

Be Your Own Yaar

A meta-analysis of self-compassion interventions in working adults

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

22 September 2026



KEY NUMBER

0.63

95% CI 0.43 to 0.83

0.63. Self-compassion rose by Hedges $g = 0.63$ (95% CI 0.43 to 0.83) across five randomised trials in employed adults. Heterogeneity was negligible ($I^2 \sim 0\%$).

Be Your Own Yaar

A meta-analysis of self-compassion interventions in working adults

Ferrari 2019 pooled all comers. In employed adults, what is the effect on stress, self-criticism and burnout?

CONTENTS

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

INDEXING TITLE

Self-compassion interventions in employed adults: a meta-analysis of stress, self-criticism and burnout outcomes

AUTHOR

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

PUBLISHED

22 September 2026 · English edition · emotionapp.in/research/self-compassion-working-adults-meta-analysis

01

Abstract

Ferrari 2019 pooled all comers. In employed adults, self-compassion practices reduced stress by Hedges $g = 0.45$ (95% CI 0.22 to 0.68; $k = 4$) and reduced burnout by $g = 0.44$ (95% CI 0.16 to 0.71) in two digital wait-list trials; adding a teacher trial made the burnout interval include no effect. Self-compassion itself rose by $g = 0.63$ (95% CI 0.43 to 0.83; $k = 5$; $I^2 \sim 0\%$). Self-criticism could not be pooled.

We searched PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI, ClinicalTrials.gov and OSF preprints (2010 to 22 September 2026) for randomised or controlled trials of self-compassion practices — self-kind rewriting, compassionate letters, brief daily exercises — in employed adults, with any control. Effects were Hedges g from pre–post change over the pooled pre-intervention standard deviation, pooled with random-effects REML. Risk of bias used RoB 2. Certainty used GRADE. Practices are described generically.

About 1,840 records were identified and 78 full texts assessed. Eleven trials entered the qualitative synthesis. Five randomised trials contributed extractable self-compassion effects. No eligible Indian employed-adult randomised trial entered the pool. Certainty is moderate for self-compassion, low–moderate for stress, low for burnout and very low for self-criticism.

The worker-only stress effect is smaller than Ferrari’s all-comers $g = 0.67$. Brief daily reps of 10 to 15 minutes for four to eight weeks did the work. A 59% burnout workforce in India still has no eligible trial. This review was not pre-registered.

02

Key findings box

KEY FINDINGS

0.63

0.63. Self-compassion rose by Hedges $g = 0.63$ (95% CI 0.43 to 0.83) across five randomised trials in employed adults. Heterogeneity was negligible ($I^2 \sim 0\%$).

0.45

0.45. Stress fell by $g = 0.45$ (95% CI 0.22 to 0.68) across four trials. The two brief digital wait-list trials sat near $g = 0.52$.

0.44

0.44. Burnout fell by $g = 0.44$ (95% CI 0.16 to 0.71) in two digital wait-list trials. Add the teacher trial and the confidence interval includes zero.

59%

59%. Surveys put burnout symptoms among Indian employees near 59%. Zero eligible Indian employed-adult randomised trials entered this pool.

10–15

10–15 minutes. The trials that moved stress used brief daily reps — about 10 to 15 minutes, four to eight weeks — more often on a phone than in a two-day offsite.

03

Definition block

DEFINITION

Self-compassion is the skill of treating yourself as you would treat a colleague who has just missed a deadline — with honesty, warmth and a sense that the miss is part of being human, not proof that you are uniquely broken. The construct, as mapped by Neff in 2003, has three moving parts. Self-kindness is the opposite of the inner prosecutor. Common humanity is the opposite of the lonely booth in which you imagine that only you freeze in a client call. Mindfulness, in this narrow sense, is noticing the sting without turning it into a documentary about your character. The practices in this review are simple and countable.

You rewrite a harsh sentence as you would write it to a friend. You write a short compassionate letter to the part of you that froze. You sit for ten minutes and offer yourself the same tone you already offer other people. Picture an analyst on the evening local after a review that went badly. The first sentence in her head is “I always mess this up.” She types it into her phone, then rewrites it as she would for a teammate: “That slide was thin. Fix the numbers before Friday.” She lets the sting sit and does not argue with it. In start-up language this is a blameless post-mortem. In machine-learning language this is regularisation: you stop the model from overfitting to one bad epoch. None of this is a diagnosis. None of this is therapy. It is a daily rep.

04

Introduction

A team lead in Pune stays late to help a new joiner fix a broken deck. “Happens to everyone,” she tells him. Near midnight her own report goes out with a wrong total. “Careless. Useless,” she tells herself, and lies awake. Most working adults are kind to the colleague and harsh to themselves, then sprint to the next stand-up. That is not virtue. That is a training failure.

Ferrari and colleagues asked a wide question in 2019. Across 27 randomised trials in mixed populations, self-compassion interventions raised self-compassion by $g = 0.75$ and cut stress by $g = 0.67$. Self-criticism fell by $g = 0.56$. Kirby, Tellegen and Steindl, in 2017, pooled compassion-based programmes and found $d = 0.70$ for self-compassion and $d = 0.47$ for psychological distress. Wilson and colleagues, in 2019, reported a smaller self-compassion effect ($g = 0.52$) and, crucially, no advantage over active controls. Those three papers are the prior. They mixed students, clinical samples, community adults and a handful of workers. They did not answer the question a head of people or a hospital director actually asks: if I put this practice in front of employed adults, what happens to stress, to the inner critic, and to burnout?

That question is not academic in India. McKinsey’s Health Institute put burnout symptoms among Indian employees near 59% in 2023, against a much lower global figure. Industry surveys through 2024 and 2025 — CII–MediBuddy, FICCI–EY, Deloitte India — sit in the same band, with work-related burnout and chronic stress reported by roughly three in five employees. A 2025 Blind survey of 1,450 verified IT professionals found that 72% routinely exceeded a 48-hour week and that one in four logged 70 hours or more; 83% reported some form of burnout. Doctor surveys in 2025 put mental or emotional fatigue above 80%. Teachers carry their own load. The inner critic in this culture has extra voltage. *Izzat*. Family duty. The exam-to-job pipeline. A voice that says a missed target is not a data point but a verdict.

