

Best Possible Self, Best Available Evidence

A Meta-Analysis of Future-Self Writing and Its Dose

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



KEY NUMBER

0.33

95% CI 0.189 to 0.461

0.33 — Carrillo's pooled effect on wellbeing ($d = 0.325$, 95% CI 0.189 to 0.461; 29 studies, $N = 2,909$). That is the number most people quote. It is a small-to-moderate lift, not a personality transplant.

Best Possible Self, Best Available Evidence

A Meta-Analysis of Future-Self Writing and Its Dose

CONTENTS

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

INDEXING TITLE

Best-possible-self writing and its dose: an updated meta-analysis

AUTHOR

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

PUBLISHED

22 September 2026 · English edition · emotionapp.in/research/best-possible-self-dose-meta-analysis

01

Abstract

Carrillo and colleagues pooled the best-possible-self exercise in 2019. Across updated randomised and controlled trials, best-possible-self practice produces a small-to-moderate effect on wellbeing (Carrillo $d+ = 0.33$), optimism ($d+ = 0.33$; Heekerens $g = 0.21$) and positive affect ($d+ = 0.51$; Heekerens $g = 0.28$); number of sittings is not a confirmed multiplier of that effect. This review asked a narrower question than Carrillo asked. Does one sitting change the effect relative to three or four sittings? Adults who wrote or imagined a best possible future self were compared with any control. Primary outcomes were optimism, positive affect and wellbeing. Searches covered PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI, ClinicalTrials.gov and OSF preprints from 2001 to 22 September 2026. Carrillo included 29 studies in 26 articles ($N = 2,909$). Heekerens and Eid later pooled 34 randomised trials ($N = 3,675$). An update to 2026 added further adult trials that isolate the exercise, including one-sitting laboratory and online tests, four-week writing packages, an app trial and a sittings-randomised anxiety study. Random-effects synthesis of the prior pools, plus sitting-coded narrative of the update, was planned a priori. One sitting reliably lifts same-day positive affect and positive future expectancies. Those mood effects fade within about a week. Three-to-four sittings is Laura King's original wellbeing protocol. It is associated with wellbeing at weeks, not with a proven larger g than one sitting on the same outcome at the same hour. Writing and imagery are equivalent in the one sitting that tested them head to head. Follow-up effects are small or absent. Certainty is moderate for immediate affect and expectancies, low for trait optimism and later wellbeing, and low for the sittings contrast itself. No adult Indian trial was found. For a workplace programme the honest product is one 15-to-20-minute rep for a same-day lift, and three further reps across a week if the aim is a wellbeing habit rather than a larger first-hour number.

Keywords: best possible self; optimism; positive affect; wellbeing; dose; sittings; meta-analysis; India

02

Key findings box

KEY FINDINGS

0.33

0.33 — Carrillo's pooled effect on wellbeing ($d+ = 0.325$, 95% CI 0.189 to 0.461; 29 studies, $N = 2,909$). That is the number most people quote. It is a small-to-moderate lift, not a personality transplant.

0.28

0.28 — Heckerens and Eid's Hedges g for positive affect across 34 randomised trials (95% CI 0.16 to 0.41). On the day of the exercise the same team's estimate rises to $g = 0.39$. A week later the follow-up effect is not substantial.

0.21

0.21 — Heckerens and Eid's Hedges g for optimism (95% CI 0.04 to 0.38). The effect is real when optimism means positive future expectancies. It is weak or absent when optimism means a standing trait score.

1

1 sitting versus 3–4 sittings — one sitting produces the same-day affect and expectancy lift. Three-to-four sittings is King's original wellbeing protocol. No stable pooled g yet shows that three-to-four sittings enlarge the same-hour effect. Carrillo's minutes moderator ran the other way: shorter total practice, larger wellbeing (trend, $p = .078$).

0

0 adult Indian randomised trials — CTRI, ClinicalTrials.gov, PubMed, PsycINFO and OSF yielded no best-possible-self trial in Indian working adults. The India gap is not a footnote. It is the main limit on how far any Western g should travel into an Indian office.

03

Definition block

DEFINITION

Best-possible-self writing is a structured future-self exercise. You pick a point in the future. You imagine that life has gone as well as it honestly could, after you have worked for it. You write that life in concrete detail — work, relationships, health, character, ordinary Tuesdays — for about 15 to 20 minutes. Many protocols then add five minutes of mental imagery of the same scene. The exercise is a practice, a technique, a countable daily or weekly rep. It is not a diagnosis and it is not a clinical programme. Laura King introduced the four-day writing version in 2001 as a cousin of expressive writing about trauma. Later laboratories shortened it to a single sitting and used it as an optimism induction. Both versions ask the same question: if the future you are building actually arrived, what would it look like from the inside?

04

Introduction

A team lead in Pune has one gap between the stand-up and the school pickup. The resilience webinar sits unwatched in a tab. She needs a rep she can finish in that gap.

The best-possible-self exercise looks, on paper, like the kind of thing a stand-up comic would mock. “Imagine your best life. Write it down. Feel better.” It sounds like a greeting-card algorithm. Then you notice that the same joke has survived a quarter-century of randomised tests. That is not nothing. In machine-learning terms, the field has run the same architecture through many datasets. The weights move. They do not explode. The question left on the table is not “does the model train?” The question is “how many epochs do you need before the loss stops falling?”

Carrillo, Rubio-Aparicio, Molinari, Enrique, Sánchez-Meca and Baños answered the first question in 2019. They pooled 29 studies in 26 articles and 2,909 participants. Best-possible-self practice improved wellbeing ($d+ = 0.325$), optimism ($d+ = 0.334$) and positive affect ($d+ = 0.511$) against controls. Negative affect and depressive symptoms moved little. Moderator tests for wellbeing were mostly quiet. Two trends whispered. Older participants gained more. Shorter practices, counted as total minutes, also gained more.

That second whisper is the plot of this paper.

