

Psychological Capital, Trained

A Meta-Analysis of PsyCap Interventions and Business Outcomes

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



KEY NUMBER · D

0.34

Prior reviews put the PsyCap-construct training effect at $d = 0.34$ (Lupşa 2020, 41 trials) to $d = 0.46$ post-intervention (Loghman 2025, 40 studies).

Psychological Capital, Trained

A Meta-Analysis of PsyCap Interventions and Business Outcomes

Lupşa 2020 pooled wellbeing; this pools performance, absenteeism and turnover — except a pooled estimate is not yet possible.

PRISMA-informed rapid evidence synthesis, 2006–2026

Narrative systematic review. No pooled g for business outcomes.

CONTENTS

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

INDEXING TITLE

Psychological-capital interventions and business outcomes: a meta-analysis of performance, absenteeism and turnover

AUTHOR

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

PUBLISHED

22 September 2026 · English edition · emotionapp.in/research/psycap-interventions-business-outcomes

01

Abstract

Psychological-capital (PsyCap) programmes train four state-like resources — hope, efficacy, resilience and optimism — in employed adults. **A pooled estimate for business outcomes is not yet possible: fewer than five eligible controlled trials report extractable, comparable effects on the same outcome, so this paper is a systematic review with narrative synthesis rather than a new meta-analytic pool.** Prior reviews found small-to-medium effects on the PsyCap construct itself (Lupşa and colleagues, 2020: $d = 0.34$ across 41 trials, $N = 3,911$; Loghman and colleagues, 2025: post-intervention $d = 0.46$ across 40 studies, $N = 4,207$). Those reviews pooled wellbeing and the construct. They did not settle performance, absenteeism or turnover.

This review asked a narrower question for working adults: do PsyCap programmes move business outcomes, not only the construct and wellbeing? Searches covered 2006–2026 across PubMed, PsycINFO-style records, Scopus-indexed sources, Web of Science-indexed sources, Google Scholar, CTRI, ClinicalTrials.gov and OSF, plus citation chasing of the two prior reviews. Eligible designs were randomised or controlled trials in employed adults. Primary outcomes were performance (self-rated separately from other-rated or objective), absenteeism, turnover or turnover intention, and PsyCap.

Findings are uneven. PsyCap itself moves. That claim is supported by two independent meta-analyses and by later trials in China, the United States, Spain, Bulgaria and India. Self-rated performance rose in one Chinese controlled reading-materials programme (Zhang and colleagues, 2014; ANCOVA partial $\eta^2 = 0.20$) and in uncontrolled pre–post manager samples. Other-rated and objective performance evidence is thin and mixed. Turnover-intention evidence rests on two Chinese trials: one fragile (Da and colleagues, 2020) and one large but single-site and unblinded (Sun, Gong and Yuan, 2026). Actual turnover was not found as an outcome. Absenteeism was not found as an outcome in any eligible PsyCap-intervention trial. One Indian PCI trial exists; it measured stress and job insecurity, not performance, absenteeism or turnover.

Certainty is moderate for the construct, low for self-rated performance and turnover intention, and very low or empty for other-rated performance, actual turnover and absenteeism. Indian workplaces should treat PsyCap practice as a trainable personal resource with a plausible, still unproven, path to business metrics. The India gap is explicit: no verified Indian randomised trial of a PsyCap programme has yet reported performance, absenteeism or turnover.

02

Key findings box

KEY FINDINGS

0.34

Prior reviews put the PsyCap-construct training effect at $d = 0.34$ (Lupşa 2020, 41 trials) to $d = 0.46$ post-intervention (Loghman 2025, 40 studies).

0

Zero eligible PsyCap-intervention trials measured absenteeism. That cell is empty, not null.

2

Two Chinese controlled trials measured turnover intention: Da 2020 was non-significant between groups at post-test; Sun 2026 reported a large drop that has not been replicated outside one emergency-nursing site.

1

One Chinese controlled trial (Zhang 2014, $n = 234$) found a self-rated job-performance advantage after a reading-materials PsyCap programme (ANCOVA partial $\eta^2 = 0.20$), still visible at three months.

1

One Indian PCI randomised trial exists (Patnaik, Mishra and Mishra, 2022, telecom, $n = 234$). It moved PsyCap, stress and job insecurity. It did not measure performance, absenteeism or turnover.

03

Definition block

DEFINITION

Psychological capital is a bundle of four trainable, state-like resources that people bring to work: hope (willpower plus pathway-planning when the first plan dies), efficacy (the grounded sense that a stretch task is doable if effort is applied), resilience (bounce-back, and sometimes bounce-forward, after a knock), and optimism (a realistic habit of attributing good events to lasting causes and bad events to temporary ones). Luthans and colleagues named the bundle HERO and argued it is more plastic than a personality trait and more durable than a mood. In plain English it is the inner working capital of a professional day: not a diagnosis, not a therapy, not a personality transplant. It is a set of practices. You can run them as short face-to-face workshops, as daily online reps, or as coaching. The claim under review is whether those practices move the numbers a business already tracks — output, days away, exits — rather than only the questionnaire that measures the bundle itself.

04

Introduction

Three new hires join a Gurugram operations team in the same week. Same laptop, same targets, same induction. By the second month one has asked for the messy client file. Another has found a second route when the first report template broke. The third takes no stretch task and calls it being careful. On paper their resources are equal. What differs is how each puts inner working capital to use: hope, efficacy, resilience and optimism. Those four are not ornaments. They are working capital. The question for an Indian professional is whether a programme that trains them moves performance, absenteeism and turnover — the metrics that show up in a unit-economics slide.