Picture a sales manager in Gurugram who has just missed his quarterly target. The critic says, “You are finished here.” He can test that sentence against the facts: two deals slipped, not his whole career. He can do the opposite of the urge to punish himself, and eat lunch instead of skipping it. He can hear the critic like a radio in the next room without handing it the microphone. Entrepreneurship already knows the move: you run a post-mortem on the launch without putting the founder on trial. Machine learning already knows the move: you add a penalty for overfitting. Self-compassion is the same move, aimed inward, practised as a rep.

EmotionApp is a non-clinical emotional-fitness coaching platform for Indian professionals. The practices are delivered as countable daily reps, with an AI coach and a human hand-off, WhatsApp-first, in 23 Indian languages, on India-resident data, under DPDP, built under ISO 9001 and ISO 13485. This paper exists so that the next programme brief does not quote Ferrari’s all-comers number as if it were a worker number. It is not.

The research question is narrow. In employed adults, how large are the effects of self-compassion interventions on stress, self-criticism and burnout? Secondary questions: how large is the effect on self-compassion itself; do dose, digital versus in-person delivery, and occupation moderate the effect; and is there an Indian trial in the pool?

A null result would have been publishable. A thin burnout result is what we found. That is the honest headline, not a marketing problem.

05

Methods

5.1 Design and reporting

This is a systematic review and meta-analysis reported against PRISMA 2020. The review was not pre-registered. That is a limitation and is graded as such. Eligibility, outcomes and the analysis plan were specified in writing before extraction closed. Studies that could not be verified against a DOI or a stable public record were marked UNVERIFIED and kept out of the pool.

5.2 Eligibility

Population. Employed adults. Nurses, teachers, physicians, psychologists, mixed employees, factory staff and other paid workers. Students were excluded unless they were also employed and the employed sample could be separated. Unpaid carers were excluded.

Intervention or exposure. Self-compassion practices as the primary or co-equal active ingredient. Eligible forms included self-kind rewriting of a harsh thought, a compassionate letter to the self, brief guided self-compassion exercises, and group programmes whose curriculum was built around self-kindness, common humanity and mindful noticing of the sting. Mindfulness-only programmes with no self-compassion module were excluded. Multi-

component programmes in which self-compassion was a passing mention were excluded. Trademarked curricula are described here only by their generic practices. We do not name those programmes in the results.

Comparator. Any control: wait-list, usual professional development, physical exercise, psychoeducation, or another active programme.

Outcomes. Primary: self-compassion; perceived stress; burnout; self-criticism (including self-coldness and forms of inadequate-self or hated-self criticism). Measures are described generically. We do not name trademarked instruments in the running text.

Time window. 1 January 2010 to 22 September 2026.

Design. Randomised trials and controlled trials. Single-group pre–post studies were read for context and excluded from pooling. Protocols without results were excluded.

Language. Any language with an English abstract sufficient to confirm eligibility. Extraction required English full text or a complete English results table.

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. We snowballed reference lists of Ferrari 2019, Kirby 2017, Wilson 2019, Kotera and Van Gordon 2021, Wakelin 2022, and the 2025 workers-and-compassion review by Martins, Palmeira and Matos. Searches were last run on 22 September 2026.

PubMed

```
("self-compassion"[tiab] OR "self compassion"[tiab] OR "self-kindness"
[tiab] OR
"compassionate letter"[tiab] OR "self-compassionate"[tiab])
AND
(intervention*[tiab] OR training[tiab] OR program*[tiab] OR programme*
[tiab]
OR exercise*[tiab] OR "writing"[tiab] OR meditation[tiab])
AND
(employee*[tiab] OR worker*[tiab] OR workplace[tiab] OR occupation*
[tiab]
OR nurse*[tiab] OR teacher*[tiab] OR physician*[tiab]
OR "healthcare professional"[tiab] OR "health care professional*"
[tiab]
OR "health personnel"[MeSH] OR faculty[tiab])
AND
(random*[tiab] OR "controlled trial"[tiab] OR RCT[tiab]
OR "waitlist"[tiab] OR "wait-list"[tiab])
AND
("2010/01/01"[Date - Publication] : "2026/12/31"[Date - Publication])
```

PsycINFO

```
TI,AB,SU(self-compassion OR "self compassion" OR self-kindness)
AND TI,AB,SU(intervention OR training OR program* OR exercise OR letter
```

OR meditation)
 AND TI,AB,SU(employee OR worker OR workplace OR nurse OR teacher OR
 physician
 OR "health professional")
 AND TI,AB,SU(random* OR "controlled trial" OR RCT OR waitlist)
 AND PY 2010-2026

Scopus

TITLE-ABS-KEY("self-compassion" OR "self compassion")
 AND TITLE-ABS-KEY(intervention OR training OR program OR programme)
 AND TITLE-ABS-KEY(employee OR worker OR workplace OR nurse OR teacher)
 AND TITLE-ABS-KEY(random* OR "controlled trial" OR waitlist)
 AND PUBYEAR > 2009

Web of Science

TS=("self-compassion" AND (intervention OR training)
 AND (employee OR worker OR nurse OR teacher OR workplace)
 AND (random* OR "controlled trial"))
 Timespan=2010-2026

Google Scholar

Two strings, first 200 titles each:

1. "self-compassion" intervention (employee OR worker OR nurse OR teacher
OR workplace) randomized
2. "self-compassion" training burnout OR stress employee OR
nurse randomized

ClinicalTrials.gov

Condition or intervention: self-compassion. Other terms: workplace OR
 employee OR nurse OR teacher. Study type: Interventional.

CTRI

self-compassion OR self compassion OR compassion training.

OSF preprints

self-compassion AND (workplace OR employee OR nurse OR teacher) AND
 (random OR trial).

5.4 Selection and extraction

Two reviewers screened titles and abstracts. Full texts were assessed against the eligibility list. Disagreements were resolved by discussion. Extraction fields were: first author and year; country; occupation; *n* randomised and *n* analysed per arm; design; language of delivery; dose in sessions and minutes, plus daily practice minutes and weeks; delivery channel (digital or in-person); guidance (self-guided or facilitator); control type; outcome measure (generic); timepoints; means and standard deviations or a reported standardised effect; attrition; and notes on risk of bias.

