

Three Good Things, Twenty Years Later

A Meta-Analysis of Gratitude Dose — Daily versus Weekly

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



POOLED EFFECT · HEDGES G

0.19

95% CI 0.15 to 0.22

The pooled well-being effect of gratitude interventions is small: Hedges $g = 0.19$ (95% CI 0.15 to 0.22), from 145 papers and 24,804 people in 28 countries. The listing, or counting-blessings, subtype sits at $g = 0.17$. Weekly practice has never beaten daily practice in a published head-to-head trial. Zero such trials exist.

Three Good Things, Twenty Years Later

A Meta-Analysis of Gratitude Dose — Daily versus Weekly

Lyubomirsky's weekly-beats-daily finding has never been pooled. Which dosing schedule produces durable gains, and at what point does gratitude go stale?

PRISMA 2020 systematic review and frequency synthesis

CONTENTS

01 Abstract

02 Key findings box

03 Definition block

04 Introduction

05 Methods

06 Results

07 Discussion

08 FAQ

09 References

INDEXING TITLE

Dose and frequency in gratitude-listing interventions: a meta-analysis of daily versus weekly practice

AUTHOR

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

PUBLISHED

22 September 2026 · English edition · emotionapp.in/research/gratitude-dose-daily-versus-weekly

01

Abstract

Lyubomirsky's weekly-beats-daily claim has been repeated for twenty years and has never been pooled. The pooled well-being effect of gratitude interventions is small: Hedges $g = 0.19$ (95% CI 0.15 to 0.22), from the largest pre-registered synthesis ($k = 145$ papers, $N = 24,804$, 28 countries). The listing, or "counting blessings", subtype sits slightly lower, at $g = 0.17$.

We asked a narrower question. Does weekly gratitude listing produce larger or more durable gains than daily listing in adults? We searched PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI, ClinicalTrials.gov and OSF from 2003 to September 2026. Eligible studies were randomised or controlled trials of gratitude listing or journaling with frequency stated in the methods — never inferred. Primary outcomes were well-being and positive affect. Depressive symptoms were secondary.

Twelve listing trials stated frequency exactly and met adult-population rules. Zero published trials randomised daily listing against weekly listing. The famous "weekly beats daily" contrast is a mis-statement of an unpublished student study that compared one session a week with three sessions a week. A prior meta-analysis already tested days of listing as a moderator of well-being and found no effect ($p = .233$).

Daily lists produced small gains in positive affect in some trials (typical published $d = 0.15$ to 0.38 versus a neutral control; larger versus a hassles control) and a null in the largest listing trial ($N = 958$). Weekly lists produced small gains in life evaluation in some trials and nulls in others. Twice-weekly listing for four weeks, in Hong Kong hospital staff, reduced depressive symptoms ($d = -0.49$) and perceived stress ($d = -0.95$) at three months. Effects appear by week one or two. Stretching a list from seven days to fourteen does not reliably enlarge the gain. Thrice-weekly prompts for three weeks faded by three months.

Certainty is moderate for a small listing-versus-control gain, and very low for any claim that weekly beats daily. No high-quality adult listing trial from India was found. For an Indian workplace programme the honest dose is not "daily forever". It is a short, prompted, language-local list, one to three times a week, with a stop-rule when the practice goes stale.

Keywords: *gratitude listing; three good things; dose; frequency; well-being; positive affect; India; PRISMA 2020.*

02

Key findings box

KEY FINDINGS

Five numbers worth repeating

0.19

0.19 — pooled Hedges g for gratitude interventions on well-being (95% CI 0.15 to 0.22; 145 papers; 24,804 people; 28 countries). Listing sits at $g = 0.17$.

0

0 — published randomised trials that pit daily listing against weekly listing in adults. The weekly-beats-daily line has never been a pooled contrast.

0.38

0.38 — Hedges g for positive affect after fourteen daily lists versus a neutral-events control in the largest extractable listing trial (Brazil; analysed $n = 287$). The same trial's happiness, life satisfaction and depression contrasts against neutral events were not larger.

-0.49

-0.49 — standardised drop in depressive symptoms at three months after twice-weekly work-gratitude lists for four weeks, in 102 Hong Kong hospital staff, versus no diary. Perceived stress fell by $d = -0.95$. This is the strongest professional-sample signal.

233

.233 — p -value for “days of listing” as a moderator of well-being in Davis and colleagues' 2016 meta-analysis. More days did not mean more gain. Four or fewer workplace lists, in a later occupational review, changed nothing.

03

Definition block

DEFINITION

Gratitude listing is a short writing practice. A person names a handful of specific good things — three is the usual number, five is the old number — and, in some versions, adds a sentence on why each one happened. It is not a mood diagnosis. It is not a clinical programme. It is an attention drill.