King's original protocol was not one sitting. It was four. Eighty-one undergraduates wrote for 20 minutes a day on four consecutive days. One arm wrote about trauma. One wrote about a best possible future self. One wrote both. One wrote about a dry control topic. The future-self writers were less upset than the trauma writers and showed higher subjective wellbeing three weeks later. Five months later the future-self, trauma, and combined arms had fewer illness visits than controls. The field then did what fields do. It compressed the protocol. Peters and colleagues showed that one sitting of 15 minutes of writing plus five minutes of imagery raised positive future expectancies the same afternoon. A Maastricht series repeated the one-sitting version until it became the laboratory default. Loveday, Lovell and Jones, reviewing 31 studies in 2018, wrote that repeating the activity “appears to enhance efficacy” and called for dosage trials. Nobody ran the clean experiment the call implied: randomise one sitting versus three-to-four sittings on optimism and wellbeing in adults, and hold the clock constant.

So the field now has two products wearing the same name. Product A is a 20-minute mood-and-expectancy induction. Product B is a four-sitting wellbeing protocol. EmotionApp ships the second product to Indian professionals as a WhatsApp-first coaching rep. The company needs to know whether the extra sittings earn their minutes.

The question matters in India for four ordinary reasons.

First, time is the scarce resource. A Bengaluru product manager will give a programme 20 minutes once. She may not give it 20 minutes on four consecutive nights. If one sitting does the job, the extra three sittings are overhead. If one sitting is only a spark, the extra three sittings are the product.

Second, optimism in an Indian office is not a poster. It is a working stance toward uncertain deals, delayed appraisals and family duty that does not pause for a journal prompt. A sales manager in Gurugram writes that, a year from now, she leads the regional team. Then she picks one small step that belongs in that scene: she asks to present the team's numbers at Friday's review. She puts it in her calendar before she closes the notebook. A future that stays on the page is a daydream. A future that gets a calendar square is a practice.

Third, most of the trials were run on Western undergraduates. Carrillo's mean age was about 24. Three-quarters of participants were women. Delivery languages were English, Dutch, Spanish, German, and later Chinese. Hindi, Tamil, Telugu, Bengali and Marathi do not appear. A 2024 mainland-China sitting found no rise in positive affect or wellbeing. A 2025 Chinese–Dutch comparison found larger positive-affect and optimism gains in the Dutch sample, with Chinese gains arriving later and landing more on negative outcomes. If culture moves the effect, India is not a footnote. India is an unrun trial.

Fourth, a team without a picture of a decent future burns out on the present. Take a support desk in Hyderabad in the last week of the quarter. Escalations are stacked, hiring is frozen, and nobody has said a word about next year. A

written scene of a decent year ahead does not clear the queue. It gives each person something to walk toward when the bad week arrives. But one net session does not build a batting technique. A small practice grows only when someone books the next session. Dose is the net schedule.

This paper therefore does one job. It updates Carrillo to 2026. It extracts sittings exactly. It asks whether one sitting and three-to-four sittings produce different effects on optimism, positive affect and wellbeing in adults. It reports a null or unstable dose difference as a result, not as a failure. It names the India gap in plain English.

05

Methods

Eligibility

Population: adults (18 years and over). Trials whose sample was early adolescence were excluded from pooling and noted as out of scope.

Intervention or exposure: best-possible-self writing, imagery, or both, in which participants describe or visualise themselves in a future that has gone as well as it could after effort. Variants that changed only the time frame (past or present “best self”) were retained for narrative context and excluded from the future-self pool when they were not a future-self arm.

Comparator: any control — typical-day writing, last-24-hours writing, gratitude writing, waiting list, or other active filler. Gratitude arms were not treated as the primary contrast.

Primary outcomes: optimism (future-expectancy scales or standing optimism scores), positive affect, and wellbeing (life satisfaction, flourishing, or composite wellbeing). Anxiety and depressive symptoms were recorded when reported and were not primary.

Years: 2001 to 22 September 2026.

Study designs: randomised trials and controlled trials with a comparator. Single-group pre–post studies were excluded from pooling.

Language of publication: English or Spanish, matching Carrillo, plus any English-indexed trial from other languages that reported extractable methods.

Search

Databases: PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI (India), ClinicalTrials.gov, OSF preprints. Carrillo’s search ran to November 2017. Heckerens and Eid extended coverage through late 2019. This update reran the core string through 22 September 2026 and searched the trial registers and preprint server for unpublished adult work.

Full search strings:

PubMed. ("best possible self"[tiab] OR "best possible selves"[tiab] OR "best-possible-self"[tiab]) AND (randomi* [tiab] OR RCT[tiab] OR "controlled trial"[tiab] OR "control group"[tiab])

PsycINFO. TI("best possible self" OR "best possible selves") OR AB("best possible self" OR "best possible selves") AND (randomi* OR "controlled trial" OR RCT)

Scopus. TITLE-ABS-KEY("best possible self" OR "best possible selves") AND TITLE-ABS-KEY(randomi* OR RCT OR "controlled trial")

Web of Science. TS=("best possible self" OR "best possible selves") AND TS=(randomi* OR RCT OR "controlled trial")

Google Scholar. "best possible self" OR "best possible selves" randomi* OR RCT (first 200 hits; cited-by search on King 2001 and Carrillo 2019).

CTRI / ClinicalTrials.gov / OSF. "best possible self" OR "best possible selves".

Core terms, as specified: ("best possible self" OR "best possible selves") AND randomi*.

Reference lists of Carrillo 2019, Loveday 2018, Heekerens and Eid 2021, and Malouff and Schutte 2017 were hand-searched. King 2001 was used as a cited-reference seed.

Study selection and extraction

Two reviewers independently screened titles and abstracts in the source reviews. For the 2018–2026 update, records were screened against the same eligibility rule. Disagreements were resolved by discussion.

Extraction fields, completed for every eligible adult trial that isolated the exercise: first author and year; *n* analysed in the best-possible-self arm and the control arm; design (randomised or controlled); country; language of delivery; dose in sittings and in minutes per sitting; total minutes when reported; delivery channel (face-to-face, online, app, mixed); guidance (self-led, researcher-led, testimonial present); outcome measure; timepoints (immediate post, days, weeks, months); sitting band (one; three-to-four; other).

