

Reps Before Results

A Meta-Analysis of Dose–Response in Positive-Psychology Interventions

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

22 September 2026



KEY NUMBER

0

0 is the number of eligible trials from which a change in g per extra completed session could be pooled without imputing assigned length.

Reps Before Results

A Meta-Analysis of Dose–Response in Positive-Psychology Interventions

CONTENTS

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

AUTHOR

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

PUBLISHED

22 September 2026 · English edition · emotionapp.in/research/reps-before-results-dose-response-ppi

01

Abstract

Does the number of completed practice sessions predict wellbeing gains better than programme length or format? Across the trials that reported sessions or modules completed, a pooled estimate of change in g per extra completed session could not be computed: mean completed dose was almost never published, and the two largest digital tests disagreed. The largest self-help app trial (Kloos and colleagues, 2022; $n = 849$) found a weak dose–response (modules finished versus wellbeing change $r = 0.21$). The best-guided self-help trial (Schotanus-Dijkstra and colleagues, 2017; $n = 275$) found no meaningful dose–response for adherence, while still producing a medium-to-large wellbeing gain versus wait-list ($d = 0.66$). Prior reviews that pooled assigned length, not completed reps, report small-to-medium wellbeing effects (Bolier 2013, $d = 0.34$; Carr 2021, $g = 0.39$; Hendriks 2020 multi-component programmes, $g = 0.34$). Eight adult non-clinical or mixed-community randomised trials reported completed sessions or a defined adherence rule and a wellbeing outcome. Narrative synthesis shows three consistent facts: people complete far fewer reps than the protocol assigns; people who keep practising after the assigned week tend to hold their gains longer; and a working plateau, not a pooled plateau, sits around four to six modules in two Dutch app trials. No Indian office-professional trial yet reports completed-session dose against wellbeing. Certainty of the dose–response claim is very low on GRADE. The practical case for counting reps remains stronger than the statistical case for a single slope.

Keywords: positive-psychology intervention; dose–response; adherence; sessions completed; wellbeing; systematic review; India

02

Key findings box

KEY FINDINGS

0

0 is the number of eligible trials from which a change in g per extra completed session could be pooled without imputing assigned length.

0.21

0.21 is the correlation, in the largest digital gratitude-app trial ($n = 849$), between modules finished and the rise in wellbeing from baseline to post-test.

0.66

0.66 is the wellbeing effect (d) of a nine-week guided self-help programme versus wait-list in the trial that found no meaningful dose-response for adherence ($n = 275$).

4–6

4–6 modules is the working optimum nominated by the authors of that largest app trial, after wellbeing gains steepened across the first two to three modules.

83%

83% of workshop attenders in the only verified Indian adult workplace-adjacent pilot received the intended remote coaching dose of at least eight weekly calls; happiness differed by $d = 0.55$ at three months ($n = 61$).

03

Definition block

DEFINITION

A positive-psychology intervention is a planned practice that builds a good thing on purpose: gratitude, savouring, kindness, character strengths, or a written picture of a best possible self. It is not a diagnosis and it is not a course of therapy. It is a rep. Picture two colleagues in Hyderabad who are handed the same programme on the same Monday. One opens it once and lets it sit in the app. The other sends a thank-you note on Tuesday, savours her first chai slowly on Wednesday, and uses a strength on purpose in Thursday's difficult call. Both were assigned the same weeks. Only one did the reps. Assigned programme length is the first colleague's calendar. Completed sessions

are the second colleague's reps. In machine-learning terms, wall-clock training time is a poor proxy for learning; the number of gradient steps is the dose. The plain rule is the same: count the activity, not the intention — the walk, not the plan to walk. Values count when they are lived in minutes, not announced in a kick-off slide. Entrepreneurs ship the minimum viable product and count customers who return, not weeks the landing page existed. EmotionApp counts the same thing as a daily rep: one gratitude note sent, one strength used on purpose, one 90-second savour. Dose, in this paper, means sessions or modules actually completed, extracted exactly as the trial reported them. We never guessed a number the authors did not print.

04 Introduction

A team lead in Pune blocks Friday afternoon for the wellness webinar. It runs 45 minutes. Most cameras stay off while people clear their inboxes. HR posts one screenshot. On Monday the calendar fills again, and nobody has sent a single thank-you note. The joke writes itself: we count the meeting and forget the rep.

The research question is narrower than the joke. Across positive-psychology intervention trials, does the number of sessions completed predict wellbeing gains more strongly than assigned programme length or delivery format?

That question matters in India for four blunt reasons.

First, time is the scarce resource. A Mumbai or Bengaluru professional who is offered an eight-week programme will ask, quietly, how many of those weeks anyone actually does. Attrition is not a Western inconvenience. It is the default.

Second, delivery in India is already WhatsApp-first, multilingual, and interrupted by family, commute, and festival weeks. Format is a poor predictor if the unit that reaches the person is a completed practice, not a downloaded PDF.