The mind already keeps a running tally of slights, delays, and unfinished work. On the drive home a Chennai accounts executive replays the late invoice and the curt email. The list is a counter-tally: the vendor who paid early, the colleague who covered her call. It trains attention to land, on purpose, on what went right. It works while each line still means something; copied out of habit, it becomes a dead rule. An entrepreneur would call it a minimum viable practice: ship three lines, measure the feeling, iterate when it goes stale.

Fresh items keep the practice alive. A team lead in Pune who writes “family, health, job” for the ninth night running is no longer looking; she is copying. Yesterday’s list, reheated, goes stale.

04

Introduction

Indian working adults do not lack advice. They lack minutes. A Bengaluru product manager closes Slack at 8.40 p.m., eats standing at the kitchen counter, and a reminder pops up: write three good things before sleep. A Pune ward nurse finishes a double shift, sits on the bus home, and finds the same homework on her phone. The question is not whether gratitude is virtuous. Most Indian households already practise thanks — at the dinner table, to the neighbour who took in a parcel, in a WhatsApp forward. The question is dose. How often should a countable rep be asked of a tired professional before the rep becomes another item on the to-do list?

The scientific story began cleanly. In 2003 Emmons and McCullough asked students to list up to five blessings. One sample did it every week for ten weeks. Another did it every day for a fortnight. A third sample, adults with neuromuscular disease, listed daily for three weeks. Blessing-listers felt better, in patches, than people who listed hassles. Positive affect was the most reliable signal.

Two years later Seligman, Steen, Park and Peterson put “three good things” and a causal explanation on the internet for one week. Happiness rose and depressive symptoms fell, and the curves were still bent at six months among people who kept going. That is where the nightly-journal rule began: a notebook on the bedside table, three lines before the light goes off.

In the same season Lyubomirsky, Sheldon and Schkade described a six-week student study. Counting blessings once a week raised happiness. Counting them three times a week did not. The interpretation was hedonic adaptation: too often, and the task becomes a chore. The finding entered popular psychology as “weekly beats daily”. That sentence is wrong in two ways. The trial compared weekly with thrice-weekly, not with daily. And the contrast was never published as a standalone paper with extractable means, standard deviations, and a Hedges g . It lives in a review, a book, and a talk. Treating it as settled dose-response is like treating a pitch deck as an audited P&L.

Meanwhile the market for gratitude did what markets do. It overfit. Apps asked for a daily streak. Schools set homework. Workplace wellness slides cited Emmons and Seligman and moved on. Four meta-analyses then poured cold water, politely.

Davis and colleagues (2016) found a small well-being gain against an alternative activity ($d = 0.17$) and a larger one against measurement-only ($d = 0.31$). Days of listing did not moderate the effect. Minutes did not either. Dickens (2017) said unique benefits were over-sold. Cregg and Cheavens (2021) put the depression-and-anxiety effect at $g = -0.29$ after the programme and $g = -0.23$ at follow-up — small, and larger against wait-lists than against active controls. Kirca, Malouff and Meynadier (2023) restricted to expressed gratitude and got $g = 0.22$ for well-being; length did not moderate.

Choi, Cha, McCullough, Coles and Oishi (2025), in the Proceedings of the National Academy of Sciences, pooled 145 papers and 24,804 people from 28 countries. Overall $g = 0.19$. Counting blessings sat at $g = 0.17$. Sessions were “less critical” than the type of outcome and the fact of randomisation. India was among the countries where the country-level effect on well-being was not significant.

None of those syntheses modelled daily versus weekly listing as a confirmatory contrast. That is the gap this paper fills. Not with a fake forest plot of a comparison that does not exist, but with a PRISMA-structured reading of every adult listing trial that actually states its frequency, a conservative alignment with the best-powered pooled number, and a plain answer for people who build programmes in India.

Why India, specifically? Because the positivity industry travels faster than the evidence. Because collectivist thanks is not the same act as an individualist journal. Because a 2025 multi-country synthesis already flagged Indian samples as a place where the imported programme did not lift well-being. And because an emotional-fitness coaching platform that delivers countable reps

on WhatsApp, in twenty-three Indian languages, under DPDP and ISO rules, should not copy a Californian daily streak if the dose that survives contact with an Indian workday is weekly.

Machine-learning people have a name for repeating the same batch until the loss stops falling: overfitting. The model memorises “family, health, job” and stops noticing the new thing that happened at 4 p.m. — the intern who fixed the build without being asked. A coach would change the prompt. An entrepreneur would kill the feature that nobody uses. This paper is the kill-or-keep memo for daily gratitude listing.

05

Methods

Eligibility

Population. Adults aged 18 and over, or mixed-age samples that were majority adult. Non-clinical or mixed. Workplace samples included. Child-only and adolescent-only trials excluded.

Intervention. Gratitude listing or journaling with frequency stated in the methods in so many words: daily, weekly, twice weekly, three times weekly, or another explicit schedule. Items per session recorded when reported. The “why it happened” prompt recorded when present. Gratitude letters and gratitude visits were excluded unless a listing arm could be separated. Multi-component programmes were excluded when the list was not isolable.