Sittings were extracted from the methods paragraph of each paper, not from Carrillo's "length in days" column. King 2001 was recoded as four sittings. Carrillo had coded it as one day and 20 total minutes. That coding error is material to any dose story and is flagged wherever King appears.

Risk of bias and certainty

Risk of bias for randomised trials was judged with RoB 2 domains: randomisation process; deviations from intended intervention; missing outcome data; measurement of the outcome; selection of the reported result. Because the exercise is a writing or imagery task, participant blinding is not available. That fact was not treated as a moral failing. It was treated as a built-in limitation of the design.

Certainty of evidence used GRADE. Starting certainty was high for randomised trials and was downgraded for risk of bias, inconsistency, indirectness, imprecision, or publication bias. Indirectness was applied generously. A student sample in Maastricht is not an Indian workplace.

Analysis plan

The pre-specified model was random-effects with REML, Hedges g with a 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. Pre-specified moderators: sittings (one versus three-to-four); writing versus imagery; follow-up (immediate versus one week or later).

Two constraints changed what could be pooled de novo. First, study-level means and standard deviations were not extractable for every eligible trial at the sitting-band level with a shared outcome and a shared timepoint. Second, Carrillo and Heekerens already supply high-quality random-effects pools on the same exercise. Re-averaging their published d and g values as if they were a new REML would pretend at precision the cells do not support.

The analysis therefore has three layers.

Layer one reports Carrillo and Heekerens as the verified overall pools, including their heterogeneity, Egger and trim-and-fill results.

Layer two classifies every verified adult trial into one sitting, three-to-four sittings, or other, using source-paper methods. Effects are described by band. Individual g values appear only where the source paper published them or where a later meta-analysis published a usable study-level estimate. Weights are not invented.

Layer three is the update narrative: trials published after Carrillo's search that change the dose story, the culture story, or the follow-up story.

The fallback rule in the protocol was: if fewer than five eligible trials are found, do not pool. That rule was not triggered. Dozens of adult trials exist. The sittings contrast itself, however, is under-powered as a pooled moderator. That fact is reported as a result.

Position against prior evidence

King 2001 is the parent protocol. Loveday 2018 is the narrative review that asked for dosage trials. Carrillo 2019 is the first comprehensive effectiveness meta-analysis. Malouff and Schutte 2017 is the optimism-intervention meta-analysis that found best-possible-self studies larger than other optimism methods ($g = 0.41$ overall for optimism training). Heekerens and Eid 2021 is the more conservative, timing-sensitive pool. This paper sits on top of those five documents and asks the dose question they left open.

06

Results

PRISMA 2020 flow

Carrillo's search to November 2017 identified 350 database records plus two extra records. After duplicates, 236 unique records were screened. Fifty-five full texts were read. Twenty-six articles reporting 29 studies were included ($N = 2,909$; 1,270 in best-possible-self arms, 1,178 in control arms, 461 in gratitude arms). Mean quality on Carrillo's nine-item scale was 6.58 out of 9 (range 4 to 8). Assessor concealment was present in one of 29 studies. Intention-to-treat analysis was present in 15 of 29.

Heekerens and Eid, searching through late 2019 with a randomised-trial filter and a tighter outcome coding rule, included 34 randomised controlled trials ($N = 3,675$).

The 2018–2026 update identified additional adult randomised or controlled trials that isolate the exercise and report at least one primary outcome. Verified additions used in the sitting-coded synthesis include Carrillo and colleagues 2021 (two trials), Boselie and colleagues 2023, the COVID-era four-week writing trial (Heekerens and colleagues 2024), Booth and colleagues 2025, Radstaak and colleagues 2025, Wu, Hanssen and Peters 2024 and 2025, Auyeung and Mo 2019, and Duffy and colleagues 2025 (dose-randomised; anxiety primary). Adolescent trials and multi-component programmes that did not isolate the exercise were excluded from the adult pool.

No adult Indian trial was identified in CTRI, ClinicalTrials.gov, PubMed, PsycINFO, Scopus, Web of Science, Google Scholar or OSF.

A de-novo count of every 2026 database hit cannot be reconstructed with Carrillo-level precision from public HTML alone. The flow above cites Carrillo's published numbers and treats the update as a verified study-level add-on, not as a second invented PRISMA box.

Study characteristics

The adult literature has a shape. It is young, female, student, short, and Western or East-Asian university-based.

Carrillo's participants had a mean age of 23.56 years (range of study means 17.83 to 35.62). Mean percentage of women was 74%. Twenty studies used university students. Four mixed students with community adults. Two used community adults only. Delivery was individual in 25 studies and group in four; face-to-face in 23 and online in six. Fifteen studies included an imagery component. Compensation was present in 21. Controls were almost always an active neutral writing task. One study used a waiting list.

Dose in the source papers ranges from one 20-minute sitting to eight weeks of weekly writing. Carrillo summarised length as 1 to 56 days (mean 14), intensity as 10 to 75 minutes per week (mean 24), and magnitude as 20 to 170 total minutes (mean 45). Those summaries mix one-sitting laboratory tests with multi-week homework packages. They are not sittings.

Language of delivery in the verified adult set is English, Dutch, Spanish, German or Chinese. No verified trial delivered the exercise in an Indian language.

Sittings, extracted exactly

One sitting. Typical recipe: 15 minutes of writing plus five minutes of imagery, or 20 minutes of writing, once. Verified adult tests include Peters and colleagues 2010 ($n = 44$ versus 38); Harrist and colleagues 2007, writing and talking arms; Hanssen and colleagues 2013 ($n = 79$); Boselie and colleagues 2014, 2016a, 2016b and 2017 (cell sizes 61 to 81); Geschwind and colleagues 2015 ($n = 50$); Meevissen and colleagues 2012 ($n = 72$); Renner and colleagues 2014 ($n = 40$); Peters and colleagues 2016; Heekerens, Eid and Heinitz 2020 ($n = 188$); Heekerens and colleagues 2022 ($n = 432$ German adults; best-possible-self versus gratitude letter versus self-compassion versus control; momentary optimism rose versus control; best-possible-self and gratitude letter did not differ); Boselie and colleagues 2023 ($n = 141$); Booth and colleagues 2025 ($n = 240$); Wu, Hanssen and Peters 2024 ($n = 70$, mainland China). This band is the evidence base for same-day positive affect and same-day positive future expectancies.

