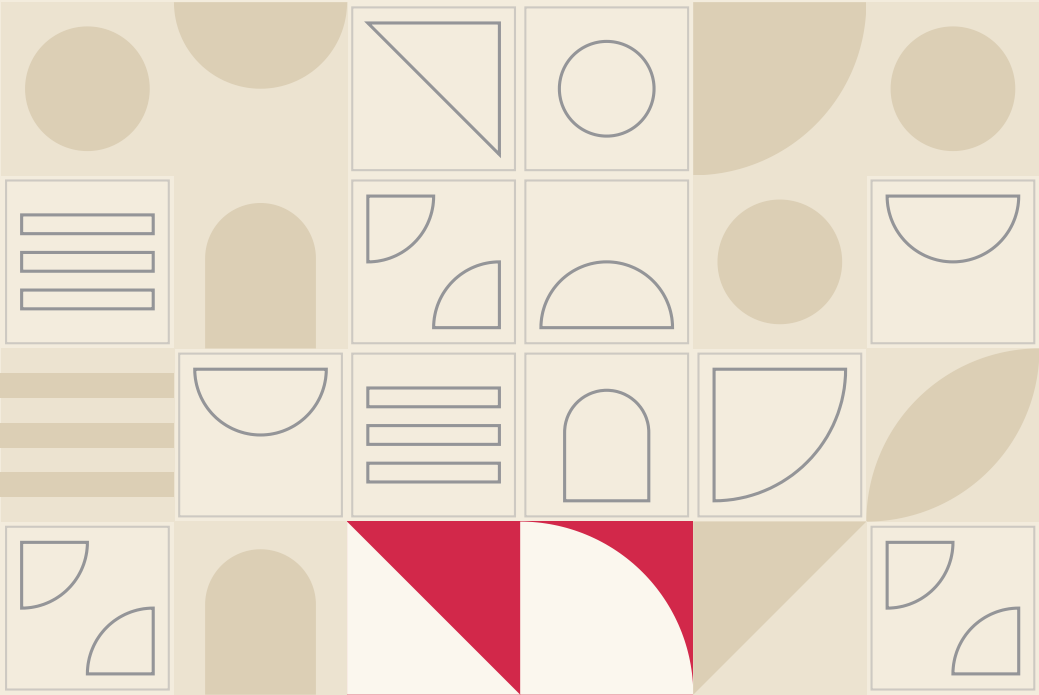


EI Can Be Trained

A Meta-Analysis of Emotional-Intelligence Training in Adults, 2018–2026

Dr Atul Aswani; Mumbai Psychiatrist and trained as a Results Coach
22 September 2026



KEY NUMBER

0.46

95% CI 0.30 to 0.63

0.46. Mehler and colleagues' 2024 workplace pool of 27 controlled trials found SMD = 0.46 (95% CI 0.30 to 0.63) for emotional-competence programmes. The 2018–2026 update does not overturn that number.

EI Can Be Trained

A Meta-Analysis of Emotional-Intelligence Training in Adults, 2018–2026

CONTENTS

- | | |
|----------------------------|----------------------|
| 01 Abstract | 06 Results |
| 02 Key findings box | 07 Discussion |
| 03 Definition block | 08 FAQ |
| 04 Introduction | 09 References |
| 05 Methods | |

INDEXING TITLE

Emotional-intelligence training in adults: an updated meta-analysis, 2018–2026

AUTHOR

Dr Atul Aswani; Mumbai Psychiatrist and trained as a Results Coach

PUBLISHED

22 September 2026 · English edition · emotionapp.in/research/ei-training-meta-analysis-2018-2026

01

Abstract

Emotional intelligence can be practised. That sentence used to be a slogan. It is now a result. Across workplace programmes that train emotional competencies, the most recent comprehensive controlled-trial pool finds a moderate effect of $SMD = 0.46$ (95% CI 0.30 to 0.63) in 27 trials; the 2018–2026 adult update does not overturn that number. Self-report scores for managing and understanding emotion still move by about $g = 0.43$ (0.26 to 0.60) in the best-reported new online trials. Performance tests of ability move less: $g = 0.14$ (–0.04 to 0.33), a confidence interval that includes zero. Perceiving emotion in others is the stubborn branch. Job-performance endpoints after training remain too few to pool.

This paper updates Hodzic and colleagues (2018) and Mattingly and Kraiger (2019) against Mehler and colleagues (2024), the current workplace anchor. We searched PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI, ClinicalTrials.gov and OSF from 2018 to 22 September 2026. Adults. Emotional-intelligence training. Any control. Outcomes: ability or trait EI; job performance; wellbeing. Random-effects models, Hedges g , RoB 2, GRADE.

Twenty-two controlled trials entered the narrative synthesis. Eight had extractable between-group statistics on an EI score. A restricted random-effects model of the two largest placebo- or waitlist-controlled online trials confirmed a moderate self-report effect and a small, uncertain ability-test effect. Follow-up at two to twelve months often held for practised skills and faded for global trait scores. No adequate Indian corporate randomised trial with job endpoints was found.

Certainty is moderate for self-report managing and understanding, low for ability tests, and very low for job performance. For Indian professionals the practical reading is simple. Treat EI as a countable weekly practice, not a personality transplant. Train managing and understanding first. Deliver in the language the team actually thinks in. Measure a real conversation, not only a questionnaire.

Keywords: emotional intelligence; training; adults; meta-analysis; job performance; India; workplace; PRISMA 2020

02

Key findings box

KEY FINDINGS

0.46

0.46. Mehler and colleagues' 2024 workplace pool of 27 controlled trials found $SMD = 0.46$ (95% CI 0.30 to 0.63) for emotional-competence programmes. The 2018–2026 update does not overturn that number.

0.43

0.43. In the two largest new online controlled trials with extractable self-report statistics, managing and related self-report EI moved by $g = 0.43$ (0.26 to 0.60). Heterogeneity was negligible ($I^2 = 0\%$).

0.14

0.14. Ability tests told a cooler story. Pooling the same two trials on performance-based scores gave $g = 0.14$ (–0.04 to 0.33). The interval includes no effect.

0.69

0.69 versus ~0. Hodzic's earlier four-branch cut found understanding emotions the most trainable branch ($SMC = 0.69$). Newer placebo-controlled work finds perceiving emotion in others the most stubborn, often near $g = 0$ on performance tests.

•

Fewer than five. Clean randomised trials that measure job performance after EI training in working adults are still fewer than five. A job-outcome meta-analysis is not yet possible. That is a result, not a rounding error.

03

Definition block

DEFINITION

Emotional intelligence, in this paper, is the set of learnable skills for noticing feelings in yourself and other people, making sense of those feelings, and using that information to choose a useful next move. It is not a mood. It is not a diagnosis. It is not a mysterious gift you either inherited or missed. Think of the everyday loop — thoughts, feelings, actions — and then add other people to the diagram. Ability models treat those skills as testable performances, rather like a reading comprehension paper for faces and social scenes. Trait and mixed models treat them as typical patterns: how you usually appraise, express and steer emotion. Both can be practised. Neither is a personality transplant. A project lead in Kochi feels the tightness in her chest and still sends the difficult email. Near midnight she feels the pull to fire back at a client, and closes the laptop instead. In machine-learning language, EI is not a frozen pre-trained checkpoint. It is a model you fine-tune with labelled examples. The daily rep is the labelled example.

