Mental Health*, 4(2), 144–151. <https://doi.org/10.1080/17542863.2010.537821>

Chowdhary, N., Jotheeswaran, A. T., Nadkarni, A., Hollon, S. D., King, M., Jordans, M. J. D., Rahman, A., Verdeli, H., Araya, R., & Patel, V. (2014). The methods and outcomes of cultural adaptations of psychological treatments for depressive disorders: A systematic review. *Psychological Medicine*, 44(6), 1131–1146. <https://doi.org/10.1017/S0033291713001785>

d'Ardenne, P., Ruaro, L., Cestari, L., Fakhoury, W., & Priebe, S. (2007). Does interpreter-mediated CBT with traumatized refugee people work? A comparison of patient outcomes in East London. *Behavioural and Cognitive Psychotherapy*, 35(3), 293–301.

<https://doi.org/10.1017/S1352465807003645>

Griner, D., & Smith, T. B. (2006). Culturally adapted mental health intervention: A meta-analytic review. *Psychotherapy: Theory, Research, Practice, Training*, 43(4), 531–548.

<https://doi.org/10.1037/0033-3204.43.4.531>

Hall, G. C. N., Ibaraki, A. Y., Huang, E. R., Marti, C. N., & Stice, E. (2016). A meta-analysis of cultural adaptations of psychological interventions. *Behavior Therapy*, 47(6), 993–1014.

<https://doi.org/10.1016/j.beth.2016.09.005>

Husain, N., et al. (2024). Efficacy of a culturally adapted, cognitive behavioural therapy-based intervention for postnatal depression in British South Asian women (ROSHNI-2). *The Lancet*.

[https://doi.org/10.1016/S0140-6736\(24\)01612-X](https://doi.org/10.1016/S0140-6736(24)01612-X)