Three-to-four sittings. Typical recipe: 15 to 20 minutes, repeated three or four times, either on consecutive days or once a week for four weeks. Verified adult tests include King 2001 (four consecutive days $\times$ 20 minutes; analysed future-self cell about 19 against a control cell of 16, inside an $n = 81$ four-arm study); Lalous, Nelson and Lyubomirsky 2013 (once weekly $\times$ four weeks; face-to-face and online arms); Heekerens and Heinitz 2019 (one laboratory sitting plus two weekly homework sittings); the COVID-era four-week trial (once weekly $\times$ four weeks; $n = 87 / 85 / 82$); Wu, Hanssen and Peters 2025 (three 20-minute sittings in one week at two-day intervals; $n = 61$ Chinese and 48 Dutch). Duffy and colleagues 2025 randomises one versus two versus three versus four sittings, but the primary outcome is generalised anxiety, not optimism or wellbeing.

Other doses. These trials are eligible for the overall story and ineligible for the one-versus-three-to-four contrast. Sheldon and Lyubomirsky 2006: one laboratory sitting plus four weeks of home practice. Meevissen, Peters and Alberts 2011: one construction sitting then five minutes of imagery daily for 14 days. Enrique and colleagues 2017: about 30 days, 170 total minutes. Liao and colleagues 2016: Singapore undergraduates, about 40 total minutes over 30 days. Boehm, Lyubomirsky and Sheldon 2011: six weeks, ten minutes a week. Lyubomirsky and colleagues 2011: eight weeks. Auyeung and Mo 2019: six consecutive days online, Chinese students, $n = 100$. Carrillo and colleagues 2021: two trials, seven days, past or present or future best-self versus last-24-hours control. Radstaak, Carrillo and colleagues 2025: two-week app, $n = 290$.

Overall effects — the verified pools

Carrillo versus controls:

- Wellbeing: $k = 29$, $d+ = 0.325$, 95% CI 0.189 to 0.461, Q significant, $I^2 = 73.8\%$.
- Optimism: $k = 13$, $d+ = 0.334$, 95% CI 0.246 to 0.422, Q not significant, $I^2 = 0\%$.

- Positive affect: $k = 13$, $d+ = 0.511$, 95% CI 0.257 to 0.765, Q significant, $I^2 = 79.3\%$.
- Negative affect: $k = 13$, $d+ = 0.192$, 95% CI -0.328 to 0.712 , $I^2 = 94.9\%$. With the one non-active control removed, $d+ = -0.047$.
- Depressive symptoms: $k = 3$, $d+ = 0.115$, 95% CI -0.272 to 0.502 .

Egger tests in Carrillo were not significant. Trim-and-fill left most estimates in place. For optimism, four studies were imputed and $d+$ fell from 0.33 to 0.28. That adjusted figure is the more honest number to carry in a pocket.

Heckerens and Eid versus controls, Hedges g :

- Positive affect: $g = 0.28$, 95% CI 0.16 to 0.41.
- Optimism: $g = 0.21$, 95% CI 0.04 to 0.38.
- Immediate-day positive affect: $g = 0.39$.
- Immediate-day optimism: $g = 0.36$.
- Follow-up at about seven days: no substantial effects.
- Optimism was significant when coded as positive future expectancies. It was not significant when coded as a general orientation in life.
- Online administration did not produce a significant positive-affect effect in their moderator cut.

Malouff and Schutte, across optimism-training trials rather than best-possible-self trials only: $g = 0.41$. Best-possible-self studies sat above other methods. Effects were larger when the last assessment fell within one day of the exercise, when delivery was in person, and when expectancies rather than a standing optimism score were used.

The three pools agree on the shape. There is a same-day lift. It is small to moderate. It is clearer for affect and for expectancies than for a trait. It shrinks when the researcher comes back a week later.

Forest-plot data table

A study-level forest with weights was not computed de novo. Weights without study-level variances are theatre. The table below reports the verified pooled rows and the sitting-banded trials that published a usable contrast. Where a source paper reported direction without a usable g , the cell says “direction only”.

Table. Verified pooled effects and sitting-banded adult trials

Study or pool	Sittings	g or d	95% CI	Weight	Outcome and timepoint
Carrillo 2019 pool	Mixed	$d+ 0.325$	$0.189, 0.461$	—	Wellbeing, first post
Carrillo 2019 pool	Mixed	$d+ 0.334$	$0.246, 0.422$	—	Optimism, first post

Study or pool	Sittings	<i>g</i> or <i>d</i>	95% CI	Weight	Outcome and timepoint
Carrillo 2019 pool, trim-and-fill	Mixed	d+ 0.28	—	—	Optimism, adjusted
Carrillo 2019 pool	Mixed	d+ 0.511	0.257, 0.765	—	Positive affect, first post
Heekerens 2021 pool	Mixed	g 0.28	0.16, 0.41	—	Positive affect
Heekerens 2021 pool	Mixed	g 0.21	0.04, 0.38	—	Optimism
Heekerens 2021, day-of	Mixed	g 0.39	—	—	Positive affect, same day
Heekerens 2021, day-of	Mixed	g 0.36	—	—	Optimism, same day
Heekerens 2021, ~7-day	Mixed	ns	—	—	Affect and optimism, follow-up
Peters 2010	1	Direction: BPS > control	—	—	Positive expectancies and PA, immediate
Boselie 2023, write	1	Within-arm <i>d</i> on PA ~0.90 to 1.12 versus own baseline; between-arm vs typical-day smaller	—	—	PA and expectancies, immediate
Boselie 2023, imagery	1	Equivalent to writing	—	—	PA and expectancies, immediate
Booth 2025	1	Large direct effects on +/- expectancies and positive mood at post; small-to-medium on +expectancies at 1 week	—	—	Expectancies, mood, anxiety (indirect)
Wu 2024 China	1	Null on PA and wellbeing; reduction in NA and negative expectancies	—	—	PA, WB, NA, expectancies
King 2001	4	SWB higher at 3 weeks vs control; less upsetting than trauma write	—	—	Subjective wellbeing, 3 weeks
Layous 2013	4 weekly	PA and flow up; testimonial enlarged the gain; online = in person	—	—	PA, flow, 4 weeks
Heekerens 2019 thriving	~3	Path $r = .11$ on thriving	—	—	Thriving, 2 weeks
COVID 4-week trial 2024	4 weekly	PA and mood up in week 1 only; optimism ns	—	—	PA, mood, optimism
Wu 2025 CN–NL	3 in 1 week	Dutch > Chinese on PA, trait optimism, life satisfaction	—	—	PA, optimism, SWL
Carrillo 2021, two RCTs	7 days (other)	All arms rose; BPS not > last-24-hours control	—	—	PA and wellbeing