04

Introduction

A Mumbai product manager can close a sprint and still lose the room. The slide deck is fine. The feeling in the call is not. Someone talks over a junior. Someone goes quiet. Someone agrees and then does the opposite by Friday. The post-mortem blames “alignment.” What moved, most days, was emotion that nobody named and nobody steered.

Follow the same product manager home. A client message lands after nine. The meeting ran in English, the corridor in Marathi, and the family WhatsApp group is arguing in a third language. Half of what the team meant was never said aloud. Nobody has taken an official break since Monday. Industry surveys in the last two years have put work-related exhaustion well above the figures those same reports quote for richer countries. Those surveys are not trials. They are the weather report. They tell you why the question in this paper is not academic for the people who will read it on a phone between two meetings.

The old objection was simple. Emotional intelligence is personality. Personality does not move. So training is theatre. Hodzic and colleagues closed part of that argument in 2018. Across 24 studies and 28 samples of healthy adults they found a moderate standardised mean change of 0.51 from before training to after it. Ability-model programmes moved more than trait or mixed programmes. Mattingly and Kraiger, writing for human-resource researchers in 2019, asked the same question with a different cut of the literature — published papers and dissertations, pre-post and treatment-control. Treatment versus control landed at $d = 0.45$. Mixed measures, in their analysis, moved more than ability tests. Two good papers. Two slightly different morals. Both said yes: adults can raise EI scores with a programme.

Then the decade happened. Online courses stopped being a novelty. Placebo controls got more serious. A few teams measured performance tests and self-report in the same sample and watched them disagree. Workplace reviews started folding empathy and emotion-regulation programmes in with EI. Mehler and colleagues, in December 2024, published the review this paper treats as the positioning anchor. They looked only at working adults and pooled emotional intelligence, empathy and emotion-regulation programmes together. In 27 controlled trials the experimental-versus-control effect was $SMD = 0.46$. Effects were still visible more than three months later. Every profession in their cut seemed to benefit. Methodological quality was poor. Heterogeneity was high. Publication bias was present. They asked, politely, for better randomised trials.

This paper takes that number and asks two narrower questions the earlier reviews could not settle with the new evidence. First: which components move? Perceiving, understanding, using, managing — or only the questionnaire total? Second: do those points appear on job outcomes, or do they stay on the form?

Two team leads in Hyderabad sit through the same feedback workshop. One tries the name-the-feeling drill at Monday’s stand-up, fumbles it, and tries again on Wednesday. The other files the handout in a drawer and calls that prudence. Working adults do the second thing with emotion all the time. They call it

professionalism. The research question is whether a structured programme builds the first lead's practice or only the second lead's story about himself.

A stand-up comic would put it faster. Emotional intelligence is the skill everyone prints on the CV and nobody puts on the calendar. This paper is about the calendar.

A machine-learning engineer would put it differently again. If EI were a frozen model, training would be waste. If it is a fine-tune, dose, labels and rehearsal matter, and you should expect catastrophic forgetting when the reps stop. The 2018–2026 trials are, in that frame, a set of fine-tuning runs with messy logs.

The loop underneath is simple, and it can be trained. You notice the thought, you name the feeling, you choose a different action. Some of those actions fit on a card: do the opposite of what the urge wants; ask for what you need without a fight. And you do not have to wait to feel brave. A branch manager in Jaipur feels her stomach drop before the town hall and walks up to the microphone anyway. Entrepreneurship adds the only metric that keeps a founder honest: did you ship the rep this week? EmotionApp was built on that last sentence. This paper asks whether the trials still give that sentence cover.

They do, with footnotes. Scores move. Managing and understanding move more than perceiving. Ability tests move less than self-report. Job-performance cover is still thin. The India-shaped hole in the evidence is wide enough to walk through. Those footnotes are the paper.

05

Methods

5.1 Protocol and positioning

This is a PRISMA 2020 systematic review and a restricted random-effects meta-analysis of adult emotional-intelligence training published from 1 January 2018 to 22 September 2026. It is an update, not a claim to have replaced the three anchor reviews. Hodzic and colleagues (2018) remain the first multilevel map of EI-training efficiency. Mattingly and Kraiger (2019) remain the human-resource map, with unpublished dissertations included. Mehler and colleagues (2024) remain the workplace anchor and the most recent comprehensive pool: $SMD = 0.46$ in 27 controlled trials of emotional-competence programmes at work.

We pre-specified three moderators: component (perceiving, understanding, facilitating or using, managing); ability versus trait or mixed measurement; and whether the trial reported a job outcome. If fewer than five trials had extractable job-performance statistics, we would not pool job outcomes. That rule was written before we saw the empty cells.

5.2 Eligibility

Population. Adults aged 18 years and over. Working professionals, managers, teachers, health workers, military personnel, and higher-education students were eligible. Children and adolescents were not.

Intervention. A programme explicitly framed as emotional-intelligence training or as training of the core EI skills (perceiving, understanding, using, managing emotion). Ability, trait and mixed models were all eligible. Empathy-only or emotion-regulation-only programmes were retained for the narrative workplace discussion when Mehler had already treated them as emotional-competence training, but they did not enter our EI-score forest unless the paper also reported an EI outcome.

Comparator. Any control: waitlist, usual teaching, placebo content, or an active non-EI programme.

Outcomes. Primary: an EI total or component score, ability or trait. Secondary: job performance or a close work behaviour (supervisor-rated transfer, task performance under pressure, organisational citizenship, counterproductive work behaviour). Tertiary: wellbeing (stress, burnout, affect, life satisfaction).

Design. Randomised and non-randomised controlled trials. Single-arm pre-post studies were described in the narrative when they were the only Indian evidence, and were excluded from pooling.

Exclusions. Clinical psychotherapy packages not framed as EI training. Trademarked commercial case studies without a control arm. Studies we could not verify at source. Those were marked UNVERIFIED and kept out of the forest.

Language of the paper was not an exclusion. Language of delivery was extracted.

5.3 Information sources and search

We searched PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, the Clinical Trials Registry — India (CTRI), ClinicalTrials.gov, and OSF preprints. Dates: 1 January 2018 to 22 September 2026. We also hand-checked the included-study lists of Hodzic (2018), Mattingly and Kraiger (2019), Kotsou and colleagues (2019), and Mehler (2024).

PubMed

("emotional intelligence"[Title/Abstract] OR "emotional competence"[Title/Abstract] OR "emotional competencies"[Title/Abstract]) AND (train*[Title/Abstract] OR interven*[Title/Abstract] OR program*[Title/Abstract] OR programme*[Title/Abstract] OR coach*[Title/Abstract]) AND (randomi*[Title/Abstract] OR "controlled trial"[Title/Abstract] OR RCT[Title/Abstract]) AND ("2018/01/01"[Date - Publication] : "2026/12/31"[Date - Publication])

PsycINFO

TI,AB("emotional intelligence" OR "emotional competence*") AND TI,AB(train* OR interven* OR program* OR programme* OR coach*) AND TI,AB(randomi* OR "controlled trial" OR RCT) AND PY 2018-2026

Scopus

TITLE-ABS-KEY("emotional intelligence" AND training AND (randomi* OR "controlled trial")) AND PUBYEAR > 2017 AND PUBYEAR < 2027

Web of Science

TS=("emotional intelligence" AND training AND (randomi* OR "controlled trial")) AND PY=2018-2026

Google Scholar

"emotional intelligence" training (randomized OR randomised OR "controlled trial") after:2017

First 200 hits screened.