- Klein Schiphorst, A. T., Levi, N., & Manley, J. (2022). Found in translation? The effectiveness of counselling using an interpreter, compared to counselling with a same language counsellor, in a non-English speaking population. *International Journal for the Advancement of Counselling*, 44, 628–645. <https://doi.org/10.1007/s10447-022-09475-z>
- Lopez-Maya, E., Olmstead, R., & Irwin, M. R. (2019). Mindfulness meditation and improvement in depressive symptoms among Spanish- and English speaking adults: A randomized, controlled, comparative efficacy trial. *PLOS ONE*, 14(7), e0219425. <https://doi.org/10.1371/journal.pone.0219425>
- Office of the Registrar General & Census Commissioner, India. (2011). Census of India 2011: Data on language and mother tongue. <https://censusindia.gov.in/>
- Patel, V., Weiss, H. A., Chowdhary, N., Naik, S., Pednekar, S., Chatterjee, S., De Silva, M. J., Bhat, B., Araya, R., King, M., Simon, G., Verdelli, H., & Kirkwood, B. R. (2010). Effectiveness of an intervention led by lay health counsellors for depressive and anxiety disorders in primary care in Goa, India (MANAS): A cluster randomised controlled trial. *The Lancet*, 376(9758), 2086–2095. [https://doi.org/10.1016/S0140-6736\(10\)61508-5](https://doi.org/10.1016/S0140-6736(10)61508-5)
- Patel, V., Weobong, B., Weiss, H. A., Anand, A., Bhat, B., Katti, B., Dimidjian, S., Araya, R., Hollon, S. D., King, M., Vijayakumar, L., Park, A. L., McDaid, D., Wilson, T., Velleman, R., Kirkwood, B. R., & Fairburn, C. G. (2017). The Healthy Activity Program (HAP), a lay counsellor-delivered brief psychological treatment for severe depression, in primary care in India: A randomised controlled trial. *The Lancet*, 389(10065), 176–185. [https://doi.org/10.1016/S0140-6736\(16\)31589-6](https://doi.org/10.1016/S0140-6736(16)31589-6)
- Pathare, S., Joag, K., Kalha, J., Pandit, D., Krishnamoorthy, S., Chauhan, A., et al. (2023). Atmiyata, a community champion led psychosocial intervention for common mental disorders: A stepped wedge cluster randomized controlled trial in rural Gujarat, India. *PLOS ONE*, 18(6), e0285385. <https://doi.org/10.1371/journal.pone.0285385>
- Sander, R., Laugesen, H., Skammeritz, S., Mortensen, E. L., & Carlsson, J. (2019). Interpreter-mediated psychotherapy with trauma-affected refugees: A retrospective cohort study. *Psychiatry Research*, 271, 684–692. <https://doi.org/10.1016/j.psychres.2018.12.058>
- Schulz, P. M., Resick, P. A., Huber, L. C., & Griffin, M. G. (2006). The effectiveness of cognitive processing therapy for PTSD with refugees in a community setting. *Cognitive and Behavioral Practice*, 13(4), 322–331. <https://doi.org/10.1016/j.cbpra.2006.04.011>
- Smith, T. B., Domenech Rodríguez, M., & Bernal, G. (2011). Culture. *Journal of Clinical Psychology*, 67(2), 166–175. <https://doi.org/10.1002/jclp.20757>
- Smith, T. B., & Trimble, J. E. (2016). Culturally adapted mental health services: An updated meta-analysis of client outcomes. In T. B. Smith & J. E. Trimble, *Foundations of multicultural psychology: Research to inform effective practice* (pp. 129–144). American Psychological Association. <https://doi.org/10.1037/14733-007>
- Soto, A., Smith, T. B., Griner, D., Domenech Rodríguez, M., & Bernal, G. (2018). Cultural adaptations and therapist multicultural competence: Two meta-analytic reviews. *Journal of Clinical Psychology*, 74(11), 1907–1923. <https://doi.org/10.1002/jclp.22679>
- Weobong, B., Weiss, H. A., McDaid, D., Singla, D. R., Hollon, S. D., Nadkarni, A., Park, A. L., Bhat, B., Katti, B., Anand, A., Dimidjian, S., Araya, R., King, M., Vijayakumar, L., Wilson, G. T., Velleman, R., Kirkwood, B. R., Fairburn, C. G., & Patel, V. (2017). Sustained effectiveness and cost-effectiveness of the Healthy Activity Programme, a brief psychological treatment for depression delivered by lay counsellors in primary care: 12-month follow-up of a randomised controlled trial. *PLOS Medicine*, 14(9), e1002385. <https://doi.org/10.1371/journal.pmed.1002385>

लेखक के बारे में

लेखक के बारे में

डॉ. अतुल असवानी; मुंबई के साइकियाट्रिस्ट और ट्रेंड रिज़ल्ट्स कोच
डॉ. अतुल असवानी मुंबई के साइकियाट्रिस्ट और ट्रेंड रिज़ल्ट्स कोच हैं। उन्होंने इमोशनऐप शुरू किया। यह भारतीय प्रोफेशनल्स का नॉन-क्लिनिकल इमोशनल फ़िटनेस प्लेटफ़ॉर्म है। प्रैक्टिस गिने रोज़ के रेप्स की तरह मिलती है। एआई कोच है। इंसान का हाथ भी। पहले व्हाट्सएप। 23 भारतीय भाषाएँ। डेटा भारत में रहता है। DPDP का पालन। काम ISO 9001 और ISO 13485 के नीचे है।

हवाला कैसे दें

Aswani A. मातृभाषा में कोचिंग: असर 0.28 बढ़ा — भारत में ट्रायल शून्य. EmotionApp Research Paper 39 (Hindi edition). Mumbai: Tranquilo Meditek Sytems Private Limited; 2026.
<https://emotionapp.in/research/language-of-delivery-meta-analysis>

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 प्रोफेशनल से या Tele-MANAS (14416) पर बात करें।