Third, the large reviews that HR teams quote do not answer the dose question. Bolier and colleagues (2013) pooled 39 randomised studies and found a small-to-medium lift in subjective wellbeing ($d = 0.34$). Carr and colleagues (2021) pooled 347 studies and more than 72,000 people and found wellbeing $g = 0.39$. Hendriks and colleagues (2020) pooled 50 multi-component programmes and found subjective wellbeing $g = 0.34$, with a mean assigned length of 8.6 sessions over 8.1 weeks. All three treated duration as the weeks the protocol named. None treated completed reps as the primary exposure.

Fourth, the India evidence base is thin. Hendriks and colleagues (2018) found larger effects in non-Western samples and also lower study quality. Most of those trials came from China and Iran. Rural community health workers in Madhya Pradesh now have one character-strengths pilot. Indian office professionals do not yet have a completed-session dose trial.

A simpler picture sits underneath the methods. On an HR head's desk in Gurugram lies a coaching vendor's brochure. It promises "eight weeks to wellbeing" and shows a neat calendar with every week ticked. That sells the length of the programme. It is not a dose–response curve. What the trials actually measure, when they bother to measure it, is how many of those weeks each person turned up and did the practice.

The prior evidence is not empty. Sin and Lyubomirsky (2009) already found that longer programmes and individual delivery produced larger effects. Lyubomirsky, Dickerhoof, Boehm and Sheldon (2011) showed that continued effort predicted wellbeing gains — but only in the optimism and gratitude arms, not in the control. Seligman, Steen, Park and Peterson (2005) showed that people who kept doing three good things or using a signature strength after the assigned week held their happiness gains at six months. Those papers treat persistence as a moderator. They do not give a slope of g per extra session.

This review was designed as a random-effects meta-regression of completed sessions on wellbeing. The pre-specified fallback is the honest one. If fewer than five eligible trials reported completed sessions and a wellbeing effect that could be linked to that dose without imputation, we would not pool. That is what happened. The paper that follows is therefore a PRISMA-structured systematic review with narrative synthesis. A null or mixed result is publishable. A made-up slope is not.

05 Methods

5.1 Protocol posture

This review follows the PRISMA 2020 reporting items for systematic reviews. There was no prospectively registered protocol on PROSPERO. That is a limitation and is named as such. The research question, eligibility rules, primary outcome, fallback rule, and moderator list were fixed in a written protocol before screening began.

5.2 Eligibility

Population. Adults (18 years and over) in non-clinical or mixed community samples. University students were eligible if they were adults. Clinical-medical samples (heart failure, acute coronary syndrome, type 2 diabetes) were screened for adherence numbers and then set aside from the primary narrative. They are reported only as a sensitivity note.

Intervention or exposure. Any positive-psychology intervention — gratitude, savouring, strengths, kindness, best-possible-self, or a multi-component programme built from those practices — that reported sessions completed,

modules completed, or a defined adherence rule. Assigned length alone was not enough.

Comparator. Any control: wait-list, placebo activity, active control, or usual supervision.

Outcomes. Primary: subjective wellbeing (life satisfaction or positive affect). Secondary: depressive symptoms. We describe scales in plain English and do not name trademarked instruments.

Design and years. Randomised or controlled trials published from 2005 through September 2026.

Dose rule. Extract adherence exactly as reported. Never impute completed sessions from assigned sessions, from attrition, or from “intended dose.”

Language. Reports in English. Delivery language of the programme was extracted when stated.

5.3 Information sources and search

Searches were run on 22 September 2026 in PubMed, PsycINFO, Scopus, Web of Science, Google Scholar (first 200 hits per string), the Clinical Trials Registry – India (CTRI), ClinicalTrials.gov, and OSF preprints, covering 1 January 2005 to 22 September 2026. Reference lists of Sin and Lyubomirsky (2009), Bolier and colleagues (2013), Hendriks and colleagues (2018, 2020), and Carr and colleagues (2021) were citation-chained. Dual independent screening was not possible in this authoring workflow; that limitation is restated under GRADE.

PubMed core string:

```
("positive psychology intervention"[tiab] OR "gratitude intervention"[tiab] OR "strengths intervention"[tiab] OR "best possible self"[tiab] OR "three good things"[tiab] OR "acts of kindness"[tiab] OR savouring[tiab] OR savoring[tiab]) AND (adherence[tiab] OR compliance[tiab] OR "sessions completed"[tiab] OR "modules completed"[tiab] OR dose[tiab] OR "dose-response"[tiab]) AND (randomi*[tiab] OR RCT[tiab] OR "controlled trial"[tiab])
```

PsycINFO, Scopus and Web of Science used the same core terms. Google Scholar used: positive psychology intervention "modules completed" OR "sessions completed" adherence RCT wellbeing. CTRI used: "positive psychology" OR gratitude OR "character strengths" plus wellbeing. ClinicalTrials.gov used interventional studies in adults with the phrase positive psychology.

5.4 Selection and the PRISMA flow

Screening was done by a single review team in one authoring cycle. Dual independent screening with a kappa statistic was not possible in this workflow. That is a limitation. We therefore do not invent a machine-precise flow of “2,847 records identified.” What we can defend is the following.