Lupşa, Vîrga, Maricuţoiu and Rusu (2020) answered a neighbouring question. Across 41 controlled trials and 3,911 employees and students, interventions designed to raise PsyCap produced a small overall effect on PsyCap variables, $d = 0.34$. Omnibus PsyCap moved by $d = 0.26$. Resilience moved most ($d = 0.49$), then efficacy ($d = 0.37$), optimism ($d = 0.36$) and hope ($d = 0.22$). The authors also reported that wellbeing and performance moved. They noted, in a scatter-plot comment that has been under-quoted, that secondary outcomes generally topped out around $d = 0.25$ even when the primary PsyCap effect ran higher. That is the first honest ceiling.

Loghman, Ramirez-Perez, Bohle and Martin (2025) updated the construct-only picture. Forty studies and 4,207 participants, searched through February 2023, produced a larger post-intervention effect on PsyCap ($d = 0.46$; outlier-adjusted $d = 0.43$) and a follow-up effect that held for the bundle, hope, resilience and optimism — not for efficacy. Their most useful surprise was not the size. It was the type. Programmes labelled “PsyCap-focused” were not the strongest way to raise PsyCap. Broader positive-psychology, mindfulness and leadership-development programmes often did more. If you are an operator, that sentence should rearrange your catalogue. The brand of the workshop is not the active ingredient.

Neither review pooled absenteeism. Neither review pooled turnover. Neither review separated self-rated performance from supervisor ratings or from sales, errors, or output. Avey, Reichard, Luthans and Mhatre (2011) had already shown the correlational backbone: PsyCap tracks performance at $r = 0.26$ across 24 samples and 6,931 people, a little higher for self-ratings ($r = 0.33$) and supervisor ratings ($r = 0.35$) than for objective metrics ($r = 0.27$), and it tracks lower turnover intention. Correlation is not a training effect. It is a map. This paper walks the causal road those maps pointed at.

Why does the question matter in India now? Because Indian professional life is a high-load, high-context, high-language operating system. Attrition in corporate India sat at 16.2 per cent in 2025 on the Aon survey series, the lowest in five years, with voluntary exits still about three-quarters of the flow. That is not a crisis headline. It is a cost. White-collar absenteeism is priced, in compilation estimates, in the tens of thousands of rupees per employee per year. Those figures are context, not outcomes of PsyCap trials. They explain why a finance lead will fund a programme only if it can be argued, without bluff, against performance, absence and exits.

A simple loop helps here, and should stay in its lane. Thoughts, feelings and actions feed each other. PsyCap practices are mostly thought-and-action drills. A sales manager whose pitch dies before lunch writes two other routes by tea-time (hope). An analyst takes one stretch task and notes each part she gets right (efficacy). A team lead who misses a deadline names the bounce: what he will do differently tomorrow (resilience). A bad week is read as a bad week, not a life sentence (optimism). Some of it is acting against the mood, or noticing a sticky story as a story and taking one small step towards what matters. None of that makes the programme a treatment. It makes it a set of reps. EmotionApp’s own design — countable daily reps, WhatsApp-first delivery, 23 Indian languages, India-resident data, DPDP-compliant, ISO 9001 and ISO 13485 — is an implementation bet on that distinction. Reps compound. One-off pep talks do not.

Machine-learning people already know the trap this literature is in. Training loss on the PsyCap questionnaire is going down. That is the construct. The test set is business outcomes. If you only report training loss you will overfit the story. Regularisation, in this paper, is the eligibility rule: employed adults, PsyCap-type programme, a control, and an outcome a business already measures. The

India office scene writes itself. A team lead in Pune runs a two-hour hope-and-efficacy workshop on a Saturday. Monday the dashboard is unchanged. Wednesday two people are on casual leave. Friday a high performer resigns for a 20 per cent bump. The workshop may have been good. The dashboard does not know that yet. This review is the attempt to read the dashboard without lying about the sample size.

The logline is therefore blunt. Lupşa pooled wellbeing. Loghman pooled the construct and its durability. This paper asked for performance, absenteeism and turnover. The short version of the answer is that the construct is trainable, the business-outcome trials are few, absenteeism is an empty cell, and India has not yet run the trial that would let a CHRO treat PsyCap practice as a proven lever on exits or days away.

05

Methods

5.1 Protocol stance

This is a rapid evidence synthesis that follows PRISMA 2020 items as far as feasible. It was not pre-registered as a new review protocol. It positions itself against two pre-existing syntheses (Lupşa et al., 2020; Loghman et al., 2025) rather than repeating their construct-only pools. The pre-specified fallback, set before extraction closed, was: if fewer than five eligible trials report extractable effects on a given business outcome, do not pool; deliver a systematic review with narrative synthesis and say so in plain English. That fallback was triggered.

5.2 Eligibility

Population. Employed adults. Students were excluded from the business-outcome narrative unless a trial also reported a separate employed subsample. Mixed student–professional samples were described and not treated as workplace evidence.

Intervention or exposure. Programmes explicitly designed to develop psychological capital or its four resources (hope, efficacy, resilience, optimism), including the original micro-intervention model and close descendants (reading-materials programmes, daily online self-learning, multi-week group programmes). Adjacent positive-psychology and coaching programmes that measured PsyCap and a business outcome were included in the narrative as adjacent evidence and labelled as such.

Comparator. Any control: wait-list, placebo activity, routine training, or decision-making exercise.

Outcomes. Primary: job or work performance; absenteeism or sick days; turnover or turnover intention; PsyCap. Performance was coded as self-rated, other-rated (supervisor, peer, employee-of-leader), or objective (sales, output, appraisal score from organisational records). These three were never collapsed in a single cell.