Study or pool	Sittings	<i>g</i> or <i>d</i>	95% CI	Weight	Outcome and timepoint
Duffy 2025	1 vs 2 vs 3 vs 4	Anxiety dropped at ≥ 2 sittings; 1 sitting null	—	—	GAD symptoms (not primary here)
Radstaak 2025 app	~14 (other)	Some mental-health benefit; engagement predicted outcome, adherence did not	—	—	Life satisfaction and mental health

Dash cells are not missing maths. They are a refusal to invent a weight.

Moderator table

Table. Pre-specified moderators

Moderator	Finding	Certainty	Source
Sittings: 1 vs 3–4	One sitting produces the same-day PA and expectancy lift. Three-to-four sittings is the King wellbeing protocol and shows mixed, often time-limited, effects. No stable pooled <i>g</i> shows 3–4 sittings beat 1 sitting on the same outcome at the same hour.	Low	Sitting-coded update; Carrillo days ns; Carrillo minutes trend negative
Total minutes	Shorter total practice, larger wellbeing (trend $p = .078$; 25% variance).	Low	Carrillo 2019
Length in days	Not significant.	Low	Carrillo 2019
Writing vs imagery	One sitting, online: writing, imagery, and writing-plus-imagery were equivalent for optimism and affect.	Moderate for immediate effects	Boselie 2023
Follow-up	Immediate effects are the reliable ones. Heekerens: no substantial ~7-day effects. Enrique: optimism up at post, not at 3 months. COVID 4-week trial: week-1 only. Booth: expectancies, not mood, survive 1 week.	Moderate	Heekerens 2021 plus update
Age	Older-within-young-adult samples trended toward larger wellbeing.	Low	Carrillo 2019, $p = .079$
Delivery channel	Online versus in-person not different in Layous 2013. Heekerens: online ns for PA.	Low	Layous 2013; Heekerens 2021
Culture	Mainland China one-sitting trial: no PA or wellbeing rise. Chinese–Dutch three-sitting trial: Dutch gains larger on positive outcomes.	Low (small <i>n</i>)	Wu 2024; Wu 2025
Testimonial / rationale	A peer testimonial enlarged wellbeing gains.	Low	Layous 2013
Engagement vs adherence	Engagement predicted app outcome. Adherence (showing up) did not.	Low	Radstaak 2025

Heterogeneity

Wellbeing and positive affect are heterogeneous. Carrillo’s wellbeing I^2 was 74%. Positive affect I^2 was 79%. Optimism in Carrillo was homogeneous ($I^2 = 0\%$). That last number is the most surprising result in the 2019 paper and the most useful. When laboratories measure future expectancies after a best-possible-self sitting, they tend to find the same small-to-moderate lift. When they measure “wellbeing” with a mixed bag of life-satisfaction, happiness and composite scales, at mixed delays, the bag sloshes.

Heekerens’s smaller overall *g* values are the price of coding time and construct more tightly. That is a feature. A mood induction that is sold as a trait change will look disappointing in a tight coding scheme. It still does the job it actually does.

A 95% prediction interval was not recomputed de novo. Given Carrillo’s wellbeing *I*² of 74%, any prediction interval for a new wellbeing trial would be wide enough to include a small negative effect. Programmes should plan for that.

Publication bias

Carrillo’s Egger tests were not significant. Trim-and-fill moved optimism from 0.33 to 0.28 and left the other primary estimates largely in place. Heekerens’s more conservative overall *g* already sits near that adjusted optimism figure. The update did not find a file drawer of large negative adult trials. It did find several published null or time-limited results (Carrillo 2021; COVID four-week trial; Wu 2024 on positive outcomes). Those papers reduce the odds that the remaining literature is only the happy tail.

Risk-of-bias summary

No verified adult trial in this set is low risk across all RoB 2 domains.

- Randomisation process: some concerns in most papers. Sequence generation and allocation concealment are often unreported.
- Deviations from intended intervention: some concerns by design. Participants know whether they are writing a glowing future or a typical Tuesday. Researchers who run the session usually know too.
- Missing outcome data: low risk in one-sitting laboratory tests; high risk in multi-week online packages, where attrition of 15% to 40% is common.
- Measurement of the outcome: some concerns. Outcomes are self-report. Assessors are rarely concealed (one of 29 in Carrillo).
- Selection of the reported result: some concerns. Pre-registration is recent. Booth 2025 is a welcome exception.

Overall judgement for the immediate one-sitting literature: some concerns. Overall judgement for the multi-week wellbeing literature: some concerns to high, driven by attrition and unblinded self-report. Duffy 2025, the only sittings-randomised trial, is not in the primary pool because its outcome is anxiety.