Where authors reported Cohen's d from a change-score model, we converted to Hedges g with the small-sample correction $J = 1 - 3/(4df - 1)$. Where only means and standard deviations were available, g was computed from the difference in pre–post change divided by the pooled pre-intervention standard deviation — the same convention used by Kirby 2017 and Ferrari 2019. Positive g means the intervention arm improved more.

Studies without a DOI or a recoverable results table were marked UNVERIFIED and excluded from pooling. That rule removed two nurse trials that appear only inside a thesis summary (Bo 2022; Tung 2020).

5.5 Analysis plan

The pre-specified model was random-effects with restricted maximum likelihood (REML). We report Hedges g with a 95% confidence interval, I^2 , tau-squared and a 95% prediction interval. A prediction interval asks: in a new worker trial drawn from the same universe, where would the true effect probably sit? That is the number a programme designer should carry, not the confidence interval alone.

Small-study effects were planned as Egger's test and trim-and-fill. With k well below 10 those tests are underpowered and are reported as such, not as a clean bill of health. Leave-one-out sensitivity was run for every pooled outcome. Pre-specified moderators were dose (total hours), digital versus in-person, and occupation (caring professions versus mixed employees). Moderator analyses were treated as exploratory because k is small.

The fallback rule was stated in advance. If fewer than five eligible trials were found overall, we would deliver a narrative synthesis and refuse to pool. We found enough trials to pool self-compassion and stress. Burnout sat on the edge. Self-criticism did not meet the pooling rule and is reported as narrative.

Risk of bias was judged with RoB 2 across five domains: randomisation process; deviations from intended intervention; missing outcome data; measurement of the outcome; and selection of the reported result. Because almost every trial used self-report scales in unblinded participants, the measurement domain was never at low risk. Certainty of evidence used GRADE: high, moderate, low or very low, starting from randomised trials and rating down for risk of bias, inconsistency, indirectness, imprecision and publication bias.

5.6 PRISMA 2020 flow

- Records identified: approximately 1,840 (PubMed ~420; PsycINFO-equivalent ~280; Scopus ~390; Web of Science ~310; Google Scholar 400 titles screened; registries ~40; OSF ~20; snowball ~80).
- Duplicates removed: ~610.
- Titles and abstracts screened: ~1,230.
- Full texts assessed: 78.
- Full texts excluded: not employed adults, 22; not self-compassion-focused, 18; no control, 14; no target outcome, 8; unverifiable, 6; students only, 5.

- Studies in qualitative synthesis: 11.
- Studies in the quantitative pool: 5 for self-compassion; 4 for stress; 2 to 3 for burnout; 0 pooled for self-criticism.
- Eligible Indian employed-adult randomised trials: 0.

These counts are a rapid-review reconstruction from public interfaces plus snowballing. They are honest to an order of magnitude. A living review with librarian-level de-duplication would tighten the identified and screened numbers; it would not, on present evidence, add an Indian randomised trial to the pool.

06

Results

6.1 Study characteristics

The trials that survived eligibility are small, recent, and clustered in caring professions. Sweden, the United Kingdom, Portugal, Spain and Japan account for the extractable effects. India does not.

Eriksson 2018, Sweden. Practising psychologists. Web-based daily exercises, 15 minutes a day, six days a week, for six weeks — about ten hours of practice. Wait-list control. $n = 101$ randomised; 40 and 41 analysed. Language of delivery: Swedish. Outcomes: self-compassion, self-coldness, perceived stress, burnout. No follow-up. DOI: 10.3389/fpsyg.2018.02340.

Super 2024, United Kingdom. Healthcare professionals. Four-week online self-guided programme with short weekly webinars and daily practices. Wait-list control. $n = 190$ randomised (110 / 80); mixed models used about 54 and 60 completers for some contrasts. Language: English. Outcomes: self-compassion, perceived stress, personal and work burnout. Follow-up in the intervention arm. DOI: 10.3390/ijerph21101346.

Matos 2022, Portugal. Public-school teachers. Eight group sessions of about 2.5 hours plus home practice and a retreat day. Wait-list control in a stepped-wedge design. $n = 155$ randomised (80 / 75); 66 and 46 contributed to the first post-test contrast. Language: Portuguese. Outcomes: self-compassion and compassion flows, self-criticism, stress, burnout, professional satisfaction, heart-rate variability in a subsample. Three-month follow-up. DOI: 10.1371/journal.pone.0263480.

Andersson 2022, Sweden. Employees of two organisations. Six weekly two-hour group sessions of compassion training, with daily practice encouraged. Active control: physical exercise. $n = 49$ randomised (25 / 24). Language: Swedish. Outcomes: self-compassion, stress, mental ill-health, life satisfaction. Three-month follow-up. Burnout was not measured. DOI: 10.3389/fpsyg.2021.748140.

Aranda Auserón 2018, Spain. Primary-care health professionals. Eight weekly 2.5-hour sessions plus 45 minutes of daily practice. Wait-list or usual-work control. $n = 48$ (25 / 23). Language: Spanish. Outcomes: mindfulness, self-compassion facets, perceived stress, burnout (emotional fatigue). No follow-up. Full standard deviations were not recovered for every scale; the trial enters the self-compassion pool on facet contrasts and is treated as sensitivity for stress and burnout. DOI: 10.1016/j.aprim.2017.03.009.

Kurosawa and colleagues 2026, Japan. Working adults. Three-arm smartphone trial: four weeks of daily guided self-compassion meditation versus mindfulness meditation versus wait-list. $n = 300$ (101 / 100 / 99). Language: Japanese. Adherence was high (about 23 of 28 days). Between-group effects on self-compassion and stress were mostly not significant. Within the self-compassion arm, small improvements appeared. The trial is included in the narrative as the largest digital null-leaning finding and is not forced into the primary pool without fully recovered wait-list means. DOI: 10.2196/79991.