Design. Randomised controlled trials and controlled trials with a comparison group. Uncontrolled pre–post samples were described as uncontrolled and excluded from any informal pooling logic.

Years. 2006 to 2026. 2006 is the year the micro-intervention paper appeared.

Exclusions. Correlational studies of trait PsyCap; clinical samples whose goal was a diagnosis-related endpoint; school-absenteeism trials in children; workplace programmes that never measured PsyCap or a PsyCap-family resource.

5.3 Information sources and search strings

The pre-specified sources were searched or queried through a combination of native interfaces, open indexes, publisher pages and registry portals: PubMed, PsycINFO-style bibliographic records, Scopus-indexed records, Web of Science-indexed records, Google Scholar, the Clinical Trials Registry — India (CTRI), ClinicalTrials.gov, and OSF preprints. Backward and forward citation chasing started from Lupşa et al. (2020), Loghman et al. (2025), Luthans, Avey, Avolio and Peterson (2010), and Avey et al. (2011).

PubMed core string

("psychological capital"[tiab] OR PsyCap[tiab] OR "psychological capital intervention"[tiab]) AND (intervention[tiab] OR training[tiab] OR programme[tiab] OR program[tiab] OR PCI[tiab]) AND (randomi*[tiab] OR "controlled trial"[tiab] OR "control group"[tiab] OR experiment*[tiab]) AND (2006:2026[dp])*

Business-outcome filter (AND)

(performance[tiab] OR "job performance"[tiab] OR "work performance"[tiab] OR absenteeism[tiab] OR "sick leave"[tiab] OR turnover[tiab] OR "intention to leave"[tiab] OR "intention to quit"[tiab] OR "turnover intention"[tiab])

PsycINFO, Scopus and Web of Science used the same Boolean, field-adapted. Google Scholar: “psychological capital” AND intervention AND (randomised OR randomized) AND (performance OR absenteeism OR turnover). CTRI: “psychological capital”. ClinicalTrials.gov: “psychological capital” intervention. OSF: “psychological capital” intervention.

Hit counts from subscription databases are not invented here. Records were identified from the included lists of the two prior reviews, from targeted database and web searches, from registries, and from citation chasing. Full texts of candidate intervention trials were screened against the PICO above. That is a limitation of this synthesis and is declared in the risk-of-bias and certainty sections.

Registry notes, not pooled: ClinicalTrials.gov entries in the family include NCT05718973 (video PsyCap), NCT06706648 (nurse PsyCap programme) and NCT07127809 (mindfulness-PsyCap in nurses). None contributed published, extractable business-outcome results that met eligibility at the search close.

- 5.4 Extraction fields** n (randomised and analysed, by arm); design; country; language of delivery; dose in sessions and minutes; delivery channel; guidance (facilitated, self-guided, coaching); outcome measure and rating source; time-points; effect-size ingredients (means, SDs, F, t, partial η^2 , or author-reported d/g). Where authors reported a standardised mean difference, it was used. Where means and SDs were available, Hedges g was computed as a descriptive, study-level quantity for the tables. Study-level g is not a pooled estimate. Conversion of partial η^2 to d was treated as crude and is flagged wherever it appears.
- 5.5 Analysis plan** The pre-specified model was a random-effects REML meta-analysis of Hedges g with 95 per cent confidence intervals, I^2 , τ^2 , a 95 per cent prediction interval, Egger's test, trim-and-fill, leave-one-out sensitivity, RoB 2, and GRADE. Pre-specified moderators were micro-intervention versus longer programme, industry, and country.
- That model was not run for business outcomes. The rule was numerical, not rhetorical: fewer than five eligible trials with extractable, comparable effects on the same outcome. Combining performance, turnover intention and absenteeism into one "business-outcome g" would have been a category error. Absenteeism had zero trials. Actual turnover had zero trials. Performance and turnover intention did not share a metric. PsyCap as a construct was not re-pooled. Two competent reviews already exist.
- 5.6 Risk of bias and certainty** Randomised trials were appraised with the RoB 2 domains in narrative form: randomisation process, deviations from intended programme, missing outcome data, measurement of the outcome, and selection of the reported result. Uncontrolled pre-post samples were labelled high risk for causal claims by design. Certainty of evidence used GRADE, outcome by outcome. Expected downgrades were risk of bias (no blinding of self-report), inconsistency, indirectness (students in prior pools; nurses-only samples), and imprecision (small k and small n).
- 5.7 PRISMA-style flow** A precise machine-generated database count is not claimed. The flow below is the honest account of this synthesis.
- Identification. Included-study lists of Lupşa et al. (2020; 41 trials) and Loghman et al. (2025; 40 studies) were taken as the backbone. Targeted searches for 2020–2026 trials, registry hits, and citation chasing added further candidates.
 - Screening. Titles and abstracts were screened for employed adults, a PsyCap-type programme, a control, and at least one primary outcome.
 - Included for narrative synthesis of business outcomes. Zhang et al. (2014); Da, He and Zhang (2020); Sun, Gong and Yuan (2026); Peláez Zuberbuhler, Salanova and Martínez (2020) as adjacent coaching evidence; Peláez, Coe and Salanova (2020) as adjacent strengths-coaching evidence; Luthans et al. (2010) main manager sample as uncontrolled supporting evidence only; Hodges (2010) dissertation as a weak-signal RCT.

- Construct-only context. Luthans, Avey and Patera (2008); Dello Russo and Stoykova (2015); Carter and Youssef-Morgan (2022); Patnaik, Mishra and Mishra (2022).
- Excluded from business-outcome causal claims. Student-only RCTs; correlational field studies of absenteeism (Avey, Patera and West, 2006); longitudinal non-intervention models (Peterson et al., 2011); delivery-mode comparisons without a no-programme control.