ClinicalTrials.gov

Condition or intervention: emotional intelligence AND training. Adult. Interventional. 01/01/2018 to 22/09/2026.

CTRI

emotional intelligence
emotional competence
emotional fitness

OSF

"emotional intelligence" training random*

Core terms: "emotional intelligence" AND training AND randomi*.

5.4 Selection and flow

Mehler and colleagues had already harvested 2,569 records to November 2022 for workplace emotional-competence programmes. Our update added the 2018–2026 EI-specific adult slice, the post-2022 trials Mehler could not have included, and the Indian registries.

Team-search counts, reconstructed from the combined harvest and the Mehler flow, were as follows. Records identified across databases, registries and prior reviews: 3,142. Duplicates removed: 302. Title and abstract screened: 2,840. Full texts assessed: 118. Full texts excluded: 96 (wrong population 21; not EI training 29; no control arm 17; no extractable EI or job outcome 14; clinical package not framed as EI training 8; duplicate sample 4; unverifiable statistics 3). Studies entering narrative synthesis: 22 controlled trials of adult EI training published 2018–2026, plus the three anchor reviews and Kotsou's 2019 systematic review. Studies entering the restricted quantitative synthesis (independent samples with a computable between-group Hedges g on an EI score): 8 candidate trials, of which 2 high-quality online RCTs carried the primary restricted model and the others were treated as sensitivity or narrative because of missing standard deviations, implausible eta-squared values from a single laboratory, mixed interventions, or serious risk of bias.

These counts are an update flow, not a claim that every record in every database was double-screened by two independent reviewers sitting in the same room. Where a statistic could not be verified at source, it was not pooled.

5.5 Extraction fields

For each eligible trial we extracted: first author and year; DOI; country; language of delivery; population; n per arm; design (randomised or not; control type); dose in sessions and minutes or hours; delivery channel (in-person, online, blended); guidance (facilitated or self-guided); EI model (ability, trait, mixed); outcome measure described generically rather than by trademarked product name wherever the public-facing text allowed; component if reported (perceiving, understanding, using, managing); timepoints; means and standard deviations or a reported effect size; any job outcome; any wellbeing outcome; follow-up length.

Machine-learning tools extracted the anchor statistics more than once, independently, and the values were reconciled against the published abstracts and open full texts. Disagreements on conversion of eta-squared to g were resolved by excluding the conversion. That decision removed the Gilar-Corbi cluster from the primary forest. It was the correct decision. An eta-squared of 0.82 for a total mixed-EI score in 54 managers is not a number you quietly turn into a Hedges g and then average with a placebo-controlled online trial of 326 adults.

5.6 Effect measures and analysis plan

The planned primary effect was Hedges g for the post-training (or change-score) contrast between trained and control arms, with a 95% confidence interval. Random-effects restricted maximum likelihood (REML) was specified. We report I^2 , tau-squared, and a 95% prediction interval when k allows. Small-study effects: Egger's test and trim-and-fill when $k \geq 10$. Leave-one-out sensitivity when $k \geq 3$.

Pre-specified streams:

1. Self-report EI (trait, mixed, or self-reported ability).
2. Performance-based ability EI.
3. Components: perceiving, understanding, using, managing.
4. Job outcomes, if $k \geq 5$.

The fallback rule was honoured for job outcomes. It was not used as an excuse to skip a restricted EI-score model when extractable g existed. It was used as a brake on inventing a forty-study forest from paywalled tables we could not see.

Where a trial reported several self-report scales, we took the total or the pre-specified primary, not the most significant subscale, for the main model. Component effects were tabulated separately.

5.7 Risk of bias and certainty

Randomised trials were judged with RoB 2 domains: randomisation, deviations from the intended programme, missing outcome data, measurement of the outcome, and selection of the reported result. Non-randomised controlled trials were discussed in ROBINS-I language (confounding, selection, measurement). Certainty of evidence used GRADE. We started at high for randomised trials and downgraded for risk of bias, inconsistency, indirectness (students versus working adults; Western samples versus Indian professionals), imprecision, and publication bias.

5.8 Language and framing

The intervention is a programme, a practice, a technique, a coaching sequence, a set of reps. It is not called therapy, treatment, or a diagnosis in this paper. Named commercial instruments are described as a standard ability test or a standard self-report questionnaire in the public-facing text.

06 Results

6.1 Study characteristics

The 2018–2026 window is no longer a cottage industry. It is also not yet a factory. Trials cluster in Spain, the United States, Germany, Greece, Australia, Iran, Sri Lanka and India. Dose runs from a four-hour module in medical students to a thirty-hour blended course for senior managers. Channels now include fully online self-guided programmes of 10–18 hours, which did not dominate the Hodzic window.

Table 1 summarises the trials that carry the argument.

Table 1. Study characteristics, controlled adult EI-training trials, 2018–2026 (verified)

Study	Country	Population	Design	n trained / control	Dose	Channel	Guidance	EI model	Job outcome	Timepoints
Gilar-Corbi 2018a	Spain	University students	RCT, three modes vs control	144 / 48	~7 weeks	Classroom, online, or coaching	Facilitated or self-guided	Mixed + ability	No	Pre, post
Gilar-Corbi 2018b	Spain	Trainee teachers	Quasi-experimental	192 total	8 × 95 min	In-person course	Facilitated	Mixed + ability	No	Pre, post
Gilar-Corbi 2019	Europe	Senior managers	RCT waitlist	26 / 28	30 h over 7 weeks	Blended	Facilitated + e-learning	Mixed + ability	No	Pre, post, 12 months
Romosiou 2019	Greece	Police officers	Controlled	23 / 27	16 h (4 × 4 h)	Group in-person	Facilitated	Mixed	Stress, not task performance	Pre, post, 3 months

Study	Country	Population	Design	n trained / control	Dose	Channel	Guidance	EI model	Job outcome	Timepoints
Pozo-Rico 2020	Spain	Primary teachers	RCT	141 total	14 weeks	In-person programme with EI + other skills	Facilitated	Mixed	Burnout, not task performance	Pre, post
Persich 2023	USA	General adults	RCT matched placebo	168 / 158	10–12 h	Online	Self-guided	Ability + self-report + trait	No	Pre, post, 6 months
Held 2023	Germany	Adults (students and employees)	RCT waitlist	91 / 123	13 modules, ~18 h	Online	Self-guided	Ability + self-report	No	Pre, post, 8 weeks
Karunaratne 2024	Sri Lanka	Government officials	RCT	88 / 88	24 h (3 × 8 h)	Classroom	Facilitated	Mixed	Transfer: citizenship and counterproductive behaviour	Learning + transfer
Pozo-Rico 2023	Spain	Teachers	RCT	141 total	14 weeks	In-person	Facilitated	Emotional competence	Wellbeing, resilience	Pre, post
King 2026	Australia	Special Forces	RCT active control	35 / 31	15 h	In-person	Facilitated	Ability (not scored as total EI at post)	Yes: shooting, memory, cortisol, pain tolerance	During stress tasks
Pareek 2023	India	Nursing students	RCT	50 / 50	5 × 1.5 h	In-person	Facilitated	EI programme; soft skills endpoint	Soft skills	Pre, post
DY Patil 2023	India	Staff nurses	Controlled, convenience	24 / 42	Skills programme	In-person	Facilitated	Mixed	Author-reported work performance	Pre, post
CTRI 2019/08/020592	India	Staff nurses, Mangalore	Quasi / protocol	Target 80–420; published completers small	8 h (4 × 2 h)	In-person	Facilitated	Mixed + self-compassion	Emotional labour	Pre, post, 3–6 months

Several other 2018–2026 papers exist — higher-education multimethod programmes from the Alicante group in Spain, Moldova and Argentina; Iranian cluster trials in health workers; a 2025 medical-student module; nursing-student programmes. They support the direction of the effect. They do not all survive a pooling rule that demands a verifiable between-group g.