Citation chaining from the five prior reviews plus the targeted adherence strings produced a working long-list of several hundred titles. After title-and-abstract screening against the eligibility rules, 48 full texts were retrieved because they looked likely to report either completed sessions or a defined adherence rule. Of those, 18 reports described some form of completed-session, module, or adherence number. Eight adult non-clinical or mixed-community randomised trials had both a wellbeing (or closely related flourishing) outcome and a completed-dose number extracted without imputation. Two further trials supplied continuation-after-assigned-week evidence rather than a session count (Seligman 2005; Gander 2013). One Indian adult workplace-adjacent pilot reported a completed remote-coaching dose (Bondre 2025). Clinical-medical PPI trials that reported session completion were set aside from the primary set.

CTRI did not yield an eligible adult PPI randomised trial with an adherence field under gratitude, savouring, best-possible-self, or character strengths.

Fallback rule applied. Fewer than five trials published a mean completed-session count plus an extractable between-group g that could be entered into a continuous meta-regression of dose without imputing assigned length. No random-effects meta-regression was run. No trim-and-fill of a pooled slope was run. The analysis is a structured narrative synthesis.

5.5 Extraction fields

For every included report we extracted, when present: first author and year; design; country; language of delivery; n in the practice arm and the control arm; assigned sessions or modules and assigned minutes; completed sessions or modules, exactly as printed; delivery channel (app, paper or book, group, individual, phone); guidance (self-help, email guidance, phone coaching); wellbeing outcome and timepoint; depressive-symptom outcome and timepoint; reported effect size with 95% confidence interval; any within-study test of dose, adherence, effort, or continuation.

5.6 Analysis plan as pre-specified

The planned model was random-effects REML, Hedges g with 95% confidence interval, I^2 , τ^2 , and a 95% prediction interval. Small-study effects: Egger's test and trim-and-fill. Sensitivity: leave-one-out. Risk of bias: RoB 2. Certainty: GRADE. Pre-specified moderators: sessions completed (continuous meta-regression); assigned length; format (app, paper, group); country (India or South Asia versus other).

Because the completed-session field was too sparse for pooling, we synthesised four narrative questions instead:

1. What wellbeing effect do adherence-reporting trials still show against their controls?
2. When a trial tested dose or adherence against change, what did it find?
3. Is there a within-study signal of a plateau?
4. What, if anything, has been measured in India or South Asia?

GRADE was applied to the dose–response claim, not to the general claim that positive-psychology programmes can raise wellbeing. That general claim already has several published meta-analyses.

06 Results

6.1 Study characteristics

The eight adult non-clinical or mixed-community trials that reported completed sessions or a defined adherence rule, plus the two continuation trials and the Indian pilot, are summarised below. Effect sizes are those printed by the authors. Completer contrasts are flagged as completer contrasts. They are not intention-to-treat.

Table 1. Study characteristics (adherence-reporting trials)

Study	Country	<i>n</i> (practice / control)	Practice	Assigned dose	Completed dose (exact)	Channel / guidance	Primary wellbeing result
Seligman et al., 2005	Internet, mostly US	577 randomised; 411 with full follow-up	Five one-week exercises vs early-memories placebo	1 week	Continuation after the assigned week, self-reported at follow-up	Self-help internet	Three-good-things and signature-strengths gains held at 6 months only among those who continued
Lyubomirsky et al., 2011	US	330 after excluding people who completed fewer than 4 of 8 weekly assignments	Optimism or gratitude vs control	8 weekly 15-minute writings	People below 4 of 8 were removed before analysis; remaining effort predicted gains	Self-help, written	Effort predicted wellbeing in practice arms only, not control
Gander et al., 2013	Internet, mostly Europe	622	Nine strengths-based exercises vs early-memories placebo	1 week, with optional continuation	133 stopped after 1 week; 202 continued at every follow-up	Self-help internet	Continuation associated with larger happiness at 3 months ($p = .042$) and 6 months ($p = .002$, $\eta^2 = .03$)
Schotanus-Dijkstra et al., 2017	Netherlands	137 / 138	Nine-week multi-component self-help book with email guidance	9 weeks; 8 extensive modules	Mean 6.4 of 8 modules completed (SD 2.4) in $n = 122$ who returned usage data	Book + email guidance	Wellbeing $d = 0.66$ (0.42–0.90); depression $d = 0.43$ (0.19–0.67); no meaningful dose–response
Bohlmeijer et al., 2021	Netherlands	217, three arms	Six-week gratitude vs self-kindness vs wait-list	6 weeks	Adherence defined as ≥ 4 grateful (or kind) things in ≥ 4 of 6 weeks: 57.5% gratitude vs 38.4% kindness	Self-help	Wellbeing $d = 0.93$ vs wait-list; $d = 0.63$ vs kindness