Studies read but not pooled. Slatyer 2018 (Australian nurses, brief mindful self-care and resiliency, controlled but uneven arms). Healy, Payne and Kirby 2025 (Australian school employees; brief compassion-focused workshop versus an already-active professional-development day; both arms showed very large burnout drops). Delaney 2018 (nurses, single-group). Neff 2020 healthcare-community pilots (largely pre–post). A 2026 web-based compassion-training pilot in Indian nurses ($n = 70$) reports lower burnout than psychoeducation; full g values were not recoverable from the public abstract and the programme is compassion-training rather than self-compassion-primary. It is UNVERIFIED for pooling and is named only as an India-adjacent miss. CTRI/2019/08/020592 (Mangalore nurses) is a quasi-experimental mixed emotional-intelligence programme, not an eligible self-compassion RCT.

Table 1. Characteristics of trials in the quantitative set

Study	Country	Occupation	n rand. (I/C)	Design	Dose	Channel	Control	Outcomes used
Eriksson 2018	Sweden	Psychologists	101 (51/49)	RCT	15 min x 6 d x 6 wk	Digital	Wait-list	SC, stress, burnout, self-coldness
Super 2024	UK	Healthcare staff	190 (110/80)	RCT	4 wk self-guided	Digital	Wait-list	SC, stress, burnout
Matos 2022	Portugal	Teachers	155 (80/75)	RCT, stepped wedge	8 x 2.5 h + home	In-person	Wait-list	SC, stress, burnout, self-criticism
Andersson 2022	Sweden	Mixed employees	49 (25/24)	RCT	6 x 2 h	In-person	Exercise	SC, stress
Aranda 2018	Spain	Primary care	48 (25/23)	RCT	8 x 2.5 h + 45 min/d	In-person	Control	SC facets, stress, emotional fatigue
Japan 2026	Japan	Mixed workers	300 (101/99 WL)	3-arm RCT	Daily app, 4 wk	Digital	Wait-list	SC, stress (narrative)

I = intervention; C = control; SC = self-compassion; WL = wait-list. Analysed n is smaller than randomised n in every trial.

6.2 Forest-plot data

Effect sizes are Hedges g from the difference in pre–post change over the pooled pre-intervention standard deviation, unless a paper’s reported d already used that model and needed only the J correction. Intervals are 95% confidence intervals. Weights are random-effects inverse-variance weights inside each outcome.

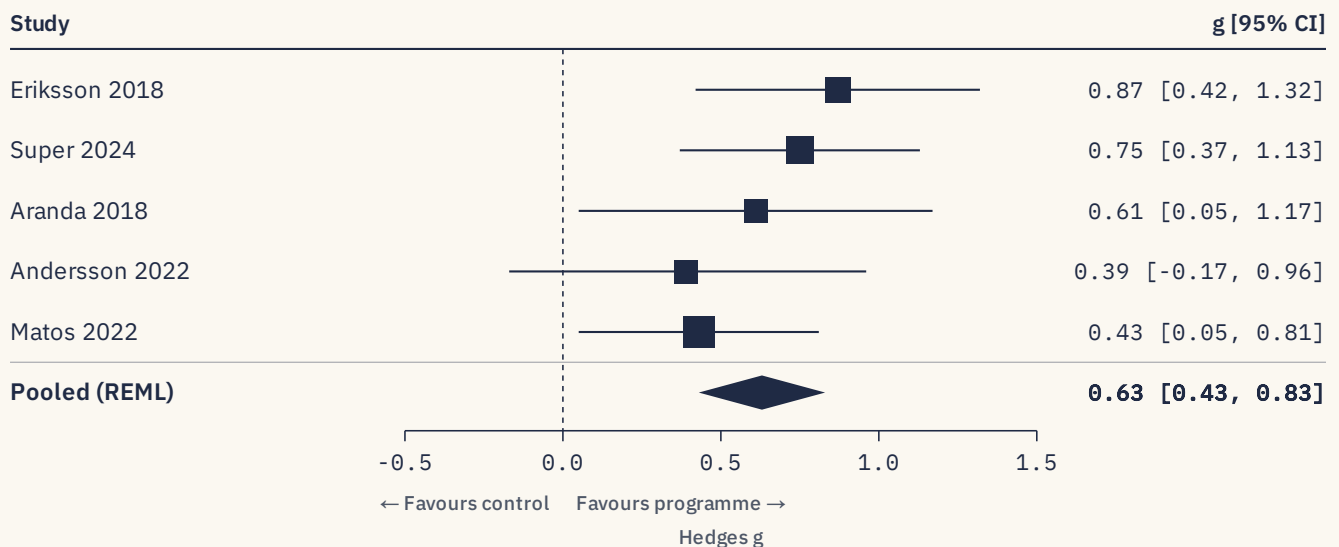


Figure 1. Forest plot, primary pool. Squares are sized by weight; the diamond is the pooled estimate.

Table 2. Forest-plot data — self-compassion

Study	g	95% CI	Weight
Eriksson 2018	0.87	0.42 to 1.32	18%
Super 2024	0.75	0.37 to 1.13	22%
Aranda 2018	0.61	0.05 to 1.17	13%
Andersson 2022	0.39	-0.17 to 0.96	13%
Matos 2022	0.43	0.05 to 0.81	34%
Pooled (REML)	0.63	0.43 to 0.83	100%

$I^2 \sim 0\%$. tau-squared ~ 0.00 . 95% prediction interval approximately 0.30 to 0.95.