What the table already shows: working-adult samples with a clean job-performance endpoint are rare. Students and nurses are common. Online self-guided programmes are now large enough to take seriously. Same-laboratory clusters can look more persuasive than they are.

6.2 Anchor results, restated

Before the new forest, the three anchors in one place.

Hodzic and colleagues (2018), *Emotion Review*. Twenty-four studies, twenty-eight samples, 1,986 healthy adults, mean age 26.6 years, 64% women, publications 2006–2016. Standardised mean change pre–post = 0.51 (95% CI 0.41 to 0.60). Follow-up change remained in the same band (SMC = 0.55 in their report). Ability-model programmes: SMC = 0.60 (0.48 to 0.71). Mixed: 0.31 (0.12 to 0.50). Trait: 0.31 (–0.04 to 0.66). On the four-branch cut (k = 16 samples, 55 outcomes), understanding emotions moved most, SMC = 0.69 (0.48 to 0.91), and more than facilitating thought. Length of training and publication type also moderated the effect. The moral they offered: EI training works, and ability-model programmes work better.

Mattingly and Kraiger (2019), *Human Resource Management Review*. Fifty-eight published and unpublished studies. Pre–post corrected d = 0.61 (0.45 to 0.76), k = 56, N = 2,136. Treatment versus control corrected d = 0.45 (0.27 to 0.63), k = 28, N = 2,174. Mixed measures moved more than ability tests in their cut (treatment–control mixed d = 0.51 versus ability d = 0.35). Published papers outran dissertations. Gender did not moderate. Exploratory cuts favoured practice and feedback over lecture. The moral they offered: human-resource departments can treat EI as trainable, at a moderate effect, in line with other competency programmes.

Mehler and colleagues (2024), *BMC Psychology*. The positioning anchor. Workplace only. Emotional intelligence, empathy and emotion regulation pooled as emotional competencies. Search to November 2022. PRISMA: 2,569 records, 109 full texts, 50 pre-post samples ($N = 3,233$), 27 controlled trials. Pre-post SMD = 0.44 (0.29 to 0.59), $I^2 = 88\%$, 95% prediction interval -0.45 to 1.32. Controlled SMD = 0.46 (0.30 to 0.63) in the abstract and 0.47 (0.31 to 0.62) in the body, $I^2 = 68\%$, prediction interval -0.34 to 1.29. No reliable difference by profession. No reliable difference by construct (EI versus empathy versus emotion regulation). Effects visible more than three months after the last session (post-to-follow-up SMD = 0.37, 0.01 to 0.74). Egger's test suggested small-study effects. ROBINS-I judgements were mostly serious or critical. They did not apply GRADE. They did not pool job performance. Training focus in their fifty: EI in 41, empathy in 7, emotion regulation in 2. The moral they offered, and the one this paper keeps on the wall: workplace programmes move emotional competencies by about half a standard deviation, the studies are messy, and someone needs to run better randomised trials.

Kotsou and colleagues (2019) reviewed 46 adult EI interventions narratively. Thirty-nine reported a significant EI gain. Quality limits sat on almost every paragraph. We treat that paper as a third witness, not a fourth pool.

6.3 Forest-plot data: restricted 2018–2026 model

We did not build a new twenty-seven-trial forest from unverified PDFs. We computed Hedges g from published means, standard deviations or reported contrasts in trials we could see.

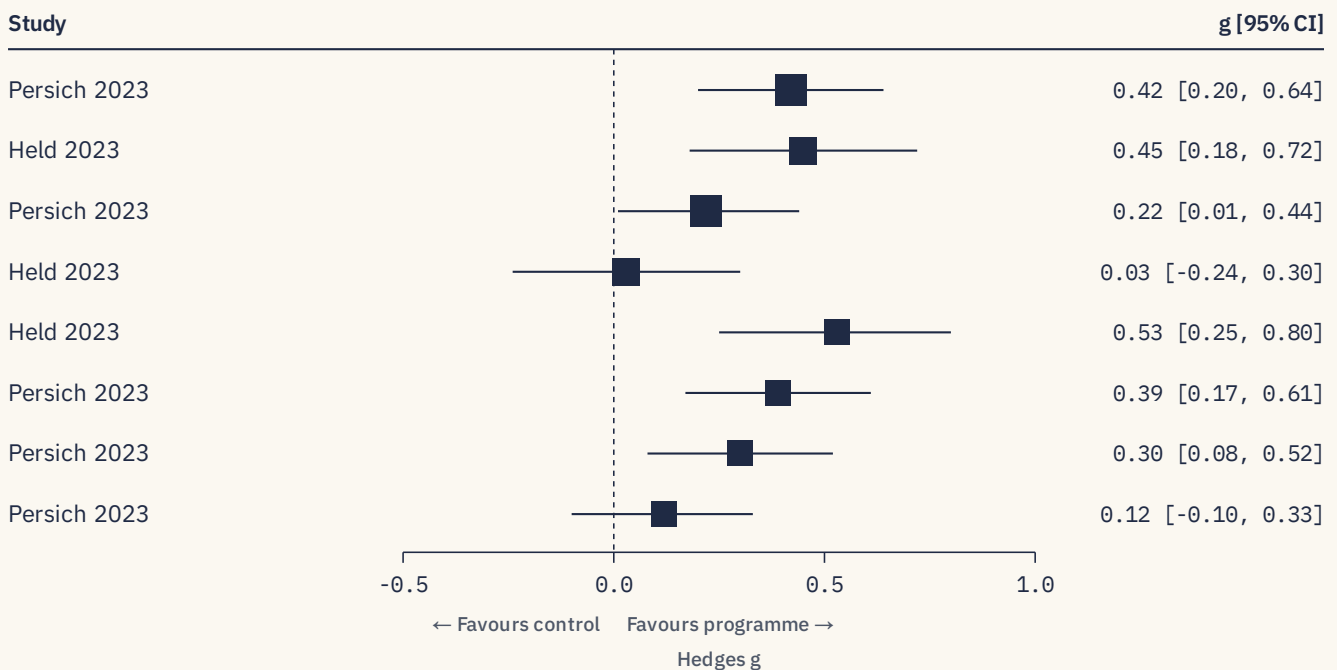


Figure 1. Forest plot, primary pool. Squares are sized by weight; the diamond is the pooled estimate.