06 Results

6.1 Study characteristics

The business-outcome evidence base is small, geographically clustered, and mostly self-report. Language of delivery was Chinese in the three China trials, English in the US trials, Spanish in the Spanish coaching trials, and not separately reported as a bilingual protocol in the Indian telecom trial. No trial in this table delivered the programme in an Indian regional language.

Table 1. Characteristics of verified intervention reports touching business outcomes or the Indian gap

Study	n	Design	Country	Dose	Business outcome	Source
Zhang et al. 2014	234 analysed	Controlled	China	Reading materials	Job performance	Self-rated
Da, He & Zhang 2020	104 (38/31/35)	3-arm RCT	China	5 × ~30 min online	Turnover intention	Self-rated
Sun, Gong & Yuan 2026	122 (61/61)	Parallel RCT	China	6 × 90 min F2F	Turnover intention	Self-rated
Peláez Zuberbuhler 2020	41 managers	Wait-list CT	Spain	Workshop + 3 sessions	In-role / extra-role	360°
Peláez, Coó & Salanova 2020	60 (35/25)	Wait-list CT	Spain	5-week micro-coaching	Performance	Self + supervisor
Luthans et al. 2010 mgrs	80	Uncontrolled	USA	1 × 2 hours	Performance	Self + manager
Hodges 2010	Mgrs + associates	RCT	USA	Micro-intervention	Performance	Self; ceiling
Patnaik et al. 2022	234 (124/110)	RCT	India	PCI	None of the three	—

6.2 PsyCap construct — already pooled elsewhere

This paper does not re-pool PsyCap. The two prior reviews are the construct result.

Lupşa et al. (2020): 41 trials, N = 3,911, employees and students. All PsyCap variables $d = 0.34$ (95 per cent CI 0.24 to 0.44), $I^2 \approx 51$ per cent. Omnibus PsyCap $d = 0.26$ (0.14 to 0.38), $k = 9$, homogeneous. PCI subset $d = 0.26$, $k = 6$. Component effects: hope 0.22, efficacy 0.37, optimism 0.36, resilience 0.49.

Loghman et al. (2025): 40 studies, $N = 4,207$, search to February 2023. Post-intervention PsyCap $d = 0.46$; outlier-adjusted $d = 0.43$ (0.27 to 0.60). Components post-intervention: hope 0.75, efficacy 0.46, resilience 0.58, optimism 0.63. Follow-up: PsyCap $d = 0.55$; hope, resilience and optimism remained positive; efficacy did not. Working-adult subgroup post-intervention $d = 0.48$, $k = 22$. Online versus face-to-face was not a reliable moderator. Type of programme was: broader positive-psychology and mindfulness programmes often outlanded branded PsyCap workshops.

Newer construct-only trials sit comfortably inside those bands. Luthans, Avey and Patera (2008) put 364 working adults through a short web programme versus a decision-making control; the change-score between-group g was about 0.28. Dello Russo and Stoykova (2015) replicated the workshop in a small Bulgarian mixed sample ($n = 40$) and held the gain at one month. Carter and Youssef-Morgan (2022) showed face-to-face, online and micro-learning delivery can all move PsyCap in US construction employees ($n = 228$), without a no-programme control. Patnaik, Mishra and Mishra (2022) showed an Indian telecom PCI can raise PsyCap and cut stress and job insecurity ($n = 234$). Da et al. (2020) found a within-programme PsyCap d of 0.50 immediately and 0.56 at one week. Zhang et al. (2014) produced a post-test PsyCap g of about 0.49 to 0.50 against control. Sun et al. (2026) produced much larger PsyCap g values (about 1.14 at post; 0.86 at three months) in one nursing sample; those numbers sit well above the prior pools and are treated as outliers pending replication.

6.3 Performance

Self-rated and other-rated performance are reported separately, as specified.

Self-rated performance

Zhang et al. (2014) is the cleanest controlled workplace trial with a job-performance questionnaire. Two hundred and thirty-four Chinese employees completed a structured reading-materials PsyCap programme or control. Intervention means: performance 4.53 (SD 0.62) before, 4.72 (0.59) after, 4.63 (0.62) at three months. Control: 4.61 (0.62), 4.61 (0.61), 4.64 (0.58). A raw post-test Hedges g of 0.18 (95 per cent CI -0.07 to 0.44) is small and crosses zero. A gain-score g of about 0.31 (0.06 to 0.57) is more favourable because the control line was flat. The authors' ANCOVA, adjusting pretest and spanning post and retest, gave $F(1, 231) = 58.08$, $p < .001$, partial $\eta^2 = 0.20$. That partial η^2 is a large association in ANCOVA language. It should not be casually rewritten as $d = 1.0$ and dropped into a forest plot. At three months the intervention group remained higher than its own pretest on hope and on performance; the control group's PsyCap retest SD of 1.47 is anomalous against a pretest SD of 0.55 and is flagged as a data-quality concern.

Luthans et al. (2010), manager sample, $n = 80$, no control: self-rated performance 7.43 to 8.41, $t = 9.14$, $d = 0.96$. That number is famous and uncontrolled. Demand characteristics and a one-week window can produce a number that large without a single extra unit of output. It is supporting colour, not confirmatory evidence.

Hodges (2010), financial-services RCT with a real control, hit a ceiling. Self-rated performance moved from 8.29 to 8.34 in the treated managers. That is a rounding error with a sample attached. The more interesting Hodges signal was contagion: associates of treated managers showed some PsyCap lift over six weeks. Contagion is a mechanism hypothesis, not a performance result.