Comparator. No-activity, wait-list, measurement-only, hassles, neutral events, everyday-activity diaries, or an alternative frequency.

Outcomes. Primary: well-being (life satisfaction, happiness, or a subjective well-being composite) and positive affect. Secondary: depressive symptoms, signed so that a positive g means improvement.

Design and dates. Randomised or controlled trials. 1 January 2003 to 22 September 2026. Any language if extractable.

Frequency rule. Extract frequency exactly. Never infer it from “kept a journal” or “over two weeks”. Jackowska and colleagues (2016) is the cautionary case: a two-week gratitude writing trial whose abstract does not lock daily versus other schedules. It is cited in the narrative and kept out of any frequency subgroup.

Information sources and search

Sources: PubMed, PsycINFO, Scopus, Web of Science, Google Scholar (first 200 hits), CTRI, ClinicalTrials.gov, OSF preprints. Prior meta-analysis bibliographies (Davis 2016, Dickens 2017, Cregg and Cheavens 2021, Kirca 2023, Choi 2025) were hand-searched. The search is an update-and-restrict of those syntheses plus a 2023–2026 electronic pass. Dual independent screening of every database hit was not possible in this review. That limit is stated, not hidden.

Full search strings, reported verbatim.

PubMed. (gratitude[Title/Abstract] AND (journal*[Title/Abstract] OR list[Title/Abstract] OR listing[Title/Abstract] OR “three good things”[Title/Abstract] OR “3 good things”[Title/Abstract] OR “counting blessings”[Title/Abstract] OR “count your blessings”[Title/Abstract]) AND (randomi*[Title/Abstract] OR “controlled trial”[Title/Abstract] OR RCT[Title/Abstract])) AND (“2003/01/01”[Date - Publication] : “2026/12/31”[Date - Publication])

PsycINFO. TI,AB(gratitude AND (journal* OR list OR listing OR “three good things” OR “counting blessings”) AND (randomi* OR “controlled trial” OR RCT)) AND PY 2003-2026

Scopus. TITLE-ABS-KEY(gratitude AND (journal* OR list OR “three good things” OR “counting blessings”) AND (randomi* OR “controlled trial” OR rct)) AND PUBYEAR > 2002 AND PUBYEAR < 2027

Web of Science. TS=(gratitude AND (journal* OR list OR “three good things” OR “counting blessings”) AND (randomi* OR “controlled trial” OR RCT)) AND PY=(2003-2026)

Google Scholar. gratitude (journal OR list OR “three good things” OR “counting blessings”) randomi*

ClinicalTrials.gov. gratitude AND (journal OR “three good things” OR blessings)

CTRI. gratitude journal; three good things; counting blessings

OSF. gratitude AND (journal OR “three good things” OR “counting blessings”) AND random

Selection and extraction

Records were de-duplicated, screened on title and abstract, then read in full when they might be listing trials with a stated schedule. Extraction fields: n randomised and analysed; design; country; language of delivery; dose in sessions, items and minutes; frequency verbatim; duration in weeks; delivery channel; guidance (self, prompted, coached); “why it happened” present or absent; outcome measure; time-points; published d or g; control type. Anything that could not be verified against a DOI or a registry record was marked UNVERIFIED and kept out of any numerical pool.

PRISMA 2020 flow

This flow is a reconstruction informed by prior metas plus a 2023–2026 update. It is not a claim that four reviewers dual-screened 2,400 records in lockstep.

Table A. Search reconstruction counts

Stage	Count
Records identified across databases, registries, OSF and prior-meta bibliographies	~2,400
After de-duplication	~1,700
Title and abstract excluded (letter or visit only; youth-only; not a trial; frequency unstated; wrong outcome)	~1,520
Full-text assessed	~180

Stage	Count
Full-text excluded (frequency unstated; listing not isolable; no control; insufficient statistics; outside population)	~150
Eligible listing trials with explicit frequency, adult or mixed, 2003–2026	28 studies / 32 samples
With published or computable g on a primary outcome	18 samples
Head-to-head daily-versus-weekly RCTs	0
India adult listing RCTs meeting the quality bar	0
Listing trials used as the core narrative-and-table set in this paper	12

Analysis plan

The pre-specified model was random-effects REML, Hedges g, 95% confidence interval, I^2 , τ^2 , and a 95% prediction interval. Moderators: frequency (daily, weekly, other); duration in weeks; “why it happened”; control type. Small-study effects: Egger’s test and trim-and-fill. Sensitivity: leave-one-out. Software intention: metafor in R.

Two constraints changed the plan in practice. First, there is no published head-to-head daily-versus-weekly trial to pool. Second, most early trials report F tests without cell means and standard deviations. Pooling invented g values would violate the rule of this review. We therefore:

- Align the main pooled number with Choi and colleagues (2025) for gratitude interventions overall ($g = 0.19$, 95% CI 0.15 to 0.22) and for counting-blessings specifically ($g = 0.17$).
- Report extractable study-level g or d in a forest-style table.
- Treat frequency as a narrative-and-subgroup synthesis, not as a confirmatory contrast with a Q statistic we cannot honestly compute from complete data.
- Grade the frequency question as very low certainty rather than dress a missing comparison in forest-plot clothing.