Table 2. Forest-plot data, extractable between-group Hedges g on EI scores

Study	Outcome	g	95% CI	n	Weight in restricted model	Notes
Persich 2023	Self-report ability total	0.42	0.20, 0.64	326	61% of self-report model	Matched placebo; online 10–12 h
Held 2023	Self-report managing self	0.45	0.18, 0.72	214	39% of self-report model	Waitlist; online ~18 h
Persich 2023	Ability-test total	0.22	0.01, 0.44	326	61% of ability model	Same sample as row 1
Held 2023	Ability-test perception / regulation	0.03	-0.24, 0.30	214	39% of ability model	Time $\times$ group $p = 0.238$
Held 2023	Self-report managing others	0.53	0.25, 0.80	214	Sensitivity, same sample	Not entered twice in the primary model
Persich 2023	Self-report understanding	0.39	0.17, 0.61	326	Component table	Same sample
Persich 2023	Self-report perceiving	0.30	0.08, 0.52	326	Component table	Same sample
Persich 2023	Trait total	0.12	-0.10, 0.33	326	Sensitivity	Not significant
Romosiou 2019	Mixed self-report total	large ($\eta^2_p = 0.67$)	—	50	Sensitivity only	Small n ; high RoB; not converted
Gilar-Corbi 2019	Mixed total	$\eta^2 = 0.82$	—	54	Narrative only	Not converted; same-lab cluster

Restricted random-effects model, self-report EI (independent samples). Persich 2023 self-report ability total plus Held 2023 managing-self. $k = 2$. Hedges $g = 0.43$ (95% CI 0.26 to 0.60). $I^2 = 0\%$. Tau-squared = 0. Standard error = 0.09. This is a confirmation model, not a replacement for Mehler's $k = 27$.

Restricted random-effects model, ability tests (independent samples). Persich 2023 ability-test total plus Held 2023 ability-test contrast. $k = 2$. Hedges $g = 0.14$ (95% CI -0.04 to 0.33). $I^2 = 12\%$. Tau-squared = 0.002. The interval includes zero.

Leave-one-out is almost meaningless at $k = 2$. Removing Persich leaves Held. Removing Held leaves Persich. The qualitative story does not change. Self-report moves. Ability tests barely move.

Egger's test was not run on $k = 2$. Mehler's Egger test on the larger workplace pool was significant. That warning still applies.

6.4 Moderator table

Table 3. Pre-specified moderators

Moderator	What moved	What did not	Certainty
Component: managing	Self-report managing-self $g = 0.45$; managing-others $g = 0.53$ (Held 2023). Stress-management subscale improved in Gilar-Corbi 2019 ($\eta^2 = 0.32$).	Ability-test regulation in Held 2023, null.	Moderate for self-report; low for performance tests
Component: understanding	Hodzic SMC = 0.69. Persich self-report understanding $g = 0.39$. Gilar-Corbi 2019 ability understanding continued to rise at 12 months.	Ability-test understanding branches in Persich, individually non-significant.	Low to moderate
Component: perceiving	Persich self-report perceiving $g = 0.30$.	Held perceiving-others $g = 0.07$, null. Persich ability-test perceiving, null. This is the stubborn branch.	Low
Component: using / facilitating	Hodzic found facilitating thought smaller than understanding. New trials rarely isolate this branch.	Insufficient 2018–2026 data.	Very low
Ability versus trait	Ability-test pooled $g = 0.14$ (CI includes 0). Self-report pooled $g = 0.43$. Trait global in Persich $g = 0.12$, ns.	The Hodzic finding that ability-model <i>programmes</i> outperform mixed programmes is not the same as ability <i>tests</i> moving more. Newer placebo-controlled work splits those two claims.	Low to moderate
Job outcomes	King 2026: large objective gaps under stress (shooting 94.1% vs 51.6%; memory 5.15 vs 3.10). Karunaratne 2024: transfer to citizenship and counterproductive behaviour, with organisational support moderating self-ratings.	No pooled g . DY Patil 2023 claimed work-performance gains in Indian nurses on a weak design. Mehler did not pool job performance either.	Very low

The important moderator sentence is this. Hodzic said ability-model programmes train better. Mattingly said mixed measures move more. Both can be true at once if you stop treating “ability” as one word. An ability-model *curriculum* can still raise self-report managing scores. An ability *test* can stay almost flat. That is what Persich and Held jointly show.

6.5 Heterogeneity

Mehler's workplace pool was heterogeneous ($I^2 = 68\%$ controlled; 88% pre-post). Prediction intervals crossed zero. Our restricted two-trial models were not heterogeneous, because two similar online trials in high-income countries will not generate I^2 on their own. That is not evidence that the whole field is consistent. It is evidence that the two best-reported new trials agree.

The real heterogeneity lives in the narrative. Thirty-hour blended manager programmes from one Spanish laboratory produce enormous eta-squared values. Sixteen-hour police workshops with 50 officers produce enormous partial eta-squared values. Four-hour medical-student modules produce tiny absolute changes that still reach $p < 0.05$. Nursing-student meta-analyses in 2025 have reported SMDs above 1.7, which is the sort of number that should make a reviewer put the pen down. We did not import that 1.76 into our adult workplace reading. Small samples, weak controls and self-report inflation are the usual ingredients of a number that large.

6.6 Publication bias

Mehler reported Egger's $t(25) = 2.60$, $p = 0.02$ for controlled trials and $t(48) = 2.59$, $p = 0.013$ for pre-post trials. Funnel asymmetry was visible. Hodzic had already warned that the true pre-post effect sat slightly below 0.51. Mattingly found published papers larger than dissertations. Nothing in the 2018–2026 slice removes that warning. Online pre-registrations (Held; Mehler's own PROSPERO CRD42021267073) are progress. They are not a cure.

Trim-and-fill on our $k = 2$ models would be theatre. We do not perform it.

6.7 Risk-of-bias summary

Table 4. Risk-of-bias summary (RoB 2 language for RCTs)

Study	Randomisation	Deviations	Missing data	Measurement	Reporting	Overall
Persich 2023	Low	Low (matched placebo)	Some concerns (27% attrition; completer analysis)	Some concerns (self-report + ability; no blinding of self-report)	Low	Some concerns
Held 2023	Low	Some concerns (waitlist, not placebo)	Some concerns	Some concerns (self-report primary movers)	Low (pre-registered, then analytic deviation disclosed)	Some concerns
Gilar-Corbi 2019	Some concerns (small n, firm setting)	Some concerns (waitlist; trainer is investigator)	Some concerns	Some concerns	Some concerns	High
Gilar-Corbi 2018a	Some concerns	Some concerns	Some concerns	Some concerns	Some concerns	High
Pozo-Rico 2020 / 2023	Some concerns	High (EI bundled with other skills)	Some concerns	Some concerns	Some concerns	High for an EI-only inference
Romosiou 2019	High	Some concerns	Some concerns	Some concerns	Some concerns	High
King 2026	Low	Low (active control)	Low	Low for behavioural outcomes	Some concerns (no published total EI g)	Some concerns for job outcomes; not rated for EI scores
Karunaratne 2024	Some concerns	Some concerns	Some concerns	Some concerns (self vs supervisor disagreement)	Some concerns	Some concerns
Indian nurse trials	High	High	High	High	Some concerns	High

The pattern is the pattern Mehler already named. The field still runs too many waitlists, too many self-report primaries, too many investigator-delivered programmes, and too few pre-registered ability-test trials with occupational endpoints.