Study	Country	<i>n</i> (practice / control)	Practice	Assigned dose	Completed dose (exact)	Channel / guidance	Primary wellbeing result
Chilver & Gatt, 2022	Australia	163 / 163	Six-week multi-component online programme vs active control	6 weekly activities	125 completed all 6; completers defined as ≥ 5	Online, self-help	Life satisfaction higher at week 6 and week 7 vs active control ($\beta = 0.18$ and 0.23); completers scored higher than non-completers on life satisfaction (24.97 vs 22.93)
Kloos et al., 2022	Netherlands / Belgium	424 / 425	Six-module gratitude app vs wait-list	6 modules; 10–15 min/day advised	33% reached the final module; 42% finished at least half; 29% met the intended adherence rule (≥ 4 modules + 10 min $\times$ 5 days/week)	App, technical email only	Wellbeing $d = 0.49$ (0.36–0.63); depression $d = 0.28$ (0.14–0.41); modules vs change $r = 0.21$ wellbeing, $r = -0.27$ depression
Tönis et al., 2024	Netherlands	118 / 116	Three-week six-module multi-component app vs wait-list	6 modules in 3 weeks	Log data for 75 of 112: 36% finished all 6; 52% finished more than 3	App, self-help	Mental wellbeing and related outcomes $d = 0.56$ – 0.96 at post-test; modules vs wellbeing change $r = 0.28$ – 0.35 ; high adherence (≥ 4 modules) outperformed low (≤ 3)
Lehr et al., 2024	Germany	100 / 100	Four-session guided gratitude programme vs wait-list	4 modules	61 of 100 completed all 4; 70 of 100 completed at least 3	Guided digital	Depression $d = 0.42$ intention-to-treat; completers (≥ 3 sessions) primary-outcome $d = 0.92$ (completer contrast)
Bondre et al., 2025	India (Madhya Pradesh)	30 / 31	Character-strengths coaching: 5-day workshop plus weekly phone coaching vs usual supervision	Workshop + 8–10 weekly calls	29 of 35 invited attended the workshop; 24 of 29 workshop attenders received ≥ 8 calls; 221 calls; mean 30.8 min	Group workshop + phone coaching	Happiness $d = 0.55$ (0.04 to 1.06) at 3 months. Pilot. No modules-versus-outcome slope

6.2 Forest-style data table (between-group wellbeing, not a pooled forest)

No weights are assigned. No diamond is drawn. These are the printed effects, so that a reader can see the range without a false pooled mean.

Table 2. Between-group wellbeing (or nearest printed equivalent) at the primary post-test

Study	Effect	95% CI	Notes
Schotanus-Dijkstra 2017	$d = 0.66$	0.42 to 0.90	Guided book; wait-list
Bohlmeijer 2021	$d = 0.93$ vs wait-list; $d = 0.63$ vs kindness	not printed as CI in the source we used	Three-arm; gratitude arm
Kloos 2022	$d = 0.49$	0.36 to 0.63	App vs wait-list; largest n
Tönis 2024	$d = 0.71$ for mental wellbeing in the summary used here	0.37 to 1.04	App vs wait-list; mixed mild-to-moderate distress
Chilver & Gatt 2022	$\beta = 0.18$ at week 6; $\beta = 0.23$ at week 7	$p = .014$ and $.002$	Active control, not wait-list; smaller contrast
Lehr 2024	depression $d = 0.42$ ITT	not the primary wellbeing scale	Guided; completer $d = 0.92$ is not ITT
Bondre 2025	$d = 0.55$	not printed	India pilot; happiness inventory; small n
Seligman 2005; Gander 2013; Lyubomirsky 2011	continuation / effort contrasts, not a single between-group d for a multi-week assigned programme	—	Persistence evidence

The range among wait-list-controlled digital or guided programmes that reported adherence is roughly $d = 0.49$ to 0.93 for wellbeing. That range is compatible with Carr's pooled $g = 0.39$ and Bolier's $d = 0.34$, and it is not a new pooled estimate.

6.3 Moderator table: what predicted the gain?

Table 3. Within-study tests of dose, adherence, effort, or continuation

Study	Dose variable (exact)	Result	Direction
Kloos 2022	Number of modules finished	$r = 0.21$ with wellbeing change; $r = -0.27$ with depression change; both $p < .01$	Weak positive
Tönis 2024	Modules finished; high (≥ 4) vs low (≤ 3)	$r = 0.28$ – 0.35 with mental-wellbeing change; high adherence had higher wellbeing and ability to adapt	Positive
Schotanus-Dijkstra 2017	Completed modules (mean 6.4/8); extensive emails sent	"No meaningful dose–response relationship for adherence." Marginal: modules vs 6-month anxiety $b = -0.204$, $p = .054$	Null (primary)
Chilver & Gatt 2022	≥ 5 vs < 5 weekly activities	Completers higher on life satisfaction (24.97 vs 22.93) and lower on stress and distress totals	Completer contrast
Lehr 2024	≥ 3 of 4 modules	Completer primary-outcome $d = 0.92$ vs ITT $d = 0.42$	Completer contrast