A null is publishable. A missing contrast is also publishable. What is not publishable is a homemade weekly-versus-daily g built from unmatched trials that also differ in duration, control type, and country.

Risk of bias and certainty

Risk of bias was judged with RoB 2 domains: randomisation process; deviations from the intended intervention; missing outcome data; outcome measurement; selection of the reported result. Certainty used GRADE. All outcomes are self-reported. Participants are almost never blinded. Attrition in internet listing trials is often 40–70%. Those two facts cap certainty before heterogeneity is even considered.

06 Results

Study characteristics The twelve core listing trials are summarised below. Frequency is taken from the methods, not from later commentary.

Table 1. Study characteristics of frequency-specified gratitude-listing trials

Study	Country	n	Frequency	Duration	Control	Primary outcome
Emmons 2003 S1	USA	201 / ~192	Weekly	10 weeks	Hassles; events	Life appraisal; affect
Emmons 2003 S2	USA	166 / ~157	Daily	13 days	Hassles; comparison	Positive affect
Emmons 2003 S3	USA	65	Daily	21 days	Measurement	PA; sleep
Seligman 2005	Internet / USA	577 (6 arms)	Daily + why	7 days	Early memories	Happiness; depression
Chan 2013	Hong Kong	78 (40/38)	Weekly + meaning	8 weeks	Three misfortunes	LS; affect
Otsuka 2012	Japan	38 (19/19)	Weekly	4 weeks	Events list	Happiness; LS
Cheng 2015	Hong Kong	102	Twice weekly	4 weeks	Hassle; none	Depression; stress
Cunha 2019	Brazil	1,337 / 410	Daily	14 days	Hassles; events	Affect; happiness; LS; dep.
Manthey 2016	Germany	~450	Weekly, guided	8 weeks	To-do list	Subjective well-being
Kaplan 2014	USA	67 subset	3× / week	2 weeks	Other work tasks	Positive affect
Gold 2023	USA	223 (116/107)	3 texts / week	3 weeks	Delayed start	PA; depression; LS
Walsh / Regan 2023	Australia	958	Daily	7 days	Daily activities	SWB; affect

Two further daily trials inform the narrative but sit just off the core table because they are replications or personality-moderated: Mongrain and Anselmo-Matthews (2012), a Seligman replication, and Sergeant and Mongrain (2011), which found that self-critical people gained and “needy” people did not.

Gander, Proyer, Ruch and Wyss (2013) tested a one-week three-good-things arm and a two-week arm inside a ten-arm internet trial (N = 622). Stretching the same list from one week to two did not produce a cleaner win. That is a duration signal, not a frequency signal, and it belongs in the plateau section.

Forest-plot data table Only published or directly computable contrasts are listed. Missing cells are missing. They are not imputed.

Table 2. Extractable standardised effects (listing versus named control)

Study	Freq.	Contrast	Outcome	g or d	n
Choi 2025 (all GI)	Mixed	GI vs control	Well-being	$g = 0.19 [0.15, 0.22]$	24,804
Choi 2025 (listing)	Mixed	Counting blessings	Well-being	$g = 0.17$	Subtype
Davis 2016	Mixed lists	vs alt-activity	Well-being	$d = 0.17 [0.09, 0.24]$	$k = 20$
Davis 2016	Mixed lists	vs measurement	Well-being	$d = 0.31 [0.04, 0.58]$	$k = 5$
Cunha 2019	Daily	vs neutral events	Positive affect	$d = 0.38$	153 vs 134
Cunha 2019	Daily	vs hassles	Positive affect	$d = 0.60$	153 vs 123
Cunha 2019	Daily	vs hassles	Depression	$d = 0.29$	153 vs 123
Cunha 2019	Daily	vs hassles	Life satisfaction	$d = 0.33$	153 vs 123
Mongrain 2012	Daily	vs early memories	Happiness 1 wk / 3 mo / 6 mo	$d = 0.15 / 0.22 / 0.16$	102 vs 81
Cheng 2015	2× / week	vs no diary, 3-mo FU	Depression	$d = -0.49$	~34 vs ~34
Cheng 2015	2× / week	vs no diary, 3-mo FU	Perceived stress	$d = -0.95$	same
Manthey 2016	Weekly	vs to-do list	SWB	$g \approx 0.18-0.26$	150 vs 150
Walsh / Regan 2023	Daily	lists vs activity diary	SWB	~0 (ns)	958, 6 arms
Chan 2013	Weekly	vs misfortunes	Positive affect	ns	40 vs 38
Otsuka 2012	Weekly	vs events	Happiness / LS	ns	19 vs 19
Gold 2023	3× / week	vs delayed start	Positive affect	small; faded by 3 mo	116 vs 107
Cregg 2021 (prior)	Mixed GI	GI vs control	Dep. / anxiety	$g = -0.29 / -0.23$ FU	$k = 27$; $N = 3,675$

The pattern is not mysterious. Against a hassles diary, listing looks good. Against a neutral or activity-matched diary, listing looks small. Against a psychologically active comparison, listing often looks like zero. That is the Davis finding, and the 2023 Australian trial is its large-sample confirmation for lists.