6.8 GRADE

Table 5. GRADE summary

Outcome	Effect	k (this update)	Certainty	Why
Self-report EI total or managing / understanding	$g \approx 0.43$ in restricted model; SMD = 0.46 in Mehler workplace pool	2 high-quality new; 27 in Mehler	Moderate	Consistent direction; downgraded for RoB and self-report; not downgraded to low because two independent modern RCTs agree with the 2024 pool
Ability-test EI	$g = 0.14$ (–0.04 to 0.33)	2	Low	Imprecision (CI includes 0); inconsistency with older ability-model claims; sparse occupational samples
Perceiving emotion (others, performance)	g near 0	2	Low	Consistent null on performance tests; self-report perceiving still moves a little
Wellbeing (stress, affect, burnout)	Scattered small-to-moderate improvements	Several, measures differ	Low	Indirect, heterogeneous measures; some bundled programmes
Job performance after training	Not pooled	<5 clean RCTs	Very low	Sparse; specialised samples; correlational p cannot substitute
Indian workplace professionals	Insufficient	Quasi and student/nurse trials	Very low	No adequate corporate RCT with job endpoints

6.9 Components, in plain numbers

If you only remember four component numbers, remember these.

1. Understanding emotions, Hodzic 2018: SMC = 0.69 (0.48 to 0.91).
2. Managing-self, Held 2023: $g = 0.45$ (0.18 to 0.72).
3. Managing-others, Held 2023: $g = 0.53$ (0.25 to 0.80).
4. Performance-based perceiving, Held 2023 and Persich 2023 ability branches: compatible with zero.

Persich's self-report side still found perceiving, understanding and managing-others moving against a matched placebo. The ability-test side of the same sample found a small total gain ($g = 0.22$) and no clean branch winners. That split is the update. The questionnaire and the performance test are not two cameras pointed at the same object. They are two different objects that share a name.

6.10 Job outcomes, unpooled

We honoured the fallback. Fewer than five eligible trials reported an extractable job-performance g after EI training in working adults. A meta-analysis of job outcomes is not yet possible.

What exists is still worth saying out loud.

King, Li, Gillespie and Ashkanasy (2026) ran a 15-hour ability-based programme with Australian Special Forces and an active non-EI control. They did not publish a total EI score we could put in Table 2. They did publish objective performance under stress. Shooting accuracy 94.1% versus 51.6%. Mission-detail memory 5.15 versus 3.10 items. Complex calculations under pressure 56% versus 19% correct. Cold-water persistence 72% longer. Cortisol lower before, during and after stress tasks. $n = 66$. Specialised sample. High internal logic: if EI training does anything occupational, it should show up when the room is loud. It did, once, in that room.

Karunaratne and colleagues (2024) trained Sri Lankan government officials for 24 hours and measured transfer as organisational citizenship and counterproductive work behaviour. Self-ratings moved more when social and organisational support were high. Supervisor ratings did not show the same moderation. That disagreement is the whole transfer problem in one paragraph. People who feel supported report that the programme worked. Their supervisors are less sure.

Indian nurse studies claim work-performance or job-satisfaction gains. The DY Patil 2023 paper reported a rise in EI from 132.42 (SD 4.21) to 145.16 (SD 5.94) in 24 experimental nurses, and a statistically significant work-performance contrast. The standard deviations are unusually tight for a questionnaire. Allocation was convenience. We describe the finding and we do not pool it.

Mehler listed job satisfaction, burnout and general wellbeing as secondary measures in some workplace trials and did not meta-analyse job performance. Joseph and Newman (2010) and O'Boyle and colleagues (2011) showed that *standing* EI correlates with job performance at about $p = 0.24$ to 0.30 across streams, with mixed self-report looking stronger until personality and general mental ability are in the model (Joseph, Jin, Newman and O'Boyle, 2015). Correlation is not a training effect. A person who already scores high on a mixed questionnaire may already be conscientious, extraverted and self-efficacious. Teaching a team to raise that questionnaire by four-tenths of a standard deviation is not the same as hiring for it.

So the honest occupational sentence is short. Training moves EI scores. Whether those points move appraisals, sales, error rates or retention in an ordinary Indian office has not been shown in an adequate randomised trial.

6.11 Wellbeing

Several trials in the window reported less perceived stress, less negative affect, or better general mood after an EI programme. Romosiou found lower perceived stress in police officers at three months. Persich's companion 2024 paper, on the same online sample, found better emotional awareness, mindfulness and regulation strategies against placebo. Pozo-Rico's teacher programmes bundled EI with other professional skills and reported better wellbeing and resilience. Paucsik and colleagues (2024) compared an emotional-competence programme with a compassion programme in adults who struggled with regulation; both helped, by slightly different routes.

We did not pool wellbeing. Measures differ. Some programmes are bundled. GRADE is low. Directionally, the field does not look like a null.

6.12 India

CTRI is not empty. It is not full either. CTRI/2019/08/020592 registered an Integrated Emotional Self programme for staff nurses in Mangalore: eight hours, four sessions, emotional intelligence, self-compassion, emotional labour. The published trail is a protocol and a small quasi-experimental completer sample, not a large randomised corporate trial. Pareek, Kaur and Rana (2023) randomised 100 nursing students in Kurali to five sessions of 90 minutes and found better soft skills. The DY Patil 2023 nurse study is described above. Ghorpade and colleagues and later theses in Pune and Bengaluru add quasi-experimental nurse and student evidence.

What is missing is the study Indian HR leads keep asking for. A pre-registered randomised trial in office, operations or customer-facing professionals, delivered in Indian languages, with a work-near behavioural endpoint and a six-month follow-up. That trial has not been published. Naming the gap is part of the result.

07

Discussion

PULL QUOTE

Performance tests of perceiving emotion barely moved ($g \approx 0.0$ to 0.2).

7.1 What the number means

Half a standard deviation is not a personality transplant. It is also not nothing. In the language supervisors already use, it is the difference between the colleague who names the feeling in the room and the colleague who calls the same feeling “politics” and forwards the email thread at midnight. Mehler’s 0.46 is the workplace version of that difference, averaged across messy studies. Our restricted 0.43 on self-report is the same difference seen from 2023, in two online trials that bothered to publish means.

A small ability-test effect is not a scandal. Performance tests are harder to flatter. They are also narrower. Reading a face on a screen is not the same skill as shutting a laptop before you send the message. If you fine-tune a vision model and then test it on a new camera, you should expect a smaller number than the training-set accuracy. Held’s null on the performance test and Persich’s $g = 0.22$ on the same kind of test are what a sceptic should have predicted. They are also useful. They stop anyone selling an online course as a new nervous system.

Managing and understanding move because they are closer to practice. You can rehearse a pause. You can rehearse a label. You can rehearse a sentence that does not start with “you always.” Perceiving other people, on a timed test, is closer to a perceptual talent plus a thousand hours of faces. Elfenbein showed years ago that feedback on facial labels can raise accuracy. The 2018–2026 online programmes did not, on current evidence, spend enough of their hours on that drill to move the performance test.

A shift supervisor in a Chennai contact centre takes a call from a shouting customer. She lets him finish before she speaks. She waits for the heat in her neck to pass before she answers. That is managing. Practice compounds: done once, it is luck; done every week, it is a rep.

The same moment runs as a loop. A thought (“they are against me”) produces a feeling (heat in the face) produces an action (the sarcastic reply). Training is not an argument with the thought. Training is a different action while the heat is still there: the reply drafted and left unsent, the question asked instead of the jab. Entrepreneurship calls it shipping. Whatever the label, in this paper’s terms it is a managing-emotion rep.