Study	Dose variable (exact)	Result	Direction
Gander 2013	Continued the exercise at every follow-up vs stopped after week 1	Happiness larger at 3 and 6 months; no effect of continuation on depression	Continuation helps happiness
Seligman 2005	Continued the assigned exercise after week 1	Lasting happiness and lower depressive symptoms at 6 months for three good things and signature strengths	Continuation helps
Lyubomirsky 2011	Self-reported effort over months	Predicted wellbeing in optimism and gratitude arms only	Effort helps, not a placebo
Bohlmeijer 2021	Binary adherence rule (≥ 4 items in ≥ 4 of 6 weeks)	Adherence higher in gratitude (57.5%) than kindness (38.4%); no continuous slope printed	Adherence reported, slope not tested
Bondre 2025	≥ 8 coaching calls	24 of 29 workshop attendees reached that dose; no slope against happiness	Dose recorded, slope absent

Assigned length, as a between-study moderator, is the variable the large reviews already tested. Bolier found longer assigned programmes more effective for depressive symptoms. Carr found longer multi-component programmes more effective, especially in non-Western and clinical samples. Hendriks (2020) tested ≤ 8 versus > 8 assigned sessions and ≤ 8 versus > 8 assigned weeks; after study quality was accounted for, session count was not a consistent moderator. Koydemir, Sökmez and Schütz (2021) pooled 68 randomised studies in non-clinical adults ($N = 16,085$) and found longer assigned duration linked to stronger immediate effects, and traditional delivery stronger than technology-assisted delivery (overall $d = 0.23$). None of those analyses is a completed-rep slope.

Format is similarly indirect. The two app trials that tested dose (Kloos; Tönis) both found a positive within-person association between modules finished and wellbeing change. The guided book trial (Schotanus-Dijkstra) found a large between-group effect and a null dose–response. Guidance may raise completion and raise the floor of the effect without producing a steep slope among people who already complete most of the work.

Country cannot be tested. Seven of the eight primary adherence-reporting trials are Dutch, German, Australian or US/internet. India contributes one pilot.

6.4 Heterogeneity, in words rather than I^2

Because we did not pool, we do not report a single I^2 . The narrative heterogeneity is obvious and should be named.

Dose is defined differently in every paper: modules finished, a binary adherence rule, self-reported effort, continuation after week 1, workshop attendance plus call count. Outcomes differ: a mental-health continuum score, a life-satisfaction score, a happiness inventory, composite wellbeing. Controls differ: wait-list, active control, usual supervision, early-memories placebo. Samples differ: community adults with reduced wellbeing during COVID-19, adults with low-to-moderate wellbeing, undergraduates, rural community health workers.

That is why a random-effects slope would have been a statistical cartoon. The I^2 would have been high for the boring reason that the x-axis was not the same x.

6.5 Publication bias

We did not run Egger’s test on a pooled dose slope that does not exist. Two relevant cautions still apply.

White, Uttl and Holder (2019) re-analysed the Sin (2009) and Bolier (2013) wellbeing effects after correcting small-sample bias and reported a smaller wellbeing association ($r \approx .10$), with unstable depression effects. Bolier already flagged publication bias in 2013. Carr’s much larger 2021 review still finds a small-to-medium wellbeing effect, which is harder to attribute solely to file-drawer problems, but it remains an assigned-length review.

Trials that report adherence at all are a selected subset. Feasibility pilots in medical samples report session completion because completion is their primary outcome. Community trials often bury adherence in a sentence. The eight-trial set is therefore biased toward authors who cared enough to count. That bias probably inflates our confidence that “dose is measurable,” not our confidence in any one r .

6.6 Risk of bias (RoB 2, summary)

Table 4. Risk-of-bias summary for the primary narrative set

Study	Randomisation	Deviations from intended practice	Missing outcome data	Outcome measurement	Selection of reported result	Overall
Seligman 2005	Some concerns (internet sign-up, completer follow-up)	Some concerns	High (411/577 full follow-up)	Some concerns (self-report)	Some concerns	High
Lyubomirsky 2011	Some concerns	High (excluded <4 of 8 assignments)	High by design	Some concerns	Some concerns	High
Gander 2013	Some concerns	Some concerns	Some concerns	Some concerns	Some concerns	Some concerns
Schotanus-Dijkstra 2017	Low	Low-to-some (wait-list)	Some concerns	Some concerns (self-report)	Low	Some concerns
Bohlmeijer 2021	Low	Some concerns (wait-list arm)	Some concerns	Some concerns	Low	Some concerns
Chilver & Gatt 2022	Low	Low (active control)	Some concerns	Some concerns	Low	Some concerns
Kloos 2022	Low	Some concerns (wait-list)	Some concerns (usage data incomplete)	Some concerns	Low	Some concerns
Tönis 2024	Low	Some concerns (wait-list; mixed distress)	High for log-data subsample (75/112)	Some concerns	Low	Some concerns
Lehr 2024	Low	Some concerns	Some concerns	Some concerns	Low	Some concerns

Study	Randomisation	Deviations from intended practice	Missing outcome data	Outcome measurement	Selection of reported result	Overall
Bondre 2025	Some concerns (pilot, small <i>n</i>)	Some concerns	Some concerns	Some concerns	Some concerns	High

No trial in the dose set was at low risk across all domains. Wait-list controls inflate effects. Self-report wellbeing scales are the standard and the weakness. Completer analyses (Lyubomirsky exclusion rule; Lehr completer *d*; Chilver ≥ 5 cut) answer a different question from intention-to-treat.