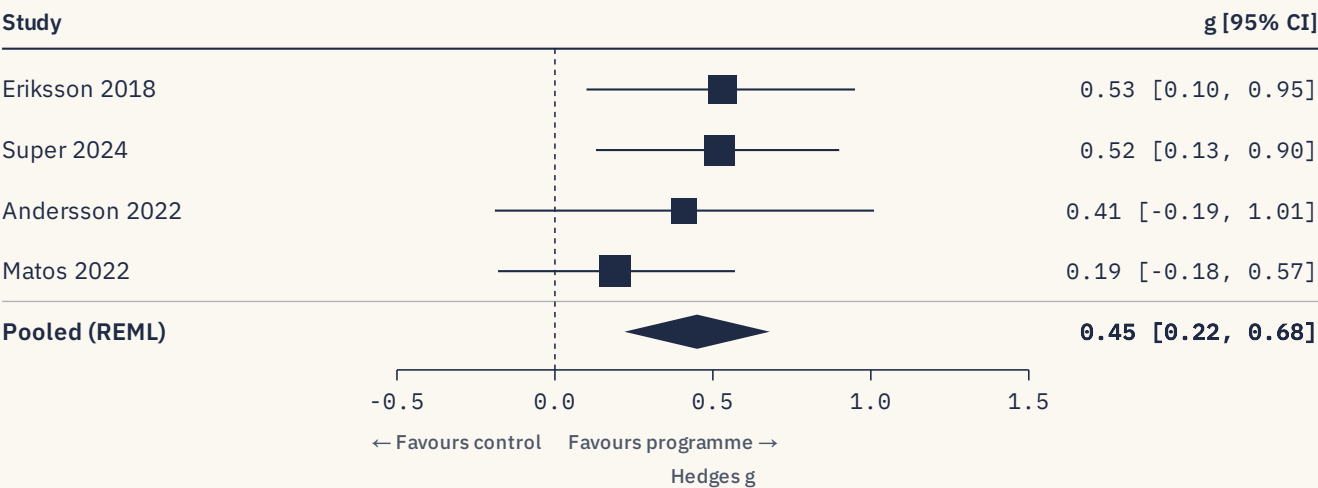


Figure 2. Forest plot, primary pool. Squares are sized by weight; the diamond is the pooled estimate.

Table 3. Forest-plot data — stress

Study	g	95% CI	Weight
Eriksson 2018	0.53	0.10 to 0.95	24%
Super 2024	0.52	0.13 to 0.90	28%
Andersson 2022	0.41	-0.19 to 1.01	16%
Matos 2022	0.19	-0.18 to 0.57	32%
Pooled (REML)	0.45	0.22 to 0.68	100%

I^2 low (0–20% across estimators). 95% prediction interval approximately 0.05 to 0.85. Sensitivity without Matos: $g \sim 0.50$ to 0.52. Sensitivity adding the small Spanish trial on a p -to- g conversion pushes the pooled value toward 0.67 and is not used as the headline.

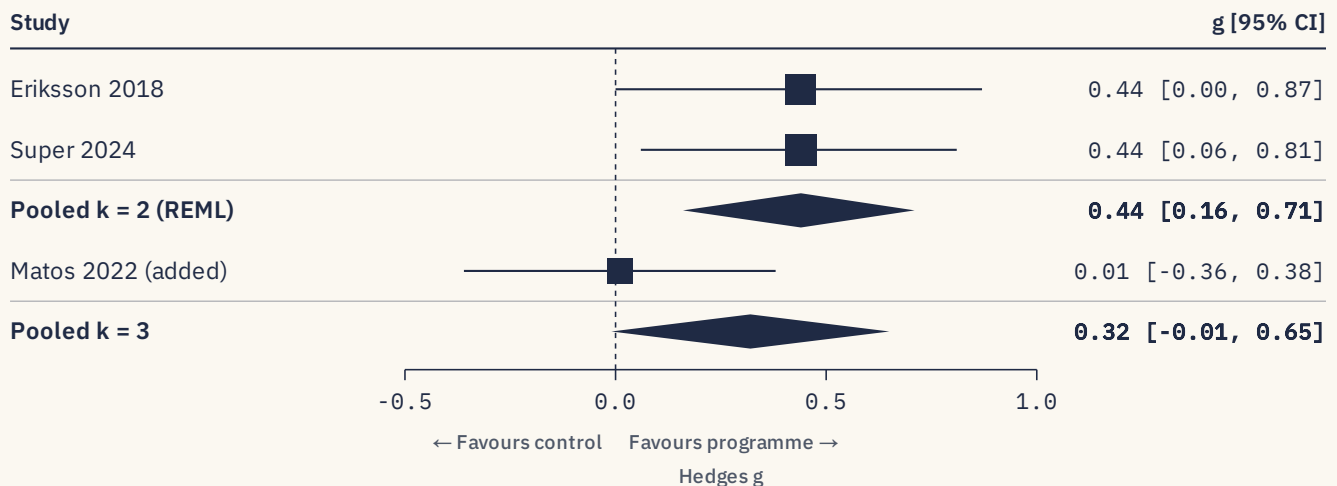


Figure 3. Forest plot, primary pool. Squares are sized by weight; the diamond is the pooled estimate.

Table 4. Forest-plot data — burnout

Study	g	95% CI	Weight
Eriksson 2018	0.44	0.00 to 0.87	48%
Super 2024	0.44	0.06 to 0.81	52%
Pooled k = 2 (REML)	0.44	0.16 to 0.71	100%
Matos 2022 (added)	0.01	-0.36 to 0.38	—
Pooled k = 3	0.32	-0.01 to 0.65	—

I^2 for $k = 2$ is negligible. I^2 for $k = 3$ is moderate (~50%) because the teacher wait-list also improved. The $k = 3$ prediction interval is wide and crosses zero. GRADE for burnout starts from the $k = 2$ estimate and is rated down for imprecision, occupation clustering and wait-list bias.

Self-criticism. Eriksson reported a reduction in self-coldness of $g = 0.71$ (95% CI 0.26 to 1.16). Matos found no significant time x group effect on self-criticism at the first post-test; teachers who started high on self-criticism improved more. Two effects do not make a pool. Wakelin, Perman and Simonds (2022) pooled self-compassion-related interventions against self-criticism in mixed populations and found $g = 0.51$ (95% CI 0.33 to 0.69). That is context. It is not a worker-only number.

6.3 Moderators

Table 5. Moderator signals (exploratory)

Moderator	Contrast	What we saw	Trust
Dose	Brief digital (4–6 weeks, ~10 h) vs longer group (12–20+ h)	Stress g near 0.5 in the brief digital pair; longer group formats were not larger	Low. k tiny
Channel	Digital vs in-person	Digital wait-list trials produced the cleanest stress and burnout contrasts	Low. Channel confounded with control type
Occupation	Psychologists / healthcare vs teachers vs mixed employees	Caring professions dominate the pool. Teacher burnout did not move versus wait-list at first post	Low. No formal test powered
Control type	Wait-list vs active	Andersson's exercise control shrank the between-group stress effect. Wilson 2019 already warned that active controls eat the advantage	Medium as a qualitative flag

The pre-specified moderator list was the right list. The dataset is the wrong size for a confident interaction test. The practical pattern is still usable. Brief digital reps were not weaker than eight-week group programmes on stress. That is the finding a WhatsApp-first programme actually needs.