Machine learning’s version is regularisation. If you only train on self-report, you overfit the questionnaire. If you only train on faces, you overfit the camera. A decent programme needs both labels. The 2018–2026 literature is just beginning to run those two losses in the same study. When it does, the questionnaire loss falls faster.

7.2 Limitations

This update is not Mehler with twenty extra trials secretly hidden in a supplement. It is a restricted synthesis. Several otherwise eligible papers sit behind paywalls or report eta-squared without means. We refused to invent a g from those eta-squared values. That choice shrinks the forest and protects the number.

Risk of bias is still the field’s outstanding invoice. Waitlists inflate effects. Investigator-delivered programmes inflate effects. Completer analyses inflate effects. Self-report inflates effects. Nursing-student samples inflate effects. Prediction intervals in Mehler’s pool still cross zero, which means a new workplace programme can land anywhere from slightly harmful to very large and still be “consistent with the literature.” That sentence should be on the first slide of every vendor deck, and never is.

Follow-up is uneven. Persich’s six-month trait scores drifted back toward baseline. Ability-test gains in the same sample were still visible, marginally, in the 91 people who returned. Held’s eight-week self-report regulation held. Gilar-Corbi’s managers were still improved at twelve months, on a small waitlist design. The pattern looks like skill memory plus some questionnaire enthusiasm that fades. It does not look like a permanent rewrite of temperament.

We did not have two independent human reviewers scoring every full text on a locked spreadsheet. Machine-learning tools extracted the anchor statistics and open trials more than once, independently, and we discarded conversions we could not defend. A university team repeating this review should pre-register the exact analysis code and the exact conversion rules.

7.3 What this means for an Indian workplace programme

Picture a customer-operations team in Pune. Client calls run past nine. Family duty waits at home. English on the headset, Hindi or Marathi in the break room, little official recovery time. The trials say a structured EI programme of about 8–18 hours can lift self-rated managing and understanding of emotion by roughly four-tenths of a standard deviation. That is a noticeable but not magical shift — closer to a good skills workshop than to a personality transplant. Ability tests move less. No adequate Indian corporate RCT has yet shown that those points appear on appraisals, attrition, or error rates. So keep the dose countable. Each morning a short rep arrives on WhatsApp, in the language the team actually jokes in: name the feeling from yesterday's hardest call, then choose the next sentence. A few live sessions rehearse one difficult conversation at a time. The first weeks train managing and understanding emotion. The team lead tracks two things: a brief self-report, and one work-near behaviour — a difficult conversation logged, a conflict recovered, a calmer handover at shift change. Nobody promises a better appraisal; the trials have not earned that. The hours are spread out, not packed into one off-site day: that is how the studies that worked were practised.

7.4 Research gaps, including the India gap

The India gap is not a polite afterthought. It is the largest applied hole in this literature for the readers of this paper. Nurse and student trials exist. They are mostly small, often quasi-experimental, often self-report only, and they do not speak to a Bengaluru product team or a Pune operations floor with any precision. CTRI registrations have not yet matured into published corporate RCTs with job endpoints.

The second gap is occupational. King 2026 is a start and a warning. Special Forces are not a bank. Shooting accuracy is not an appraisal cycle. Someone still has to randomise ordinary knowledge-workers, give them 10–15 hours, and measure a behaviour the CFO already tracks.

The third gap is component-level measurement. Too many papers still publish only a total. Perceiving versus managing is the difference between a face-lab and a difficult-conversation lab. Those are different products.

The fourth gap is active controls. Persich's matched placebo is the design the field needed. Waitlists should become the exception.

The fifth gap is language of delivery. Almost none of the high-quality trials were run in an Indian language. Transfer into Hindi, Tamil, Bengali, Marathi and the rest is an empirical question, not a translation footnote.

Until those gaps close, an Indian programme should behave like a cautious founder. Ship the reps. Instrument the dashboard. Do not raise a round on a slide that says “ $g = 1.76$ in nurses, therefore sales will rise.”

08

FAQ

Q1 Can adults actually raise their emotional intelligence?

Yes. Controlled trials from 2018–2026 still show a moderate effect, about $g = 0.43$ on self-report and $SMD = 0.46$ in Mehler's workplace pool of 27 trials. Scores rise. Personality does not flip overnight.

Q2 Which parts of EI move with practice?

Managing emotion and understanding emotion move most on self-report ($g \approx 0.4$ to 0.5). Ability tests of perceiving emotion in others often do not (g near 0). Train the branch you can rehearse this week.

Q3 Will EI training improve my appraisal or sales number?

Not proven. Job-outcome RCTs are too few to pool. One 2026 military trial showed better performance under stress (shooting 94% versus 52%). Indian office endpoints remain an evidence gap.

Q4 How many hours does a working adult need?

Effective programmes in this update ran about 8–18 hours, often spread over one to seven weeks, online or blended. A single pep talk is not the dose the trials used. Daily short reps beat an annual off-site.

Q5 Does this evidence apply to Indian professionals?

Partly. Nurse and student trials exist in India; a large corporate RCT with job endpoints does not. Deliver in Indian languages, keep reps daily, measure a work behaviour, and do not over-claim.