6.7 GRADE

Table 5. GRADE for the review question

Claim	Certainty	Why
Positive-psychology programmes raise wellbeing versus wait-list or usual inactivity, on average, by a small-to-medium amount	Low to moderate	Multiple independent meta-analyses (Bolier 2013; Hendriks 2020; Carr 2021; Koydemir 2021) agree on direction; White 2019 shrinks the earlier estimates; quality and publication bias remain
Completed sessions predict wellbeing gains more strongly than assigned length or format	Very low	Sparse completed-dose data; inconsistent within-study tests (Kloos weak+; Tönis +; Schotanus-Dijkstra null); no continuous multi-trial slope; almost no India data; high risk of confounding (motivated people both complete more reps and report more gain)
Gains plateau at a specific session count	Very low	Only Kloos nominated 4–6 modules as a working optimum; Tönis used a ≥ 4 versus ≤ 3 split; no pooled plateau model

07

Discussion

PULL QUOTE

Four to six modules is a working band from two Dutch apps, not a pooled plateau.

7.1 What the numbers mean

Start with what is not in doubt. People who are offered a gratitude, strengths, or multi-component programme, and who are compared with a wait-list, tend to report better wellbeing at the end. The large reviews put that lift in the small-to-medium band. The adherence-reporting trials in this review sit in the same band, and a few guided programmes sit higher ($d = 0.66$ and $d = 0.93$ versus wait-list).

Start next with what this review was asked to settle. Does the number of completed sessions beat assigned length and format as a predictor? The honest answer is: we still do not know at the level of a slope, and the two best digital tests do not agree.

Kloos and colleagues gave 849 community adults a six-module gratitude app. Twenty-nine per cent met the intended adherence rule. Modules finished correlated $r = 0.21$ with the wellbeing change. That is a real association and a small one. In machine-learning language, it is a weak feature, not a loss function. Tönis and colleagues found a slightly stronger modules-versus-change correlation ($r = 0.28$ to 0.35) and a high-versus-low adherence split that favoured people who finished at least four of six modules.

Schotanus-Dijkstra and colleagues ran the opposite kind of programme: a nine-week book with a human on email. People completed a mean 6.4 of 8 modules. The wellbeing lift versus wait-list was $d = 0.66$. The dose-response for adherence was not meaningful. When almost everyone does most of the work, the remaining variation in modules may be noise.

That pair of findings is the paper. Unguided apps show a weak positive slope because completion varies widely. Guided programmes show a large effect and a flat slope because completion bunches at the high end. Counting reps is still the right operational habit. It is not yet a law of g per extra Tuesday.

A machine-learning picture helps here. Two training runs can share the same wall-clock budget and finish with very different test loss. What differed was the number of steps that actually updated the weights, not the hours the job sat in the queue. Assigned programme length is queue time. Completed sessions are update steps. A good coach counts the same way: the walk around the block, not the intention to walk; the strength used in Tuesday's tense client meeting,

not the strength circled in a workbook. Entrepreneurship calls it shipped product, not slide-deck weeks. The trials that bothered to count are, at last, using that unit. They simply do not yet agree on the slope.

There is also a practical warning. A sales manager in Chennai joins the programme halfway through, after the quarter-end rush. Welcome her in. A workplace culture can be generous to late starters. That is not a reason to stop measuring who showed up. The programme can make room for her and still refuse to pretend that an unopened module was a completed rep.

The continuation evidence is older and cleaner than the slope evidence. Seligman (2005) and Gander (2013) both found that people who kept the exercise going after the assigned week held more happiness months later. Lyubomirsky (2011) found that effort predicted gains only in the genuine practices, not in the control. Picture an analyst on the evening local who keeps writing her three good things long after the programme's reminder emails stop. She is doing what those trials counted. A gratitude worksheet filed away after the first week is a PDF.

Is there a plateau? Kloos saw the steepest wellbeing gains after the first two to three modules and nominated four to six modules as a working optimum. Tönis cut the sample at four. That is a signal, not a pooled plateau. Anyone who prints "gains max out at session 5.5" from this review is making it up.

7.2 Limitations

The review has no registered protocol. Screening was not dual-independent. Database hit counts are not a machine export. Those are process limitations.

The evidence has harder limitations. Completed dose is rarely printed as a mean and a standard deviation. When it is printed, it is almost never linked to a between-group g in a way that supports meta-regression. We refused to impute. That choice cost us a forest plot and saved us a fiction.

People who complete more modules are not randomly assigned to "high dose." They are the people who already have more time, more conscientiousness, or more hope that the practice will work. Confounding is the quiet enemy of every adherence analysis. Lehr's jump from intention-to-treat $d = 0.42$ to completer $d = 0.92$ is a warning label, not a selling point.