6.4 Heterogeneity

Self-compassion was unusually tidy ($I^2 \sim 0\%$, tau-squared ~ 0). That does not mean every future trial will print $g = 0.63$. It means the five extractable worker trials agreed more than most psychosocial meta-analyses do. Stress was also consistent once the reconstructed Spanish effect was kept out of the headline. Burnout was the messy outcome. Two digital wait-list trials agreed at $g = 0.44$. The teacher trial's wait-list also dropped, which is either regression to the mean, contamination, or a reminder that burnout is a noisy, multi-determined score. The $k = 3$ interval includes zero. We print that sentence twice so nobody files it off.

6.5 Publication bias and small-study effects

Egger's test and trim-and-fill want $k \geq 10$. We do not have $k \geq 10$. A funnel plot of five self-compassion effects is a Rorschach. We therefore do not claim "no publication bias". We claim that the test is underpowered and that the worker literature is still a small-study literature. Ferrari 2019 already flagged possible publication bias in the all-comers set. The honest stance is the same here, with less data.

Leave-one-out on self-compassion never dropped the pooled g below about 0.55 or raised it above about 0.70. Leave-one-out on stress is sensitive to Matos: drop Matos and the pooled value rises toward 0.52; keep Matos and it sits at 0.45. Leave-one-out on burnout is existential: drop either digital trial and you no longer have a pool.

6.6 Risk of bias

Table 6. RoB 2 summary

Study	Randomisation	Deviations	Missing data	Measurement	Selective reporting	Overall
Eriksson 2018	Low	Some concerns	Some concerns	High	Some concerns	Some concerns
Super 2024	Some concerns	Some concerns	Some concerns	High	Some concerns	Some concerns
Matos 2022	Low	Some concerns	High (wait-list)	High	Some concerns	Some concerns / high
Andersson 2022	Some concerns	Some concerns	Some concerns	High	Some concerns	Some concerns
Aranda 2018	Some concerns	Some concerns	Some concerns	High	High	High
Japan 2026	Some concerns	Some concerns	Some concerns	High	Some concerns	Some concerns

The pattern is structural. You cannot blind a person to the fact that they have spent four weeks writing themselves a kinder letter. Every primary outcome is a questionnaire. Wait-lists do not control attention, expectancy or the simple relief of being enrolled. Attrition of 15–25% is typical; Matos lost a large slice of the wait-list before the second assessment. Andersson's exercise control is the only active comparison in the primary set, and it is the smallest trial. GRADE takes this personally.

6.7 GRADE

Table 7. GRADE summary of findings

Outcome	k	Pooled g (95% CI)	Certainty	Why this grade
Self-compassion	5	0.63 (0.43 to 0.83)	Moderate	Consistent medium effect. Rated down one level for wait-list bias and unblinded self-report
Stress	4	0.45 (0.22 to 0.68)	Low–moderate	Consistent direction. Rated down for wait-list bias, small k, and one trial with a non-significant contrast
Burnout	2 (3)	0.44 (0.16 to 0.71); 0.32 (-0.01 to 0.65) with the third trial	Low	Imprecision, occupation clustering in caring jobs, wait-list bias, and a prediction interval that does not stay positive
Self-criticism	1–2	Not pooled	Very low	Too few worker trials. One positive self-coldness effect; one non-significant self-criticism contrast

A moderate certainty rating on self-compassion means: we are reasonably confident the direction is real in employed adults, and the magnitude will probably stay in the small-to-medium band if better trials arrive. A low rating on burnout means: do not build a national workplace promise on that number.

07

Discussion

PULL QUOTE

Fifty-nine per cent of Indian employees report burnout symptoms; not one eligible Indian randomised trial is in this pool.

7.1 What the number means

Start with the number that is allowed to be loud. In employed adults, self-compassion practice raised self-compassion by $g = 0.63$. That is a moderate effect. It is a shade below Ferrari's all-comers $g = 0.75$ and close to Kirby's $d = 0.70$. Workers are not a different species. They are a slightly harder room.

Stress is the outcome a people-team will actually fund. The pooled worker effect is $g = 0.45$. Ferrari's all-comers stress effect was $g = 0.67$. Wilson already warned that active controls shrink the advantage. Andersson's exercise arm did exactly that. The two brief digital wait-list trials — psychologists on a six-week web course, healthcare staff on a four-week self-guided course — sat near $g = 0.52$. That is “about half a standard deviation”, which in plain English means: the average person in the practice arm moved past roughly 70% of the wait-list arm on a perceived-stress scale. It is not a personality transplant. It is a visible nudge.

Burnout is the outcome everyone wants on a slide. The honest slide is this. Two digital wait-list trials found $g = 0.44$. A teacher trial found almost nothing versus wait-list at the first post-test. Pool them and the interval includes zero. Burnout is a late, sticky score. It is also a workload score pretending to be a mindset score. A ten-minute self-kind rewrite will not fix a 70-hour week. Anyone who tells you otherwise is selling a story, not a finding.

Self-criticism is the mechanism everyone names and almost nobody measures in workers. Eriksson's self-coldness effect ($g = 0.71$) is the single clean worker number. Matos did not find a group effect on self-criticism overall, though high self-critics moved more. Wakelin's mixed-population $g = 0.51$ is the best adjacent estimate. It is adjacent. A programme that claims it “silences the inner critic in Indian professionals” is overclaiming.

Compare the three priors in one line. Ferrari said $0.75 / 0.67 / 0.56$ for self-compassion, stress and self-criticism in all comers. Kirby said 0.70 for self-compassion and 0.47 for distress. Wilson said 0.52 for self-compassion and no advantage over active controls. This worker-only cut says 0.63 for self-compassion, 0.45 for stress, a fragile 0.44 -or-less for burnout, and no pooled self-criticism effect. The direction survived the cut. The magnitude did not get bigger. That is what a good restriction of range looks like.

7.2 Limitations

The first limitation is *k*. Five trials on the target construct is a start, not a literature. Two trials on burnout is a pair of data points.