Peláez Zuberbuhler et al. (2020) reported large within-group self-rated in-role gains in a small coaching-based leadership sample. The programme is adjacent — coaching that includes PsyCap — not a branded micro-intervention.

Other-rated and objective performance

Luthans et al. (2010) manager-rated performance, still uncontrolled: 7.66 to 8.20, $t = 2.34$, $p = .025$, $d = 0.43$ (95 per cent CI about 0.10 to 1.00). The interval is wide. The design has no control.

Peláez Zuberbuhler et al. (2020): supervisor in-role performance, group $\times$ time interaction not significant immediately ($F = 1.88$, $p = .17$). Some follow-up contrasts favoured the coached managers, including extra-role behaviour. Employee-rated in-role effects were small at post ($d \approx 0.32$) and moderate at follow-up ($d \approx 0.57$). Sample: 41 managers. GRADE cannot climb on $n = 41$.

No eligible PsyCap-intervention RCT in this search used sales revenue, error rates, utilisation, or another organisational-record metric as the primary between-group outcome. Peterson et al. (2011) linked within-person PsyCap change to supervisor ratings and to individual sales revenue over time in a financial-services firm ($n = 179$). That is a latent-growth model, not a randomised programme. It belongs in the mechanism paragraph, not in the intervention table.

Table 2. Study-level effects on performance (not pooled)

Study	Source	Contrast	Effect	95% CI	Note
Zhang 2014	Self-rated	Post, int vs ctrl	$g = 0.18$	-0.07 to 0.44	Crosses zero
Zhang 2014	Self-rated	Gain-score	$g \approx 0.31$	0.06 to 0.57	Control flat
Zhang 2014	Self-rated	ANCOVA post+retest	$\eta^2 p = 0.20$	—	Do not convert to forest g
Luthans 2010	Self-rated	Pre-post, no ctrl	$d = 0.96$	0.76 to 1.19	Uncontrolled
Luthans 2010	Manager-rated	Pre-post, no ctrl	$d = 0.43$	0.10 to 1.00	Uncontrolled
Hodges 2010	Self-rated	RCT	~ 0	—	Ceiling
Peláez 2020	Supervisor in-role	Group $\times$ time	$F = 1.88$, $p = .17$	—	Immediate NS
Peláez 2020	Employee in-role	Post	$d \approx 0.32$	—	Small n

No pooled weight column is shown. Weights without a model are theatre.

6.4 Absenteeism

No eligible PsyCap-intervention controlled trial measured absenteeism, sick days, or an organisational absence register.

Avey, Patera and West (2006) showed that PsyCap levels tracked both voluntary and involuntary absence in a field sample. That is association. It is not a programme effect. Workplace-intervention reviews that do pool absenteeism

(for example general counselling or problem-solving programmes) are a different literature and were not grafted onto this one.

This is the emptiest, most publishable cell in the paper. The absence of trials is the finding.

6.5 Turnover and turnover intention

No eligible trial measured actual exits. Two Chinese RCTs measured turnover intention.

Da, He and Zhang (2020): 104 full-time employees, three arms, five days of online self-learning. The programme moved PsyCap (ANCOVA group effect at T2 $F = 4.55$, $p = .013$; at T3 $F = 5.02$, $p = .008$; within-programme $d = 0.50$ then 0.56). Turnover intention did not show a clean between-group win. ANCOVA group effect on turnover intention at T2 $F = 1.54$, $p = .219$. Within-programme change $d = -0.27$ and not significant. A descriptive T2 Hedges g of about -0.11 (-0.56 to 0.35) is compatible with noise. The authors' abstract is more optimistic than their ANCOVA table. This review follows the table.

Sun, Gong and Yuan (2026): 122 emergency nurses, one Zhejiang tertiary hospital, six weekly 90-minute facilitated sessions versus routine emergency training. Turnover-intention scale means: intervention 16.48 (3.06) at baseline, 10.61 (2.00) at post, 10.90 (2.11) at one month, 11.22 (2.15) at three months. Control: 16.45 (3.08), 16.31 (3.02), 16.22 (3.00), 16.17 (3.03). Group $\times$ time $F = 89.62$, $p < .001$, partial $\eta^2 = 0.43$. Study-level Hedges g at post is about -2.21 (-2.66 to -1.76); at three months about -1.87 (-2.30 to -1.45). Work engagement moved in parallel (partial $\eta^2 = 0.34$). PsyCap total moved (partial $\eta^2 = 0.29$; $g \approx 1.14$ at post, 0.86 at three months).

Those Sun effects are far larger than anything in Lupşa or Loghman. Possible readings, in order of charity: a high-stress nursing context with more headroom than office samples; a well-run six-week facilitated programme versus a thin five-day online module; expectation effects in an unblinded self-report trial at a single site; or an over-estimate. GRADE will not let a single-site nursing trial set the industry number.

Table 3. Study-level effects on turnover intention (not pooled)

Study	Contrast	Effect	95% CI	Note
Da et al. 2020	T2 int vs ctrl	$g \approx -0.11$	-0.56 to 0.35	ANCOVA $p = .219$
Da et al. 2020	Within int	$d = -0.27$	NS	Authors still claimed a reduction
Sun et al. 2026	T1 int vs ctrl	$g \approx -2.21$	-2.66 to -1.76	Single site, unblinded, self-report
Sun et al. 2026	T3 int vs ctrl	$g \approx -1.87$	-2.30 to -1.45	Held at 3 months; unreplicated

6.6 Moderator table

Table 4. Pre-specified moderators

Moderator	Pre-specified question	What the evidence can support
Micro vs longer	2-hour PCI vs multi-week	Construct: both work. Business outcomes: the only large TI effect is the 9-hour nursing programme. The 5-day online module did not move TI. k too small to test.
Industry	Cross-industry	Nursing, mixed employees, automotive coaching, financial services, telecom India. No test.
Country	Cross-country	Business-outcome RCTs are China-heavy. India has a PCI RCT without business outcomes. No clean US/EU PsyCap RCT with extractable performance g in employed adults survived the cut.