GRADE table

Table. Certainty of evidence

Outcome	Effect	<i>k</i> / source	GRADE	Why
Immediate positive affect	Small-to-moderate lift (<i>g</i> 0.28 overall; <i>g</i> 0.39 day-of; <i>d</i> + 0.51 in Carrillo)	Heekerens 34 RCTs; Carrillo 13	Moderate	Consistent direction, student samples, unblinded self-report, same-day clock
Immediate optimism as future expectancies	Small-to-moderate lift (<i>g</i> 0.21 overall; <i>g</i> 0.36 day-of; <i>d</i> + 0.33 in Carrillo; 0.28 trim-and-fill)	Heekerens; Carrillo 13	Moderate	Homogeneous in Carrillo; construct-sensitive in Heekerens
Trait optimism (standing score)	Small or absent	Heekerens construct cut	Low	Effect depends on how the construct is scored

Outcome	Effect	k / source	GRADE	Why
Wellbeing	Small-to-moderate ($d+ 0.33$), heterogeneous	Carrillo 29	Low to moderate	I^2 74%; mixed scales; mixed delays
Follow-up ≥ 1 week	Small or absent	Heekerens; update trials	Low	Consistent fade
Writing vs imagery, one sitting	Equivalent	Boselie 2023, $n = 141$	Moderate for immediate effects	One well-run equivalence test
Sittings: 1 vs 3–4 on optimism or wellbeing	No confirmed multiplier	Sitting-coded update	Low	No prior pooled test of this split; timing confounded with dose; only sittings-randomised RCT measured anxiety
Applicability to Indian working adults	Unknown	Zero adult Indian trials	Very low	Indirectness is the story

07 Discussion

PULL QUOTE

One sitting is the MVP. Three-to-four sittings is version 1.0. Confusing the two is how programmes burn the only resource Indian professionals will not give you twice: Tuesday night.

What the number means

Start with the number that survives contact with the more careful pool. Heekerens and Eid's $g = 0.28$ for positive affect and $g = 0.21$ for optimism. On the day you write, those figures rise to about 0.39 and 0.36. A week later they slump. Carrillo's friendlier $d+ = 0.33$ for wellbeing and 0.51 for positive affect is the same animal photographed with a wider lens.

In stand-up terms: the exercise is a good opening joke. It is not a two-hour special. The room laughs. The room does not change its tax status.

In machine-learning terms: one epoch moves the weights. Extra epochs do not automatically drop the loss. If you keep training on the same short prompt, you can overfit a mood. The validation set is next Tuesday morning, when the future you wrote is still a paragraph and the inbox is not.

In plain coaching terms: the sitting is a written answer to a hopeless forecast, with a handle attached. The scene is real. The handle is the next small act that belongs in it. An analyst on the morning local writes that next year she presents

to the client instead of taking notes. Before she reaches her desk she sends one message: can she open Thursday's call? Without that message, she has written the scene and skipped the step.

The same sitting does something quieter too. The writer sees the gloomy forecast as a forecast, not a fact. She is not only the tired person on the train. She is also the one who can describe a future self without having to be finished yet. Writing against the gloom for 20 minutes is a rep. Forcing cheer all day is a different product.

In entrepreneurship terms: one sitting is the MVP. Three-to-four sittings is version 1.0. Shipping the MVP is how you learn whether the customer will use the product. Shipping version 1.0 is how you learn whether the customer will keep it. Confusing the two is how start-ups burn runway. Carrillo's minutes trend — shorter total practice, larger wellbeing — is the rare methods paper that sounds like a product manager. Do not pad the onboarding.

The sittings question, asked cleanly, has an unglamorous answer. One sitting does the same-day job. Three-to-four sittings is the original wellbeing protocol and the version that can be scheduled as a workplace habit. Extra sittings have not been shown to multiply the same-hour *g*. They have been shown, in King's hands, to sit next to a three-week wellbeing gain, and in later hands, to fade if the only engine is repetition of the same paragraph.

Loveday's 2018 line that repeating "appears to enhance efficacy" was a fair reading of a narrative pile. It was not a pooled test. The trials that arrived after 2018 do not rescue a simple more-is-better rule. Carrillo 2021 found that past, present and future best-self arms, and the last-24-hours control, all improved. The COVID four-week trial found a week-1 mood lift that did not ride the next three weeks. Wu's mainland-China one-sitting trial found no positive-affect or wellbeing rise. Duffy found that two or more sittings reduced anxiety and one sitting did not — a dose signal, on a different outcome, that programmes should file and not over-claim.

King was miscoded in the 2019 pool as a one-day, 20-minute exercise. She ran four sittings. Any dose story that treats King as a one-sitting wellbeing win starts from the wrong data.

Limitations

This is an updated systematic review with quantitative synthesis of prior pools and a sitting-coded narrative of the new trials. It is not a de-novo REML of every study-level *g* with a new forest and invented weights. That limit is methodological, not rhetorical.

Most participants were students. Most were women. Most were paid. Most wrote in a laboratory or on a university platform. Most outcomes were immediate self-reports. Follow-up is thin. Pre-registration is recent. Cultural range is narrow. The only sittings-randomised trial measured anxiety. Engagement, not adherence, predicted outcome in the app trial — which means "they opened the app" is the wrong process metric.

Carrillo's days-versus-minutes coding cannot be used as a sittings code. We recoded from source methods. Residual misclassification is possible where a paper said "over two weeks" and did not say how many times the person actually sat down.

Publication bias is not absent. It is also not the whole story. Enough null and time-limited papers now exist that a reader can see the fade without needing a funnel plot.

WHAT THIS MEANS

What this means for an Indian workplace programme

Here is the rollout in a Pune sales team. On Monday a WhatsApp message arrives, in Marathi for those who think in Marathi, in Hindi or English for the rest. It offers one 15-to-20-minute sitting: write, or simply picture, a decent year ahead. That sitting is the finding with moderate certainty. It gives a same-day lift in mood and in the sense that a decent future is imaginable. Anyone who wants a wellbeing habit gets three further sittings across the week, booked in the calendar like a client meeting. That follows King's protocol, and a habit needs a calendar. The extra sittings have not been shown to enlarge the first-hour number. Some people write in a notebook. Some picture the scene on the train home. Boselie 2023 says those two doors open onto the same room. Nobody is promised a new personality. The exercise is offered as a rep. When one person's page turns into a problem the model should not hold, a human takes over the conversation. And nobody shows the regional head a Maastricht gas local evidence. Local evidence does not exist yet. That is the brief for the next trial, not a reason to withhold a 20-minute practice that has been safe, cheap and directionally helpful for 25 years.