Moderator table**Table 3. Pre-specified moderators**

Moderator	What the eligible trials show	Poolable?
Frequency: daily	Small PA gains in Emmons S2, Cunha, Seligman/Mongrain; high-powered null for lists in Walsh/Regan (N = 958); durability in Seligman tied to self-continuation after seven assigned days.	Subgroup adequate for narrative; no head-to-head
Frequency: weekly	Life appraisal in Emmons S1; mixed LS in Chan; ns happiness/LS in Otsuka; small SWB in Manthey.	k of extractable g < 5 — narrative only
Frequency: 2–3× / week	Cheng twice-weekly: clear depression and stress drop at 3 months. Gold thrice-weekly: small PA bump that faded. Kaplan: PA up, NA not. Lyubomirsky 3× / week: no happiness gain (unpublished).	Do not collapse into “daily”
Duration in weeks	Daily trials last 1–3 weeks. Weekly trials last 4–10 weeks. Duration is confounded with frequency. Gander 2013: 14 days of three-good-things not cleaner than 7.	Cannot separate from frequency
“Why it happened”	Present in Seligman 2005 and Chan 2013. Seligman is the durability showcase. Not isolated in a factorial trial of listing.	Unpooled
Control type	Hassles inflate g. Neutral events shrink it. Activity-matched shrinks it further. Wait-lists inflate depression effects (Cregg).	Consistent across priors

Heterogeneity

Choi and colleagues reported a small overall effect with between-country differences that their moderators did not explain. Davis reported I^2 around 41% in the listing-dose subset. Cregg reported small, reasonably consistent depression effects that grew when the control was a wait-list. For our frequency question the honest heterogeneity statement is this: the trials do not estimate the same contrast. A fourteen-day Brazilian email list versus hassles is not a ten-week American class list versus events, and neither is a four-week Hong Kong work diary. A prediction interval around a homemade daily-versus-weekly difference would be theatre.

Publication bias

Egger tests on the large priors have not overturned the small positive mean. Choi reported robustness to small-study methods. The bias that matters here is not a missing funnel of tiny journals. It is a missing head-to-head trial, plus a popular book finding that never acquired a results table. Trim-and-fill cannot invent a weekly-versus-daily RCT.

Risk-of-bias summary

Table 4. RoB 2 pattern across the core set

Domain	Typical judgement	Why
Randomisation process	Low in post-2015 trials; some concerns in Emmons and Seligman	Early papers describe assignment thinly
Deviations from intended intervention	Some concerns	Self-guided writing; adherence self-reported
Missing outcome data	High in internet trials; some concerns elsewhere	Cunha 1,337 → 410; Seligman 577 → 411 full-follow-up completers; Gold enrolled 48% of those invited
Outcome measurement	Some concerns, universal	Unblinded self-report of mood and well-being
Selection of reported result	Some concerns	Multiple scales; frequency not always pre-specified as the contrast
Overall	Some concerns to high	No core trial is low risk across domains

Cheng and colleagues is the cleanest workplace trial: hospital RCT, three arms, follow-up to three months, a double-blind claim. Even there the outcomes are self-reported. Gold and colleagues used delayed-start randomisation and text prompts, which is closer to how a WhatsApp programme actually runs, and their between-group story was modest.

GRADE table

Table 5. Certainty of evidence

Question	Estimate	Downgrades	GRADE
Does listing raise well-being versus control?	$g \approx 0.17$ to 0.22	RoB (−1). Inconsistency mild.	Moderate vs wait-list/hassles; Low vs activity-matched
Does listing reduce depressive symptoms?	$g \approx -0.29$ post; -0.23 FU (Cregg, mixed GI)	RoB (−1). Indirectness for listing-only (−1). Wait-list inflation.	Low
Does weekly listing beat daily listing?	No published head-to-head. Davis days $p = .233$. Walsh/Regan daily lists null.	RoB (−1). Indirectness (−1). Imprecision (−1). Missing contrast (−1).	Very low
Can we generalise to Indian professionals?	Choi 2025: India among ns countries. CTRI: no adult listing RCT.	Indirectness (−2). RoB (−1).	Very low

07

Discussion

PULL QUOTE

Twice-weekly work lists for four weeks cut depressive symptoms by $d = -0.49$ in Hong Kong hospital staff. That is the strongest professional-sample signal we have.

What the number means

A Hedges g of 0.19 is a small effect. In plain language: if you line up 100 people who listed blessings and 100 who listed yesterday's errands, only a modest extra slice of the first line will stand on the happier side of the median. Exercise, for depression, is larger. A good night's sleep is larger. A raise, a friend, and an hour less of commute are larger. Gratitude listing is a cheap attention drill that beats doing nothing by a little, beats listing your grievances by a bit more, and does not reliably beat another structured writing task.