Heterogeneity statistics (I^2 , τ^2 , prediction interval) are not reported for a model that was not fitted. Informal dispersion is obvious: Zhang's raw performance $g = 0.18$ and Sun's turnover-intention $g \approx -2.21$ do not belong in one average.

6.7 Publication bias

Egger's test and trim-and-fill were pre-specified for a pool. They were not run. Funnel-plot theatre around four study-level g values would be decorative. Prior reviews did examine small-study effects on the construct and did not find a result-changing bias for PsyCap itself (Loghman et al., 2025). That reassurance does not transfer automatically to performance or turnover intention.

6.8 Risk-of-bias summary

Table 5. RoB 2-style narrative

Study	Randomisation	Measurement	Overall for business claim
Zhang 2014	Random split after pretest; details thin	Self-report performance	High
Da 2020	RCT, three arms, IRB stated	Self-report TI; abstract warmer than ANCOVA	Some concerns to high
Sun 2026	Computer-generated, opaque envelopes	Self-report TI; single site; huge effects	High for the size of the claim
Peláez 2020	Wait-list, $n = 41$	360° helps	Some concerns; imprecise
Luthans 2010 mgrs	No random control	Self and manager	High (design)
Hodges 2010	RCT	Self; ceiling	Some concerns; uninformative
Patnaik 2022	RCT vs placebo	No business outcome	Not applicable

Common pattern: no trial blinded participants to a PsyCap-looking programme when the outcome was a questionnaire. Self-report plus an unblinded programme is the standard risk in this literature. It is also the reason other-rated and objective metrics were pre-specified.

6.9 GRADE

Table 6. Certainty of evidence

Outcome	Estimate available	k	Certainty	Why
PsyCap construct	d = 0.34 to 0.46 from prior MAs	Many	Moderate	Two MAs; mixed samples; delivery heterogeneity
Self-rated performance	Zhang $g = 0.18$; $\eta^2 p = 0.20$	1–2	Low	Few trials; self-report; raw vs ANCOVA diverge
Other-rated / objective	Luthans $d = 0.43$ uncontrolled; Peláez NS	<5	Very low	Design and imprecision
Turnover intention	Da NS; Sun very large	2	Low	Inconsistency; both Chinese; both self-report
Actual turnover	None	0	Empty	No trial
Absenteeism	None	0	Empty	No trial

07

Discussion

PULL QUOTE

Two Chinese RCTs measured turnover intention and did not agree: one g near zero, one g near -2.2 in a single nursing site.

7.1 What the numbers mean

Start with what is solid. PsyCap is trainable. Two meta-analyses, a decade apart, say so. The effect on the bundle is small to medium. It is larger than a pep talk and smaller than a personality rewrite. Resilience and hope often move more than the marketing department expects. Efficacy moves in the room and fades at follow-up more than the other three (Loghman et al., 2025). That fade is an implementation clue. Efficacy is built by mastery reps after the workshop, not by the workshop applause.

Now the test set. Business outcomes were supposed to be the generalisation check. They are not ready to average.

Self-rated performance has one respectable Chinese controlled trial and a cluster of uncontrolled or adjacent results. If you are a hopeful reader you will quote Zhang's partial $\eta^2 = 0.20$ and Luthans's self-rated $d = 0.96$. If you are a sceptical operator you will quote Zhang's raw post-test $g = 0.18$ and Hodges's 8.29 to 8.34. Both readings are in the file. The sceptical reading is the one that survives contact with other-rated data.

Other-rated performance is the adult in the room and it is quiet. Manager-rated change in an uncontrolled two-hour workshop was $d = 0.43$. Supervisor in-role change in a three-month coaching programme did not clear the group $\times$ time test at post. Objective output in a randomised PsyCap programme is, as of this search, a rumour.

Turnover intention is a tale of two Chinese trials. Five days online: no reliable between-group effect. Six weeks face-to-face with nurses: a very large effect that looks like a different species from the rest of the literature. The temptation is to build the whole slide around the biggest number in the deck. Hold it lightly. Look at the context. Emergency nursing is a high-base-rate setting for wanting to leave. A six-week group programme with a named facilitator, video fidelity checks and WeChat catch-up sessions is not a five-day PDF. External validity to an Indian IT shift, a bank operations floor, or a gig-work cohort is not free.

Absenteeism is the blank column of this literature. Everyone cites the 2006 correlational paper. No one has run the intervention trial with the absence register as the outcome. That is not a philosophical gap. It is a missing spreadsheet.

Picture the budget review. The L&D head shows a PsyCap score that rose after the workshop. The finance lead asks which line on her own sheet moved: billed hours, error rate, sick-leave days, resignations accepted. A questionnaire filled in after an unblinded workshop is soft ground. A pre-registered outcome the finance lead already believes is firm ground. This literature has built a lot on soft ground that is, to be fair, well measured. PsyCap psychometrics are not the problem. The problem is the outcome diet.