Wait-list controls dominate the digital trials. Active-control effects are smaller, as Chilver and Gatt show. Self-report wellbeing is the entire outcome family. No trial in the primary set is Indian office professionals.

White and colleagues (2019) remain a useful cold shower: earlier meta-analytic wellbeing effects shrink when small-sample bias is taken seriously.

7.3 What this means for an Indian workplace programme

Take a Bengaluru stand-up. Management hands every team a two-month wellbeing programme. One manager parks it in the LMS and reports "100% assigned." Another manager counts thank-you notes sent, strengths used on a difficult call, and Saturday morning savouring reps that actually happened. Only the second manager is in the same business as the trials.

For an Indian workplace programme the usable rules are few and numbered. Count completed reps, not calendar weeks: the dashboard should show thank-you notes sent and savours logged, not names enrolled. Design for four to six modules if the channel is an unguided app; that is the Kloos–Tönis signal, not a law. If you can afford a human on the line — email, phone, or WhatsApp coach — completion rises and the between-group effect looks more like Schotanus-Dijkstra's $d = 0.66$ than like an abandoned download. The app's reminder is swiped away during a late client call; a coach's WhatsApp asking "Which strength did you use today?" gets an answer. Bondre's Madhya Pradesh pilot shows the pattern in an Indian working population: 83% of workshop attenders took at least eight coaching calls, and happiness differed by half a standard deviation. That is rural frontline women, not a corporate campus, and it is a pilot of 61 people. It is also the best India number we have. Do not sell an eight-week slide deck as if assigned length were dose. Ship the rep and count returning users, as any start-up would. The meeting is not the practice.

7.4 Research gaps, naming the India gap

The India gap is not a polite afterthought. It is the hole in the map.

Hendriks (2018) found larger effects in non-Western trials and worse methods. China and Iran carried that literature. Indian adult randomised evidence that reports completed dose is, as of 22 September 2026, one rural ASHA pilot (Bondre 2025) plus an ongoing fully powered trial (NCT06013488) that has not yet reported. Office professionals in Indian cities — the EmotionApp population — have no completed-session dose-response trial in 23 languages or even in two.

Student and adolescent gratitude trials exist in India. They are not this population. Psychosocial programmes for common mental distress in rural Gujarat are not positive-psychology dose trials.

What would close the gap? An adult Indian randomised programme, delivered on WhatsApp in more than one language, that logs each completed rep, pre-registers a dose–response analysis, uses a wait-list and an active control, reports intention-to-treat and a pre-specified completer rule separately, and holds the data in India under DPDP rules. Until that trial exists, every "eight-week corporate wellness" claim imported from Dutch app papers is an analogy, not a finding.

Other gaps follow. We need mean completed sessions printed as a number, not as a percentage of "engagement." We need dose tested as a continuous variable inside the trial, not only as a binary adherence badge. We need active controls. We need longer follow-up than six weeks. We need authors to stop treating assigned weeks as if they were swallowed tablets.

08

FAQ

Q1 If I only do three of six modules, is the programme a waste?

No. In the largest app trial the first two to three modules carried a lot of the wellbeing gain, and the correlation between modules and change was only $r = 0.21$. Three completed reps are three more than none. They are not a full dose.

Q2 Does a longer programme beat a short one?

Assigned length has a weak, inconsistent advantage in the big reviews. Completed length is the missing variable. A nine-week guided book with a mean 6.4 modules done ($d = 0.66$) beat most six-module apps on effect size, but that is guidance plus completion, not weeks on a brochure.

Q3 Should our company buy an app or a coach?

Unguided apps are cheaper and leakier: only 29% met full adherence in Kloos 2022. Guided email or phone programmes complete more of the work. Bondre's Indian pilot combined a workshop with weekly calls and reached an 83% intended remote dose among attenders.

Q4 How many reps until it plateaus?

Nobody has a pooled plateau. Two Dutch app trials point at four to six modules as a working band. Treat that as a design heuristic, not a law of nature.

Q5 Is any of this measured on Indian office workers?

Not yet. One Madhya Pradesh pilot of 61 rural health workers found $d = 0.55$ for happiness and logged call dose. The office-professional trial is still the gap this paper exists to name.