That is not a reason to throw the practice out. Small effects that cost three minutes and no equipment are how public-health programmes work. It is a reason to stop selling the practice as a personality transplant.

The frequency question is sharper, and the answer is more disappointing. Weekly listing has never beaten daily listing in a published randomised trial, because that trial does not exist. The sentence everyone remembers — weekly beats daily — is a popular compression of an unpublished student contrast between one session a week and three. Davis and colleagues already asked whether more days of listing meant more well-being. They got $p = .233$. Choi and colleagues, with an order of magnitude more data, said sessions were not the lever. Gander and colleagues stretched three-good-things from one week to two and did not find a dose ladder. Gold and colleagues prompted three times a week for three weeks and watched the affect bump fade by month three. Komase and colleagues, reviewing workplace gratitude programmes, noted that lists done four times or fewer changed no outcome.

Put those facts in one room and a picture appears. It is not “weekly good, daily bad”. It is “a few prompted, specific lists over a few weeks beat a hassles diary and beat nothing; repeating the same list until it becomes homework does not buy more well-being; and three prompted sessions a week is already close to the point where the practice goes stale for a typical adult.”

Seligman's six-month curve is the exception that explains the rule. People were asked to list for seven nights and to write why each good thing happened. The ones whose happiness kept rising were the ones who continued of their

own accord. Continuation, not the assigned daily forever, predicted durability. In entrepreneurial language: the users who retained, retained. Forcing retention with a streak counter is not the same study.

Hedonic adaptation is the mechanism everyone reaches for, and it is not a bad one. The brain stops tagging a repeated stimulus as news. In machine-learning terms the model overfits the training set of “my family, my health, my job” and stops generalising to the new event at 4 p.m. By the third week of a daily streak, a Gurugram analyst types the same three items without looking up from her tea. The fixes are plain. Change the prompt — “one good thing a colleague did today” — so the list has to look somewhere new. Ask whether the list still means something or has become a rule to obey. Keep the core move, looking for what went well, and drop the rigidity. A list kept only to protect a streak breeds boredom.

There is a second mechanism, less quoted. Listing hassles makes people feel worse. A slice of the classic “gratitude works” literature is “rumination on burdens does not”. Cunha and colleagues are honest about this. Against hassles, daily listing won on positive affect at $d = 0.60$. Against neutral events, the same contrast was $d = 0.38$, and happiness, life satisfaction and depression did not separate cleanly from the neutral diary. Martínez-Martí, Avia and Hernández-Lloreda, replicating Emmons in Spain, watched the post-test affect gap collapse once they modelled the pre-test and the follow-up. The hassles group had dropped. The blessing group had not soared.

A third mechanism is interpersonal. Walsh, Regan and Lyubomirsky assigned 958 Australian adults to letters, essays, constrained lists, unconstrained lists, or an activity diary, every day for a week. Long-form writing beat lists. Constrained lists did not beat the activity diary on well-being. Unconstrained lists beat the diary on gratitude and positive affect. Letters produced the most elevation and the most indebtedness. If the goal is a feeling of thanks, a free list can work. If the goal is a well-being bump large enough to see in a noisy week, a list is a weak instrument next to a letter. That is a format finding, not a frequency finding, but it should humble anyone whose entire programme is three bullets on a phone.

Limitations

This review is not a dual-screened Cochrane product. The flow counts are a reconstruction. Several classic trials do not publish the cell statistics needed for a clean Hedges g . The weekly subgroup of extractable effects is smaller than five on any single outcome, so we did not force a weekly-versus-daily Q test. Lyubomirsky’s frequency contrast remains unpublished as a results paper and was not pooled. Workplace samples are few. Indian adult listing samples that meet the quality bar are none. Self-report and attrition bias sit on every estimate. We may have missed a 2025–2026 preprint. We did not name commercial instruments in the running text; the underlying papers used standard well-being and depression scales, and those choices are a source of heterogeneity we cannot fully model.

WHAT THIS MEANS**What this means for an Indian workplace programme**

A Noida sales manager survives the commute, clears the family WhatsApp group, and sits down to dinner. The app pings: time for the daily ten-minute journal. It loses to all three. The imported playbook says “streak”. The evidence says “a few specific lists, for a few weeks, then a stop-rule”.

The nearest professional analogue is Cheng’s Hong Kong hospital staff, who listed good things from work twice a week for four weeks. Three months later their depressive symptoms and stress had measurably dropped. Otsuka’s Japanese local-government staff wrote a weekly list; their happiness and life satisfaction did not move. Gold’s American health-care workers got a text prompt three times a week; their positive affect rose a little, then faded.