Research gaps, naming the India gap

The India gap is explicit. There is no adult Indian randomised trial of best-possible-self writing on optimism, positive affect or wellbeing. CTRI does not hold one. ClinicalTrials.gov does not hold one with an Indian site. PubMed, PsycINFO and OSF do not hold one. The closest Asian adult data are Singaporean (Liau 2016) and Chinese (Auyeung and Mo 2019; Wu 2024; Wu 2025). Wu's two papers are the warning label. An individualistic future-self script can land as a positive-affect exercise in a Dutch sample and as a negative-affect reducer in a Chinese sample. Indian workplaces hold both individual ambition and family duty in the same sentence. A protocol that asks only "what is the best possible you?" and never "what is the best possible us?" is an untested translation.

Other gaps follow. Someone still needs to randomise one sitting versus three-to-four sittings on optimism and wellbeing in working adults, with the clock held constant, with a one-week and one-month follow-up, and with a process

measure of engagement rather than login counts. Someone needs to test Hindi, Tamil, Telugu, Bengali and Marathi scripts. Someone needs to test whether a values-and-family frame changes the Chinese-style pattern Wu reported. Someone needs to stop treating LOT-R change after 20 minutes as a personality result. Future expectancies are the honest outcome. Trait scores are the vanity metric.

One net session is still not a technique, and overfitting is still a risk. Book the next session. Do not train the same prompt until it memorises a mood.

08

FAQ

Q1 Does writing about my best future self actually do anything?

Yes, a little, the same day. Pooled effects sit around $g = 0.28$ for positive affect and $g = 0.21$ for optimism, rising to about 0.36 – 0.39 on the day you write. That is a real lift, not a personality change.

Q2 Is one sitting enough, or do I need three or four?

One sitting is enough for a same-day lift in mood and in positive future expectancies. Three-to-four sittings is King's original wellbeing protocol. Extra sittings have not been shown to enlarge the same-hour effect. They are how you turn a spark into a habit.

Q3 Should I write it or just picture it?

Either. A 2023 online trial found writing, imagery, and both were equivalent for optimism and affect after one sitting. Pick the door you will actually walk through. Twenty honest minutes beat a perfect method you skip.

Q4 Will this still work next week?

Often, no. The careful 2021 pool found no substantial effects about seven days later. One 2025 trial found that positive expectancies, not mood, survived a week. Plan a rep, not a one-off transformation.

Q5

Has anyone tested this with Indian working adults?

No. Zero adult Indian randomised trials were found. Singapore and mainland China are the closest published neighbours, and the Chinese data suggest the script does not travel unchanged. That gap is the next study, not a reason to forbid a 20-minute practice.

09

References

- Auyeung, L., & Mo, P. K. H. (2019). The efficacy and mechanism of online positive psychological intervention (PPI) on improving well-being among Chinese university students: A pilot study of the best possible self (BPS) intervention. *Journal of Happiness Studies*, 20, 2525–2550. <https://doi.org/10.1007/s10902-018-0054-4>
- Boehm, J. K., Lyubomirsky, S., & Sheldon, K. M. (2011). A longitudinal experimental study comparing the effectiveness of happiness-enhancing strategies in Anglo Americans and Asian Americans. *Cognition and Emotion*, 25(7), 1263–1272. <https://doi.org/10.1080/02699931.2010.541227>
- Booth, R. W., Erhan, K., Erkocaoğlu, O., Kuşpınar, H., & Yaldirak, K. (2025). The best possible self task has direct effects on expectancies and mood, and an indirect effect on anxiety symptom severity. *Emotion*, 25(4), 964–971. <https://doi.org/10.1037/emo0001481>
- Boselie, J. J. L. M., Vancleef, L. M. G., van Hooren, S., & Peters, M. L. (2023). The effectiveness and equivalence of different versions of a brief online Best Possible Self (BPS) manipulation to temporary increase optimism and affect. *Journal of Behavior Therapy and Experimental Psychiatry*, 79, 101837. <https://doi.org/10.1016/j.jbtep.2023.101837>
- Carrillo, A., Martínez-Sanchis, M., Etchemendy, E., & Baños, R. M. (2019). Qualitative analysis of the Best Possible Self intervention: Underlying mechanisms that influence its efficacy. *PLOS ONE*, 14(5), e0216896. <https://doi.org/10.1371/journal.pone.0216896>
- Carrillo, A., Rubio-Aparicio, M., Molinari, G., Enrique, Á., Sánchez-Meca, J., & Baños, R. M. (2019). Effects of the Best Possible Self intervention: A systematic review and meta-analysis. *PLOS ONE*, 14(9), e0222386. <https://doi.org/10.1371/journal.pone.0222386>
- Carrillo, A., Etchemendy, E., & Baños, R. M. (2021). My best self in the past, present or future: Results of two randomized controlled trials. *Journal of Happiness Studies*, 22, 955–980. <https://doi.org/10.1007/s10902-020-00259-z>
- Duffy, J., et al. (2025). Efficacy of the best possible self intervention for generalised anxiety: Exploration of mediators and moderators. *The Journal of Positive Psychology*. <https://doi.org/10.1080/17439760.2025.2487442>
- Enrique, Á., Bretón-López, J., Molinari, G., Roca, P., Llorca, G., Guillén, V., Fernández-Aranda, F., Baños, R. M., & Botella, C. (2018). Implementation of a positive psychology group program in an inpatient eating disorders service: A pilot study. Related trial record NCT02321605; see also the e-BPS effectiveness papers from this group. *Applied Research in Quality of Life* line: <https://doi.org/10.1007/s11482-017-9552-5>
- Hanssen, M. M., Peters, M. L., Vlaeyen, J. W. S., Meevissen, Y. M. C., & Vancleef, L. M. G. (2013). Optimism lowers pain: Evidence of the causal status and underlying mechanisms. *Pain*, 154(1), 53–58. <https://doi.org/10.1016/j.pain.2012.08.010>
- Harrist, S., Carlozzi, B. L., McGovern, A. R., & Harrist, A. W. (2007). Benefits of expressive writing and expressive talking about life goals. *Journal of Research in Personality*, 41(4), 923–930. <https://doi.org/10.1016/j.jrp.2006.09.002>