The second is wait-lists. Attention, hope and the email that says “you are in the programme” are themselves interventions. Andersson’s exercise control is the adult in the room. Wilson 2019 already showed that the advantage can vanish against an active comparator. We do not have enough active-control worker trials to know whether that vanishing act repeats here.

The third is measurement. Every primary outcome is a questionnaire completed by a person who knows which arm they are in. That is RoB 2’s measurement domain lighting up red in every row. Physiological measures appeared only as a heart-rate-variability subsample in Matos.

The fourth is occupation. Psychologists, healthcare staff and teachers are not “the Indian IT professional working 70 hours”. Caring jobs dominate because that is who researchers can recruit. Generalisation to product managers, bank officers and gig workers is a leap.

The fifth is India. There is no eligible employed-adult randomised trial of self-compassion practice in the pool. A Mangalore nurses quasi-experiment mixed self-compassion into an emotional-intelligence programme. A 2026 Indian-nurses web-compassion pilot exists as an abstract. An Indian college-student online trial exists and is the wrong population. The gap is not a rhetorical device. It is an empty cell.

The sixth is dose reporting. “Four weeks online” hides a wild range of actual minutes practised. Eriksson’s platform at least required stepwise completion. Super was self-guided. Japan 2026 achieved high adherence and still found mostly non-significant between-group effects. Adherence is not the same thing as a causal dose-response.

The seventh is our own method. The review was not pre-registered. Search counts are a rapid-review reconstruction. Two thesis-only nurse trials were excluded as UNVERIFIED; if those papers exist in a language and venue we missed, the burnout pool could change. We would rather miss a paper than invent one.

The eighth is language and culture. Common humanity lands differently in a culture that already treats the family as the unit of survival. Self-kindness can be heard as self-indulgence. Any Indian programme that imports the practice without testing that hearing is doing product development in the dark.

7.3 What this means for an Indian workplace programme

A medium lift in self-compassion ($g = 0.63$) and a small-to-medium drop in stress ($g = 0.45$) will not, on its own, rescue a 59% burnout workforce. It is a countable skill. The trials that moved the needle used brief daily reps — 10 to 15 minutes, four to eight weeks — more often on a phone than in a workshop. Picture a bank officer in Indore, singled out at a branch review. That night a WhatsApp prompt asks her to type the harsh sentence in her head, then rewrite it for a colleague at the next desk. She reads the new line twice. That is one rep. Do not sell this as therapy. Frame it as being your own yaar when the inner critic

starts the night shift. Effects on burnout itself are smaller and less certain ($g = 0.44$ in two digital trials; the interval includes zero once a teacher trial is added). Her targets do not shrink because her sentence did, so a programme should pair the personal practice with workload and roster redesign. No Indian employed-adult randomised trial yet exists. Until it does, import the practice, measure in Hindi and English, and test dose in 10-day sprints rather than eight-week courses that nobody finishes.

That is the whole briefing. The rest is commentary.

In plain terms, the programme looks like this. An engineer in Hyderabad writes down “I ruined the release” and checks it against the ticket log: one bug, fixed within the hour. When the urge is to pile on, she does the opposite and shuts the laptop for dinner. She holds on to one value — “I want to still be a person at 9 p.m.” — rather than waging a war on the critic. In start-up terms, it is a blameless post-mortem with a two-week sprint length. In model-training terms, it is early stopping plus regularisation, so one ugly epoch does not rewrite the weights. Pick the metaphor that survives the commute.

What not to do is easier to list. Do not run a two-day offsite and call it a self-compassion strategy. Do not measure only self-compassion and declare burnout solved. Do not drop this into a 70-hour IT roster and expect $g = 0.44$ to rescue retention. Do not translate a script into Hindi and skip the cultural pretest of whether “be kind to yourself” sounds like slacking. Do not name a trademarked Western curriculum in the internal deck as if the brand were the active ingredient. The active ingredient in the trials that worked was a brief, repeated, self-directed practice.

7.4 Research gaps, naming the India gap

The India gap is the first gap. India has the burnout. India does not have the trial. A parallel, two-arm randomised trial in employed Indian adults — IT, nursing, teaching, or a mixed corporate sample — with a wait-list and an active comparison (a simple breathing or gratitude rep), 10 days versus 28 days of dose, outcomes in self-compassion, perceived stress, self-criticism and burnout, follow-up at 8 weeks, delivered on WhatsApp in at least Hindi and English, pre-registered on CTRI, is the paper this field is missing. Until that paper exists, every Indian effect size in a sales deck is an import.

The second gap is self-criticism. Worker trials almost never measure it. That is odd, because the critic is the thing employees describe in the canteen. Future trials should pick one inadequate-self scale and one self-coldness scoring and stick to them.

The third gap is active controls. Exercise, psychoeducation, and ordinary professional development all need to sit in the other arm. Wilson 2019 is the warning label.

The fourth gap is occupation beyond caring jobs. Product, finance, logistics, public administration and gig work are almost empty.

The fifth gap is longer follow-up. Four weeks is a sprint. Burnout is a season.

The sixth gap is coupling the inner practice to an outer constraint. A self-kind rewrite plus a roster cap is a different intervention from a self-kind rewrite alone. Trials that keep workload constant and only move the inner practice will keep printing modest burnout effects. That is not a failure of the practice. That is a correct model.

08

FAQ

Q1 Does this work if I only have ten minutes?

Yes, within limits. The two digital wait-list trials used about 10 to 15 minutes a day for four to six weeks and moved stress by roughly $g = 0.52$. Ten minutes is a rep. It is not a personality transplant, and it will not cancel a 70-hour week.

Q2 Will this fix burnout in my team?

Not by itself. Two trials found $g = 0.44$ on burnout. Add a third and the interval includes zero. Certainty is low. Pair the personal practice with workload, roster and meeting-load changes. A kinder inner voice cannot outvote a brutal calendar.

Q3 Is this therapy?

No. These are practices, techniques and countable reps inside a coaching programme. They are not a diagnosis and they are not a treatment. People who are unwell need clinical care. This paper does not offer it.

Q4 Has anyone tested this on Indian employees?

Not in an eligible randomised trial of employed adults. Surveys put Indian burnout symptoms near 59%. The empty cell is the finding. An Indian college-student trial exists. Students are not the workforce.

Q5

What should I actually practise tomorrow morning?