Choi’s 2025 country analysis put India among the places where gratitude programmes, as currently designed, did not lift well-being. So do not copy a Californian daily streak. Prompt one to three lists a week. Picture a shift supervisor in Nagpur. Twice a week a WhatsApp prompt arrives in Marathi: three good things, and an optional line on why one happened. She types them on the bus home. After four weeks, a short well-being check. When her list starts repeating “family, health, job”, the prompt changes or pauses, and a human coach is a tap away. That is a programme. That is not a diagnosis.

Research gaps, and the India gap named

The field does not need another student sample listing five blessings against a hassles diary. It needs three things.

First, a registered, adequately powered trial that randomises daily listing, thrice-weekly listing, and weekly listing, with the same item count, the same “why” instruction, the same control, and follow-up to three and six months. Until that trial exists, every “weekly beats daily” slide is a citation of a ghost.

Second, a factorial test of “why it happened”. Seligman’s durability may live in the causal sentence, not in the nightly schedule. Nobody has isolated that sentence inside a listing trial of professionals.

Third, India. CTRI returned no adult listing-only RCT that meets this review’s bar. Classroom and adolescent studies exist. An IT-sector correlational paper exists. A mixed gratitude programme for rural ASHAs exists and is not listing. Choi and colleagues’ country-level null for India is an empirical warning, not a cultural insult. Collectivist thanks is often duty-tinged. A 2024 Indian paper on target-person found that listing any person of one’s choice beat listing a family member — family gratitude can flatten affect when it feels like obligation.

An EmotionApp-style trial should therefore randomise frequency (weekly versus three times weekly) inside Indian working adults, deliver on WhatsApp, write in the person's language, keep India-resident data, and pre-specify well-being, positive affect and depressive symptoms at four weeks and three months. Until that trial is run, every dose recommendation for Indian professionals is an extrapolation. Extrapolation is allowed. Pretending it is local evidence is not.

08

FAQ

Q1 How often should I write three good things if I want a mood lift?

Three times a week, for four weeks, is the honest workplace dose. Daily lists help some people and bore others. Weekly lists help some people and under-dose others. No published trial has crowned a winner. Gains show up by week two. $g \approx 0.19$ is the well-being bump you should expect, not a personality change.

Q2 Does weekly gratitude beat a daily journal?

Not in any published head-to-head trial. Zero such trials exist. The famous finding compared one session a week with three, in students, and was never issued as a standalone results paper. Davis and colleagues found that extra days of listing did not raise well-being ($p = .233$).

Q3 Will this help if I already feel low?

A little, on average. Cregg and Cheavens pooled gratitude programmes and found $g = -0.29$ on depressive and anxiety symptoms after the programme, shrinking to $g = -0.23$ later. That is a small shift. It is not a substitute for clinical care. If mood is stuck or unsafe, speak to a clinician. In India, iCall (9152987821) and the Tele-MANAS line (14416) are starting points.

Q4

Why did a big Australian trial find that gratitude lists did nothing?

Because the control was another daily writing task, not a blank page. Nine hundred and fifty-eight adults wrote every day for a week. Constrained lists did not beat an activity diary on well-being. Letters did. Format beat frequency. A list of three is a weak instrument next to a paragraph to a real person.

Q5

What should an Indian company actually run?

A four-week, WhatsApp-first block. One to three prompted lists a week. Three specific items. Optional “why”. The person’s language. A pause button when the list repeats itself. A human coach if the practice goes stale. Measure well-being at week four. Do not install a 365-day streak and call it science. Choi and colleagues already found no country-level well-being lift for gratitude programmes in Indian samples.

09

References

Primary sources only. Each entry carries a DOI.

- Chan, D. W. (2013). Counting blessings versus misfortunes: Positive interventions and subjective well-being of Chinese school teachers in Hong Kong. *Educational Psychology*, 33(4), 504–519. <https://doi.org/10.1080/01443410.2013.785046>
- Cheng, S.-T., Tsui, P. K., & Lam, J. H. M. (2015). Improving mental health in health care practitioners: Randomized controlled trial of a gratitude intervention. *Journal of Consulting and Clinical Psychology*, 83(1), 177–186. <https://doi.org/10.1037/a0037895>
- Choi, H., Cha, Y., McCullough, M. E., Coles, N. A., & Oishi, S. (2025). A meta-analysis of the effectiveness of gratitude interventions on well-being across cultures. *Proceedings of the National Academy of Sciences*, 122(28), e2425193122. <https://doi.org/10.1073/pnas.2425193122>
- Cregg, D. R., & Cheavens, J. S. (2021). Gratitude interventions: Effective self-help? A meta-analysis of the impact on symptoms of depression and anxiety. *Journal of Happiness Studies*, 22, 413–445. <https://doi.org/10.1007/s10902-020-00236-6>
- Cunha, L. F., Pellanda, L. C., & Reppold, C. T. (2019). Positive psychology and gratitude interventions: A randomized clinical trial. *Frontiers in Psychology*, 10, 584. <https://doi.org/10.3389/fpsyg.2019.00584>
- Davis, D. E., Choe, E., Meyers, J., Wade, N., Varjas, K., Gifford, A., Quinn, A., Hook, J. N., Van Tongeren, D. R., Griffin, B. J., & Worthington, E. L., Jr. (2016). Thankful for the little things: A meta-analysis of gratitude interventions. *Journal of Counseling Psychology*, 63(1), 20–31. <https://doi.org/10.1037/cou0000107>
- Dickens, L. R. (2017). Using gratitude to promote positive change: A series of meta-analyses investigating the effectiveness of gratitude interventions. *Basic and Applied Social Psychology*, 39(4), 193–208. <https://doi.org/10.1080/01973533.2017.1323638>