Entrepreneurship language helps. Unit economics of a PsyCap programme are attractive on paper: two hours, or five daily reps, or six weekly sessions; no licence for a patented molecule; delivery that can ride WhatsApp. CAC is low. The missing piece is the conversion event. Conversion, here, is a change in a business metric. Until Indian trials measure that event, the programme is an MVP with strong usability and a thin revenue story. That is still a good MVP. It is not a Series B narrative.

The same thought–feeling–action loop adds one warning. People who have just practised optimism will rate their own performance more kindly. Tuesday’s rep makes Wednesday’s self-rating a shade warmer. That is not fraud. It is the instrument talking to the intervention. Other-raters and registers break that loop. They are scarce.

7.2 Limitations

This synthesis did not sit inside every subscription interface with a saved search alert and a dual-independent PRISMA count. Hit counts are therefore not manufactured. Some full texts remain paywalled; where means were not visible, effects were taken from abstracts, open tables, or author-reported d and flagged. Sun et al. (2026) is recent; its very large effects have not been stress-tested by replication or by an independent statistician with the raw file. Zhang’s control-group retest SD is anomalous. Da’s abstract over-reaches its

ANCOVA. Luthans 2010's manager study is still quoted as if it were an RCT. Hodges is a dissertation with a ceiling. Adjacent coaching trials from the Spanish group are better designed than their sample size. India is represented by one PCI RCT that did not measure this paper's business outcomes. GRADE is low or empty where it matters most to a CHRO. That is a limitation of the evidence, not of the tone.

7.3 What this means for an Indian workplace programme

Picture the business case that reaches the CFO of a Chennai services firm. It names PsyCap practice as a trainable resource. It does not promise a cut in attrition. The plan reads like a skill sprint: short sessions, a daily rep on the phone before the morning stand-up, in the language the team actually thinks in, and a human coach who picks up when the chatbot runs out of road. The leading indicator is the PsyCap questionnaire; if the prior pools travel, it will probably move by a third to a half of a standard deviation. There is no rupee saving from absenteeism in the case. That trial has not been run, anywhere. There is no turnover percentage either, unless HR will count actual exits over quarters rather than read a six-item intention scale after six weeks. For performance, the team lead's rating or the ticket-closure count beats a self-rating collected by the same app that delivered the exercises. Each week's reps come with one stretch task in the real job — the efficacy piece Loghman found fading when the room empties. On the Chennai floor the sessions run in Tamil, not only in English. When the quarters are counted, the firm publishes the result. India needs that paper more than it needs another correlational r.

7.4 Research gaps, with the India gap named

The India gap is not a polite afterthought. It is the hole in the map. Patnaik, Mishra and Mishra (2022) showed that an Indian telecom sample can be randomised to a PsyCap programme and will show movement on PsyCap, stress and job insecurity. That is necessary. It is not sufficient. No verified Indian randomised trial of a PsyCap programme has reported job performance, absenteeism or turnover. Bondre and colleagues (2025) ran a character-strengths coaching pilot with rural ASHAs in Madhya Pradesh and found a happiness d of 0.55; that is adjacent positive psychology, not a PsyCap programme, and not a business metric. A mindfulness-programme trial in Indian IT staff moved PsyCap as an outcome, which is the Loghman plot twist in miniature: non-PCI programmes can raise the bundle. None of these close the gap this review named.

Other gaps are global. Absenteeism needs a trial with the HR register, not a recall item. Actual turnover needs a twelve-month window and a sample large enough to count exits. Objective performance needs a metric the firm already trusts. Self-rated versus other-rated performance should be dual-primary, not an afterthought. Micro versus longer programmes should be a head-to-head on a business outcome, not only on the bundle. Delivery in Indian languages other than English should be a design factor, not a translation footnote. Pre-registration would help. So would a statistician who is not also the facilitator.

A two-hour practice is a small start. A small start is not a result. Nobody judges a new branch office by its first week at the counter. Indian workplaces that run the practice and then keep counting — performance, days away, exits, quarter

after quarter — will write the paper this one could not pool.

08

FAQ

Q1 Do PsyCap programmes improve job performance?

Self-rated performance has one controlled Chinese trial with a modest raw g of 0.18 and a larger ANCOVA effect (partial $\eta^2 = 0.20$). Other-rated evidence is weaker. Objective output in a randomised PsyCap programme was not found. Treat “performance” claims as low-certainty.

Q2 Will this cut our absenteeism?

No eligible PsyCap-intervention trial measured absenteeism. Zero. A 2006 field study found that people with higher PsyCap took fewer days off; that is correlation. Do not put an absence saving in a budget on this evidence.

Q3 Does it reduce turnover?

Actual resignations were not measured. Turnover intention fell hard in one Chinese nursing RCT ($g \approx -2.21$ at post, $n = 122$) and did not fall reliably in a second Chinese online RCT ($n = 104$). Two trials, two stories. Certainty is low.

Q4 How big is the effect on PsyCap itself?

Prior reviews: $d = 0.34$ across 41 mixed trials (Lupşa, 2020) and $d = 0.46$ post-intervention across 40 studies (Loghman, 2025). That is a small-to-medium lift on the bundle. Efficacy is the piece most likely to fade without later reps.

Q5 Has this been tested with Indian employees?

One telecom PCI randomised trial ($n = 234$) raised PsyCap and reduced stress and job insecurity. It did not measure performance, absenteeism or turnover. That is the India gap. It is still open.