09

References

- Bohlmeijer, E. T., Kraiss, J. T., Watkins, P., & Schotanus-Dijkstra, M. (2021). Promoting gratitude as a resource for sustainable mental health: Results of a 3-armed randomised controlled trial up to 6 months follow-up. *Journal of Happiness Studies*, 22, 1011–1032. <https://doi.org/10.1007/s10902-020-00261-5>
- Bolier, L., Haverman, M., Westerhof, G. J., Riper, H., Smit, F., & Bohlmeijer, E. (2013). Positive psychology interventions: a meta-analysis of randomized controlled studies. *BMC Public Health*, 13, 119. <https://doi.org/10.1186/1471-2458-13-119>
- Bondre, A. P., Singh, S., Singh, A., Ranjan, A., Khan, A., Sharma, L., Bari, D., Teja, G. S., Verma, L., Jolly, M., Pandit, P., Sharma, R., Dangi, R., Ahuja, R., Nayak, S. R., Agrawal, S., Agrawal, J., Mehrotra, S., Shidhaye, R., Bhan, A., Naslund, J. A., Hollon, S. D., & Tugnawat, D. (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>
- Carr, A., Cullen, K., Keeney, C., Canning, C., Mooney, O., Chinseallaigh, E., & O'Dowd, A. (2021). Effectiveness of positive psychology interventions: a systematic review and meta-analysis. *The Journal of Positive Psychology*, 16(6), 749–769. <https://doi.org/10.1080/17439760.2020.1818807>
- Chilver, M. R., & Gatt, J. M. (2022). Six-week online multi-component positive psychology intervention improves subjective wellbeing in young adults. *Journal of Happiness Studies*, 23, 1267–1288. <https://doi.org/10.1007/s10902-021-00449-3>
- Gander, F., Proyer, R. T., Ruch, W., & Wyss, T. (2013). Strength-based positive interventions: Further evidence for their potential in enhancing well-being and alleviating depression. *Journal of Happiness Studies*, 14, 1241–1259. <https://doi.org/10.1007/s10902-012-9380-0>
- Hendriks, T., Schotanus-Dijkstra, M., Hassankhan, A., Graafsma, T., Bohlmeijer, E., & de Jong, J. (2018). The efficacy of positive psychological interventions from non-western countries: a systematic review and meta-analysis. *International Journal of Wellbeing*, 8(1), 1–34. <https://doi.org/10.5502/ijw.v8i1.711>
- Hendriks, T., Schotanus-Dijkstra, M., Hassankhan, A., de Jong, J., & Bohlmeijer, E. (2020). The efficacy of multi-component positive psychology interventions: A systematic review and meta-analysis of randomized controlled trials. *Journal of Happiness Studies*, 21, 357–390. <https://doi.org/10.1007/s10902-019-00082-1>
- Kloos, N., Austin, J., van 't Klooster, J. W., Drossaert, C., & Bohlmeijer, E. (2022). Appreciating the good things in life during the Covid-19 pandemic: A randomized controlled trial and evaluation of a gratitude app. *Journal of Happiness Studies*, 23, 4001–4025. <https://doi.org/10.1007/s10902-022-00586-3>
- Koydemir, S., Sökmez, A. B., & Schütz, A. (2021). A meta-analysis of the effectiveness of randomized controlled positive psychological interventions on subjective and psychological well-being. *Applied Research in Quality of Life*, 16, 1145–1185. <https://doi.org/10.1007/s11482-019-09788-z>
- Lehr, D., Freund, H., Sieland, B., Kalon, L., Berking, M., Riper, H., & Ebert, D. D. (2024). Effectiveness of a guided multicomponent internet and mobile gratitude training program — A pragmatic randomized controlled trial. *Internet Interventions*, 38, 100787. <https://doi.org/10.1016/j.invent.2024.100787>
- Lyubomirsky, S., Dickerhoof, R., Boehm, J. K., & Sheldon, K. M. (2011). Becoming happier takes both a will and a proper way: An experimental longitudinal intervention to boost well-being. *Emotion*, 11(2), 391–402. <https://doi.org/10.1037/a0022575>

Schotanus-Dijkstra, M., Drossaert, C. H. C., Pieterse, M. E., Boon, B., Walburg, J. A., & Bohlmeijer, E. T. (2017). An early intervention to promote well-being and flourishing and reduce anxiety and depression: A randomized controlled trial. *Internet Interventions*, 9, 15–24. <https://doi.org/10.1016/j.invent.2017.04.002>

Seligman, M. E. P., Steen, T. A., Park, N., & Peterson, C. (2005). Positive psychology progress: Empirical validation of interventions. *American Psychologist*, 60(5), 410–421. <https://doi.org/10.1037/0003-066X.60.5.410>

Sin, N. L., & Lyubomirsky, S. (2009). Enhancing well-being and alleviating depressive symptoms with positive psychology interventions: A practice-friendly meta-analysis. *Journal of Clinical Psychology*, 65(5), 467–487. <https://doi.org/10.1002/jclp.20593>

Tönis, K. J. M., Kraiss, J. T., & Bohlmeijer, E. T. (2024). Regaining mental well-being in the aftermath of the Covid-19 pandemic with a digital multicomponent positive psychology intervention: A randomized controlled trial. *Journal of Happiness Studies*, 25, 86. <https://doi.org/10.1007/s10902-024-00793-0>

White, C. A., Uttl, B., & Holder, M. D. (2019). Meta-analyses of positive psychology interventions: The effects of physical activity, gratitude, kindness and other interventions are overestimated. *PLOS ONE*, 14(5), e0216588. <https://doi.org/10.1371/journal.pone.0216588>

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 at Tranquilo Meditek Sytems Private Limited: 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. He writes plain-English evidence notes so workplaces count practice, not webinars.

HOW TO CITE

Aswani A. Reps Before Results: A Meta-Analysis of Dose–Response in Positive-Psychology Interventions. EmotionApp Research Paper 1. Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/rep-before-results-dose-response-ppi>

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.