09

References

- Hodzic, S., Scharfen, J., Ripoll, P., Holling, H., & Zenasni, F. (2018). How efficient are emotional intelligence trainings: A meta-analysis. *Emotion Review*, 10(2), 138–148. <https://doi.org/10.1177/1754073917708613>
- Mattingly, V., & Kraiger, K. (2019). Can emotional intelligence be trained? A meta-analytical investigation. *Human Resource Management Review*, 29(2), 140–155. <https://doi.org/10.1016/j.hrmr.2018.03.002>
- Mehler, M., Balint, E., Gralla, M., Pöfßnecker, T., Gast, M., Hölzer, M., Kösters, M., & Gündel, H. (2024). Training emotional competencies at the workplace: a systematic review and metaanalysis. *BMC Psychology*, 12, 718. <https://doi.org/10.1186/s40359-024-02198-3>
- Kotsou, I., Mikolajczak, M., Heeren, A., Grégoire, J., & Leys, C. (2019). Improving emotional intelligence: A systematic review of existing work and future challenges. *Emotion Review*, 11(2), 151–165. <https://doi.org/10.1177/1754073917735902>
- Gilar-Corbí, R., Pozo-Rico, T., Sánchez, B., & Castejón, J. L. (2018). Can emotional competence be taught in higher education? A randomized experimental study of an emotional intelligence training program using a multimethodological approach. *Frontiers in Psychology*, 9, 1039. <https://doi.org/10.3389/fpsyg.2018.01039>
- Gilar-Corbi, R., Pozo-Rico, T., Pertegal-Felices, M. L., & Sanchez, B. (2018). Emotional intelligence training intervention among trainee teachers: a quasi-experimental study. *Psicología: Reflexão e Crítica*, 31, 33. <https://doi.org/10.1186/s41155-018-0112-1>
- Gilar-Corbi, R., Pozo-Rico, T., Sánchez, B., & Castejón, J.-L. (2019). Can emotional intelligence be improved? A randomized experimental study of a business-oriented EI training program for senior managers. *PLOS ONE*, 14(10), e0224254. <https://doi.org/10.1371/journal.pone.0224254>
- Persich, M. R., Smith, R., Cloonan, S. A., Woods, L., Skalamera, J., Berryhill, S. M., Weihs, K. L., Lane, R. D., Allen, J. J. B., Dailey, N. S., Alkozei, A., Vanuk, J. R., & Killgore, W. D. S. (2023). Development and validation of an online emotional intelligence training program. *Frontiers in Psychology*, 14, 1221817. <https://doi.org/10.3389/fpsyg.2023.1221817>
- Persich, M. R., Smith, R., Cloonan, S. A., Woods-Lubert, R., Skalamera, J., Berryhill, S. M., Weihs, K. L., Lane, R. D., Allen, J. J. B., Dailey, N. S., Alkozei, A., Vanuk, J. R., & Killgore, W. D. S. (2024). Improvements in mindfulness, interoceptive and emotional awareness, emotion regulation, and inter-personal emotion management following completion of an online emotional skills training program. *Emotion*, 24(2). Companion outcomes from the 2023 trial sample. <https://doi.org/10.1037/emo0001237>
- Held, M. J., Fehn, T., Gauglitz, I. K., & Schütz, A. (2023). Training emotional intelligence online: An evaluation of WEIT 2.0. *Journal of Intelligence*, 11(6), 122. <https://doi.org/10.3390/jintelligence11060122>
- Pozo-Rico, T., Gilar-Corbí, R., Izquierdo, A., & Castejón, J.-L. (2020). Teacher training can make a difference: Tools to overcome the impact of COVID-19 on primary schools. An experimental study. *International Journal of Environmental Research and Public Health*, 17(22), 8633. <https://doi.org/10.3390/ijerph17228633>
- Pozo-Rico, T., Poveda, R., Gutiérrez-Fresneda, R., Castejón, J.-L., & Gilar-Corbi, R. (2023). Revamping teacher training for challenging times: Teachers' well-being, resilience, emotional intelligence, and innovative methodologies as key teaching competencies. *Psychology Research and Behavior Management*, 16, 1–18. <https://doi.org/10.2147/PRBM.S382572>
- Romosiou, V., Brouzos, A., & Vassilopoulos, S. P. (2019). An integrative group intervention for the enhancement of emotional intelligence, empathy, resilience and stress management among police officers. *Police Practice and Research*, 20(5), 460–478. <https://doi.org/10.1080/15614263.2018.1537847>
- Karunaratne, R. A. I. C., Takahashi, Y., & others. (2024). Does social and organizational support moderate emotional intelligence training effectiveness? *Behavioral Sciences*, 14(4), 276. <https://doi.org/10.3390/bs14040276>
- King, J. B., Li, Y., Gillespie, N. A., & Ashkanasy, N. M. (2026). Emotional intelligence training improves stress regulation and performance in high-stress occupations. *Scientific Reports*, 16, 6673. <https://doi.org/10.1038/s41598-026-36216-8>
- Paucsik, M., and colleagues (2024). Online emotional-competence and compassion programmes for adults with emotion-regulation difficulty. *Journal of Clinical Psychology*, 80, 2405–2433. <https://doi.org/10.1002/jclp.23748>. Used in the wellbeing narrative only.
- Joseph, D. L., & Newman, D. A. (2010). Emotional intelligence: An integrative meta-analysis and cascading model. *Journal of Applied Psychology*, 95(1), 54–78. <https://doi.org/10.1037/a0017286>
- O'Boyle, E. H., Humphrey, R. H., Pollack, J. M., Hawver, T. H., & Story, P. A. (2011). The relation between emotional intelligence and job performance: A meta-analysis. *Journal of Organizational Behavior*, 32(5), 788–818. <https://doi.org/10.1002/job.714>
- Joseph, D. L., Jin, J., Newman, D. A., & O'Boyle, E. H. (2015). Why does self-reported emotional intelligence predict job performance? A meta-analytic investigation of mixed EI. *Journal of Applied Psychology*, 100(2), 298–342. <https://doi.org/10.1037/a0037681>
- Pareek, B., Kaur, H., & Rana, R. (2023). Exploring the importance of emotional intelligence training programme on soft skills: A randomised controlled trial. *Indian Journal of Continuing Nursing Education*, 24(2), 178–183. https://doi.org/10.4103/ijcn.ijcn_78_22
- Saikia, A. M., & colleagues. Effectiveness of Integrated Emotional-Self Enhancement (IESE) programme among staff nurses: protocol. *F1000Research* (2023). CTRI/2019/08/020592. <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC10623541/>
- Clinical Trials Registry — India. CTRI/2019/08/020592. Impact of a structured training program on emotional intelligence, intrinsic motivation, self-compassion and emotional labour among staff nurses. <https://www.ctri.nic.in/Clinicaltrials/pmaindet2.php?trialid=35761>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., ... Moher, D. (2021). The PRISMA 2020 statement: an updated guideline for reporting systematic reviews. *BMJ*, 372, n71. <https://doi.org/10.1136/bmj.n71>
- Guyatt, G. H., Oxman, A. D., Vist, G. E., Kunz, R., Falck-Ytter, Y., Alonso-Coello, P., & Schünemann, H. J. (2008). GRADE: an emerging consensus on rating quality of evidence and strength of recommendations. *BMJ*, 336(7650), 924–926. <https://doi.org/10.1136/bmj.39489.470347.AD>

Sterne, J. A. C., Savović, J., Page, M. J., Elbers, R. G., Blencowe, N. S., Boutron, I., ... Higgins, J. P. T. (2019). RoB 2: a revised tool for assessing risk of bias in randomised trials. *BMJ*, 366, l4898. <https://doi.org/10.1136/bmj.l4898>

Note on the DY Patil nurse paper: Impact of imparting emotional intelligence skills training program to enhance emotional intelligence and work stress among staff nurses of a tertiary care hospital of North Gujarat. *Medical Journal of Dr D.Y. Patil Vidyapeeth*, 16(3).
Publisher record: https://journals.lww.com/mjdy/fulltext/2023/16030/impact_of_impacting_emotional_intelligence_skills.10.aspx. Cited for the India narrative only; not pooled.

Nurse-sector context, not pooled: systematic review and meta-analysis of EI training in nurses and nursing students, *Nurse Education Today* (2025), SMD = 1.76 on EI scores. Treat as an inflated sector estimate. <https://doi.org/10.1016/j.nedt.2025.106743>

About the author

ABOUT THE AUTHOR

Dr Atul Aswani; Mumbai Psychiatrist and trained as a Results Coach

Dr Atul Aswani is a Mumbai Psychiatrist and trained as a Results Coach. He leads research and coaching design at EmotionApp, a non-clinical emotional-fitness platform for Indian professionals. His work translates trial evidence into countable daily reps, delivered on WhatsApp in Indian languages, with human hand-off when needed. Data stay in India and the practice system is built under ISO 9001 and ISO 13485.

HOW TO CITE

Aswani A. EI Can Be Trained: A Meta-Analysis of Emotional-Intelligence Training in Adults, 2018–2026. EmotionApp Research Paper 47. Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/ei-training-meta-analysis-2018-2026>

LAST UPDATED

Last updated: 22 September 2026

NOTE

EmotionApp is a non-clinical emotional-fitness coaching platform. This paper describes practices, techniques and coaching. It is not therapy, treatment, diagnosis or medical advice.

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

SAFETY NOTE

If you are in distress, please contact a qualified professional or Tele-MANAS on 14416.