- Emmons, R. A., & McCullough, M. E. (2003). Counting blessings versus burdens: An experimental investigation of gratitude and subjective well-being in daily life. *Journal of Personality and Social Psychology*, 84(2), 377–389. <https://doi.org/10.1037/0022-3514.84.2.377>
- 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>
- Gold, K. J., Dobson, M. L., & Sen, A. (2023). “Three good things” digital intervention among health care workers: A randomized controlled trial. *Annals of Family Medicine*, 21(3), 220–226. <https://doi.org/10.1370/afm.2963>
- Jackowska, M., Brown, J., Ronaldson, A., & Steptoe, A. (2016). The impact of a brief gratitude intervention on subjective well-being, biology and sleep. *Journal of Health Psychology*, 21(10), 2207–2217. <https://doi.org/10.1177/1359105315572455>
- Kaplan, S., Bradley-Geist, J. C., Ahmad, A., Anderson, A., Hargrove, A. K., & Lindsey, A. (2014). A test of two positive psychology interventions to increase employee well-being. *Journal of Business and Psychology*, 29, 367–380. <https://doi.org/10.1007/s10869-013-9319-4>
- Kirca, A., Malouff, J. M., & Meynadier, J. (2023). The effect of expressed gratitude interventions on psychological wellbeing: A meta-analysis of randomised controlled studies. *International Journal of Applied Positive Psychology*, 8, 63–86. <https://doi.org/10.1007/s41042-023-00086-6>
- Komase, Y., Watanabe, K., Hori, D., Nozawa, K., Hidaka, Y., Iida, M., Imamura, K., & Kawakami, N. (2021). Effects of gratitude intervention on mental health and well-being among workers: A systematic review. *Journal of Occupational Health*, 63(1), e12290. <https://doi.org/10.1002/1348-9585.12290>
- Lyubomirsky, S., Sheldon, K. M., & Schkade, D. (2005). Pursuing happiness: The architecture of sustainable change. *Review of General Psychology*, 9(2), 111–131. <https://doi.org/10.1037/1089-2680.9.2.111>
- Manthey, L., Vehreschild, V., & Renner, K.-H. (2016). Effectiveness of two cognitive interventions promoting happiness with video-based online instructions. *Journal of Happiness Studies*, 17, 319–339. <https://doi.org/10.1007/s10902-014-9596-2>
- Martínez-Martí, M. L., Avia, M. D., & Hernández-Lloreda, M. J. (2010). The effects of counting blessings on subjective well-being: A gratitude intervention in a Spanish sample. *Spanish Journal of Psychology*, 13(2), 886–896. <https://doi.org/10.1017/S1138741600002535>
- Mongrain, M., & Anselmo-Matthews, T. (2012). Do positive psychology exercises work? A replication of Seligman et al. (2005). *Journal of Clinical Psychology*, 68(4), 382–389. <https://doi.org/10.1002/jclp.21839>
- Otsuka, Y., Hori, M., & Kawahito, J. (2012). Improving well-being with a gratitude exercise in Japanese workers: A randomized controlled trial. *International Journal of Psychology and Counselling*, 4(7), 86–91. <https://doi.org/10.5897/IJPC11.031>
- Regan, A., Walsh, L. C., & Lyubomirsky, S. (2023). Are some ways of expressing gratitude more beneficial than others? Results from a randomized controlled experiment. *Affective Science*, 4, 72–81. <https://doi.org/10.1007/s42761-022-00157-y>
- 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>
- Sergeant, S., & Mongrain, M. (2011). Are positive psychology exercises helpful for people with depressive personality styles? *Journal of Positive Psychology*, 6(4), 260–272. <https://doi.org/10.1080/17439760.2011.553824>

About the author

ABOUT THE AUTHOR

Dr Atul Aswani; Mumbai Psychiatrist and trained as a Results Coach
Atul Aswani is a Mumbai Psychiatrist and trained as a Results Coach. He builds EmotionApp, a non-clinical emotional-fitness coaching 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, under ISO 9001 and ISO 13485. He writes evidence reviews so programmes ship the dose the studies actually support.

HOW TO CITE

Aswani A. Three Good Things, Twenty Years Later: A Meta-Analysis of Gratitude Dose — Daily versus Weekly. EmotionApp Research Paper 9. Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/gratitude-dose-daily-versus-weekly>

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.