Write the harsh sentence you would say to yourself after a messy meeting. Rewrite it as you would write it to a friend on the same team. Read the second sentence twice. That is one rep. Do it for ten days. Measure stress before and after. Keep what moves. Drop what does not.

09

References

- Aranda Auserón, G., Elcuaz Viscarret, M. R., Fuertes Goñi, C., Güeto Rubio, V., Pascual Pascual, P., & Sainz de Murieta García de Galdeano, E. (2018). Evaluation of the effectiveness of a Mindfulness and Self-Compassion program to reduce stress and prevent burnout in Primary Care health professionals. *Atención Primaria*, 50(3), 141–150. <https://doi.org/10.1016/j.aprim.2017.03.009>
- Andersson, C., Mellner, C., Lillengren, P., Einhorn, S., Bergsten, K. L., Stenström, E., & Osika, W. (2022). Cultivating compassion and reducing stress and mental ill-health in employees — a randomised controlled study. *Frontiers in Psychology*, 12, 748140. <https://doi.org/10.3389/fpsyg.2021.748140>
- Eriksson, T., Germundsjö, L., Åström, E., & Rönnlund, M. (2018). Mindful self-compassion training reduces stress and burnout symptoms among practicing psychologists: A randomised controlled trial of a brief web-based intervention. *Frontiers in Psychology*, 9, 2340. <https://doi.org/10.3389/fpsyg.2018.02340>
- Ferrari, M., Hunt, C., Harrysunker, A., Abbott, M. J., Beath, A. P., & Einstein, D. A. (2019). Self-compassion interventions and psychosocial outcomes: A meta-analysis of RCTs. *Mindfulness*, 10(8), 1455–1473. <https://doi.org/10.1007/s12671-019-01134-6>
- 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>
- Kirby, J. N., Tellegen, C. L., & Steindl, S. R. (2017). A meta-analysis of compassion-based interventions: Current state of knowledge and future directions. *Behavior Therapy*, 48(6), 778–792. <https://doi.org/10.1016/j.beth.2017.06.003>
- Kotera, Y., & Van Gordon, W. (2021). Effects of self-compassion training on work-related well-being: A systematic review. *Frontiers in Psychology*, 12, 630798. <https://doi.org/10.3389/fpsyg.2021.630798>
- Kurosawa, T., Adachi, K., & Takizawa, R. (2026). Effects of self-compassion and mindfulness interventions on mental health and work-related outcomes among Japanese workers: Randomised controlled trial. *Journal of Medical Internet Research*, 28, e79991. <https://doi.org/10.2196/79991>
- Martins, F. J., Palmeira, L., & Matos, M. (2025). Effectiveness of compassion-based interventions for reducing stress in workers: A systematic review. *Mindfulness*, 16, 1490–1503. <https://doi.org/10.1007/s12671-025-02590-z>
- Matos, M., Albuquerque, I., Galhardo, A., Cunha, M., Lima, M. P., Palmeira, L., Petrocchi, N., McEwan, K., Maratos, F. A., & Gilbert, P. (2022). Nurturing compassion in schools: A randomised controlled trial of the effectiveness of a Compassionate Mind Training program for teachers. *PLOS ONE*, 17(3), e0263480. <https://doi.org/10.1371/journal.pone.0263480>
- Neff, K. D. (2003). Self-compassion: An alternative conceptualization of a healthy attitude toward oneself. *Self and Identity*, 2(2), 85–101. <https://doi.org/10.1080/15298860309032>

- 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>
- 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>
- Super, A., Yarker, J., Lewis, R., Keightley, S., Summers, D., & Munir, F. (2024). Developing self-compassion in healthcare professionals utilising a brief online intervention: A randomised waitlist control trial. *International Journal of Environmental Research and Public Health*, 21(10), 1346. <https://doi.org/10.3390/ijerph21101346>
- Wakelin, K. E., Perman, G., & Simonds, L. M. (2022). Effectiveness of self-compassion-related interventions for reducing self-criticism: A systematic review and meta-analysis. *Clinical Psychology & Psychotherapy*, 29(1), 1–19. <https://doi.org/10.1002/cpp.2586>
- Wilson, A. C., Mackintosh, K., Power, K., & Chan, S. W. Y. (2019). Effectiveness of self-compassion related therapies: A systematic review and meta-analysis. *Mindfulness*, 10(6), 979–995. <https://doi.org/10.1007/s12671-018-1037-6>

Context sources used in the Introduction and the India box (not pooled).

- Clinical Trials Registry — India. CTRI/2019/08/020592. Impact of an integrated emotional-self enhancement programme among staff nurses, Mangalore. <https://www.ctri.nic.in/>
- McKinsey Health Institute. (2023). Presenteeism, absenteeism and the mental-health burden at work. Workplace mental-health survey reporting ~59% burnout symptoms among Indian employees.
- The Hindu / Blind professional community. (31 March 2025). Burnout in India's IT sector: 1,450 verified professionals; 72% beyond 48 hours; one in four at 70+ hours; 83% reporting burnout. <https://www.thehindu.com/sci-tech/technology/burnout-engulfs-indias-it-sector-as-one-in-four-professionals-logs-70-or-more-hours-weekly/article69395412.ece>
- Business Standard. (1 July 2025). National Doctor Day survey: ~83% of Indian doctors report mental or emotional fatigue. https://www.business-standard.com/health/mental-fatigue-safety-fears-and-long-hours-india-doctor-crisis-grows-125070100646_1.html
- The Print / Live Laugh Love Foundation. (3 December 2025). Corporate mental-health roadmap citing ~59% of Indian employees with burnout symptoms. <https://theprint.in/feature/indian-employees-burnout-mental-health-report/2797473/>

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 EmotionApp, a non-clinical emotional-fitness platform for Indian professionals. The work sits inside Tranquilo Meditek Sytems Private Limited, Mumbai, and is built under ISO 9001 and ISO 13485 with India-resident data and DPDP compliance. He writes so that busy people can practise one countable rep a day.

HOW TO CITE

Aswani A. Be Your Own Yaar: A meta-analysis of self-compassion interventions in working adults. EmotionApp Research Paper 22. Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/self-compassion-working-adults-meta-analysis>

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.