- Heekerens, J. B., & Heinitz, K. (2019). Looking forward: The effect of the best-possible-self intervention on thriving through relative intrinsic goal pursuits. *Journal of Happiness Studies*, 20, 1379–1395. <https://doi.org/10.1007/s10902-018-9999-6>
- Heekerens, J. B., & Eid, M. (2021). Inducing positive affect and positive future expectations using the best-possible-self intervention: A systematic review and meta-analysis. *The Journal of Positive Psychology*, 16(3), 322–347. <https://doi.org/10.1080/17439760.2020.1716052>
- Heekerens, J. B., Eid, M., & Heinitz, K. (2020). Dealing with conflict: Reducing goal ambivalence using the best-possible-self intervention. *The Journal of Positive Psychology*, 15(3), 325–337. <https://doi.org/10.1080/17439760.2019.1610479>
- Heekerens, J. B., Heinitz, K., Eid, M., & Riemann, R. (2022). Cognitive-affective responses to online positive-psychological interventions: The effects of optimistic, grateful, and self-compassionate writing. *Applied Psychology: Health and Well-Being*. <https://doi.org/10.1111/aphw.12326>
- Heekerens, J. B., et al. (2024). Best possible selves in times of crisis: Randomized controlled trial of best possible self-interventions during the COVID-19 pandemic. *The Journal of Positive Psychology*, 19(6), 1080–1090. <https://doi.org/10.1080/17439760.2023.2297213>
- King, L. A. (2001). The health benefits of writing about life goals. *Personality and Social Psychology Bulletin*, 27(7), 798–807. <https://doi.org/10.1177/0146167201277003>
- Layouts, K., Nelson, S. K., & Lyubomirsky, S. (2013). What is the optimal way to deliver a positive activity intervention? The case of writing about one's best possible selves. *Journal of Happiness Studies*, 14, 635–654. <https://doi.org/10.1007/s10902-012-9346-2>
- Liau, A. K., Neihart, M. F., Teo, C. T., & Lo, C. H. M. (2016). Effects of the best possible self activity on subjective well-being among Singaporean adolescents. *The Asia-Pacific Education Researcher*, 25, 473–481. <https://doi.org/10.1007/s40299-015-0273-y>
- Loveday, P. M., Lovell, G. P., & Jones, C. M. (2018). The best possible selves intervention: A review of the literature to evaluate efficacy and guide future research. *Journal of Happiness Studies*, 19, 607–628. <https://doi.org/10.1007/s10902-016-9824-z>
- 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>
- Malouff, J. M., & Schutte, N. S. (2017). Can psychological interventions increase optimism? A meta-analysis. *The Journal of Positive Psychology*, 12(6), 594–604. <https://doi.org/10.1080/17439760.2016.1221122>
- Meevissen, Y. M. C., Peters, M. L., & Alberts, H. J. E. M. (2011). Become more optimistic by imagining a best possible self: Effects of a two week intervention. *Journal of Behavior Therapy and Experimental Psychiatry*, 42(3), 371–378. <https://doi.org/10.1016/j.jbtep.2011.02.012>
- Peters, M. L., Flink, I. K., Boersma, K., & Linton, S. J. (2010). Manipulating optimism: Can imagining a best possible self be used to increase positive future expectancies? *The Journal of Positive Psychology*, 5(3), 204–211. <https://doi.org/10.1080/17439761003790963>
- Radstaak, M., Carrillo, A., Etchemendy, E., Baños, R. M., & Bohlmeijer, E. T. (2025). My best self in the past or future: A randomized controlled trial examining adherence, engagement, age and mental health in a mobile-based BPS intervention. *Journal of Happiness Studies*, 26, 131. <https://doi.org/10.1007/s10902-025-00962-9>
- Renner, F., Schwarz, P., Peters, M. L., & Huibers, M. J. H. (2014). Effects of a best-possible-self mental imagery exercise on mood and dysfunctional attitudes. *Psychiatry Research*, 215(1), 105–110. <https://doi.org/10.1016/j.psychres.2013.10.033>
- Sheldon, K. M., & Lyubomirsky, S. (2006). How to increase and sustain positive emotion: The effects of expressing gratitude and visualizing best possible selves. *The Journal of Positive Psychology*, 1(2), 73–82. <https://doi.org/10.1080/17439760500510676>

Wu, L., Hanssen, M. M., & Peters, M. L. (2024). The effectiveness and mechanisms of a brief online best-possible-self intervention among young adults from mainland China. *The Journal of Positive Psychology, 19*(6), 1066–1079. <https://doi.org/10.1080/17439760.2023.2297204>

Wu, L., Hanssen, M. M., & Peters, M. L. (2025). Best possible self and culture: A Chinese–Dutch comparison of three sittings. *Journal of Happiness Studies, 26*, 22. <https://doi.org/10.1007/s10902-024-00855-3>

About the author

ABOUT THE AUTHOR

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

Dr Atul Aswani is a Mumbai Psychiatrist and trained as a Results Coach. He builds EmotionApp, a non-clinical emotional-fitness platform for Indian professionals, under Tranquilo Meditek Sytems Private Limited, Mumbai. The work packages positive-psychology techniques as countable daily reps, with an AI coach, human hand-off, WhatsApp-first delivery and 23 Indian languages. Data stay in India. The system is DPDP-compliant and built under ISO 9001 and ISO 13485.

HOW TO CITE

Aswani A. Best Possible Self, Best Available Evidence: A Meta-Analysis of Future-Self Writing and Its Dose. EmotionApp Research Paper 26. Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/best-possible-self-dose-meta-analysis>

LAST UPDATED

Last updated: 22 September 2026

NOTE

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

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

SAFETY NOTE

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