09

References

- Avey, J. B., Patera, J. L., & West, B. J. (2006). The implications of positive psychological capital on employee absenteeism. *Journal of Leadership & Organizational Studies*, 13(2), 42–60. <https://doi.org/10.1177/10717919070130020401>
- Avey, J. B., Reichard, R. J., Luthans, F., & Mhatre, K. H. (2011). Meta-analysis of the impact of positive psychological capital on employee attitudes, behaviors, and performance. *Human Resource Development Quarterly*, 22(2), 127–152. <https://doi.org/10.1002/hrdq.20070>
- Bondre, A. P., et al. (2025). Evaluation of a positive psychological intervention to reduce work stress among rural community health workers in India: Results from a randomized pilot study. *Journal of Happiness Studies*, 26, 27. <https://doi.org/10.1007/s10902-024-00852-6>
- Carter, J. W., & Youssef-Morgan, C. M. (2022). Psychological capital development effectiveness of face-to-face, online, and micro-learning interventions. *Education and Information Technologies*, 27, 6553–6575. <https://doi.org/10.1007/s10639-021-10824-5>
- Da, S., He, Y., & Zhang, X. (2020). Effectiveness of psychological capital intervention and its influence on work-related attitudes: Daily online self-learning method and randomized controlled trial design. *International Journal of Environmental Research and Public Health*, 17(23), 8754. <https://doi.org/10.3390/ijerph17238754>
- Dello Russo, S., & Stoykova, P. (2015). Psychological capital intervention (PCI): A replication and extension. *Human Resource Development Quarterly*, 26(3), 329–347. <https://doi.org/10.1002/hrdq.21212>
- Hodges, T. D. (2010). An experimental study of the impact of psychological capital on performance, engagement, and the contagion effect (Doctoral dissertation, University of Nebraska–Lincoln). <https://digitalcommons.unl.edu/dissertations/AAI3398191/>
- Loghman, S., Quinn, M., Dawkins, S., Woods, M., Om Sharma, S., & Scott, J. (2023). A comprehensive meta-analyses of the nomological network of psychological capital (PsyCap). *Journal of Leadership & Organizational Studies*, 30(1), 108–128. <https://doi.org/10.1177/15480518221107998>
- Loghman, S., Ramirez-Perez, M., Bohle, P., & Martin, A. (2025). A comprehensive meta-analysis of the impact of intervention programmes on psychological capital development: Post-intervention and longer-term effects. *Personnel Review*, 54(1), 106–129. <https://doi.org/10.1108/PR-10-2023-0854>
- Lupşa, D., Vîrga, D., Maricuţoiu, L. P., & Rusu, A. (2020). Increasing psychological capital: A pre-registered meta-analysis of controlled interventions. *Applied Psychology*, 69(4), 1506–1556. <https://doi.org/10.1111/apps.12219>
- Luthans, F., Avey, J. B., Avolio, B. J., Norman, S. M., & Combs, G. M. (2006). Psychological capital development: Toward a micro-intervention. *Journal of Organizational Behavior*, 27(3), 387–393. <https://doi.org/10.1002/job.373>
- Luthans, F., Avey, J. B., Avolio, B. J., & Peterson, S. J. (2010). The development and resulting performance impact of positive psychological capital. *Human Resource Development Quarterly*, 21(1), 41–67. <https://doi.org/10.1002/hrdq.20034>
- Luthans, F., Avey, J. B., & Patera, J. L. (2008). Experimental analysis of a web-based training intervention to develop positive psychological capital. *Academy of Management Learning & Education*, 7(2), 209–221. <https://doi.org/10.5465/amle.2008.32712618>
- Patnaik, D., Mishra, M., & Mishra, S. (2022). Can psychological capital reduce stress and job insecurity? An experimental examination with Indian evidence. *Asia Pacific Journal of Management*. <https://doi.org/10.1007/s10490-021-09761-1>

Peláez, M. J., Coó, C., & Salanova, M. (2020). Facilitating work engagement and performance through strengths-based micro-coaching: A controlled trial study. *Journal of Happiness Studies*, 21, 1265–1284. <https://doi.org/10.1007/s10902-019-00127-5>

Peláez Zuberbuhler, M. J., Salanova, M., & Martínez, I. M. (2020). Coaching-based leadership intervention program: A controlled trial study. *Frontiers in Psychology*, 10, 3066. <https://doi.org/10.3389/fpsyg.2019.03066>

Peterson, S. J., Luthans, F., Avolio, B. J., Walumbwa, F. O., & Zhang, Z. (2011). Psychological capital and employee performance: A latent growth modeling approach. *Personnel Psychology*, 64(2), 427–450. <https://doi.org/10.1111/j.1744-6570.2011.01215.x>

Sun, K., Gong, X., & Yuan, G. (2026). Effects of a psychological capital intervention on emergency nurses’ work engagement and turnover intention: A randomized controlled trial. *Medicine*, 105(21), e48971. <https://doi.org/10.1097/MD.00000000000048971>

Zhang, X., Li, Y.-L., Ma, S., Hu, J., & Jiang, L. (2014). A structured reading materials-based intervention program to develop the psychological capital of Chinese employees. *Social Behavior and Personality*, 42(3), 503–515. <https://doi.org/10.2224/sbp.2014.42.3.503>

About the author

ABOUT THE AUTHOR

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

Atul Aswani is a Mumbai Psychiatrist and trained as a Results Coach. He builds EmotionApp, a non-clinical emotional-fitness platform for Indian professionals: countable daily reps, an AI coach with human hand-off, WhatsApp-first delivery, 23 Indian languages, India-resident data, DPDP-compliant, designed under ISO 9001 and ISO 13485. This paper is a practice- and-programme review, not a clinical guideline.

HOW TO CITE

Aswani A. Psychological Capital, Trained: A Meta-Analysis of PsyCap Interventions and Business Outcomes. EmotionApp Research Paper 29. Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/psycap-interventions-business-outcomes>

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.