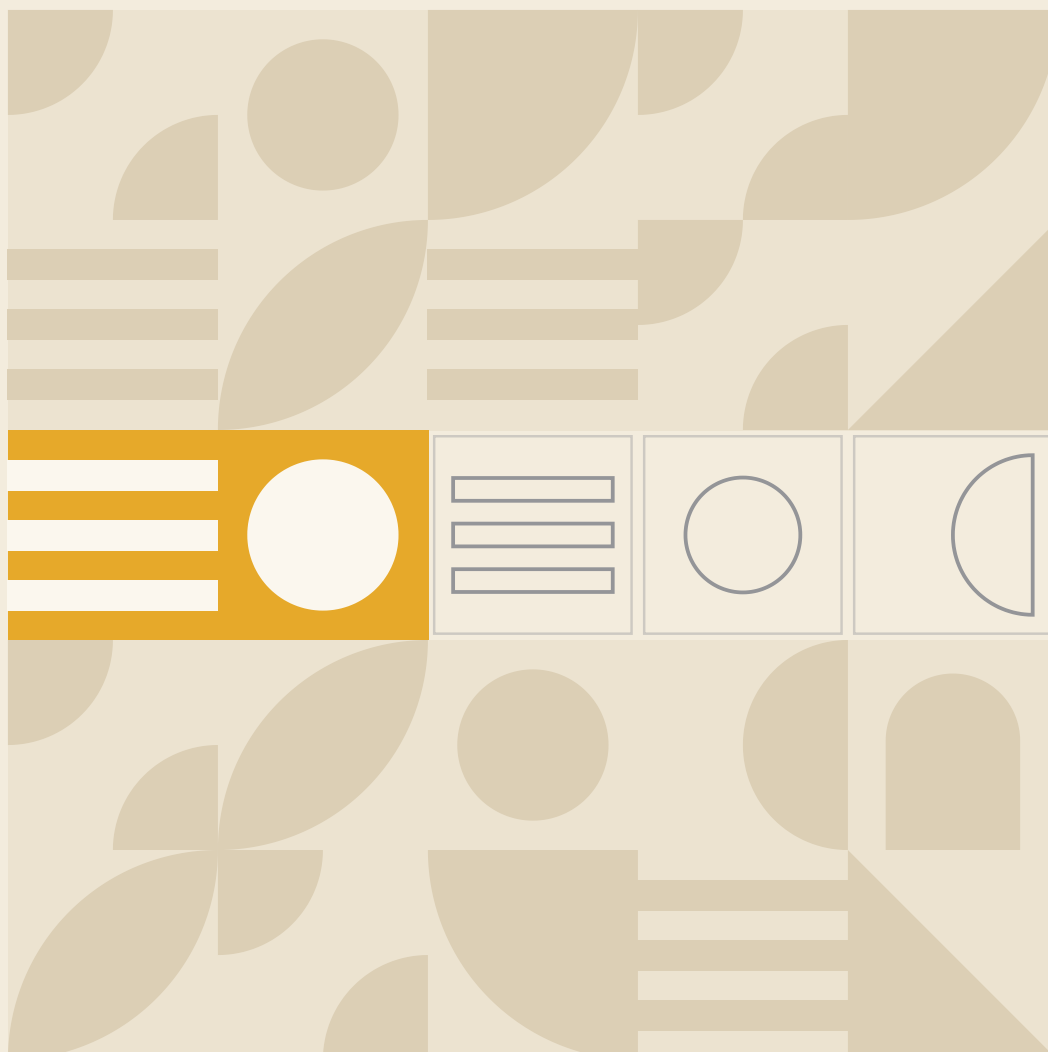


# Emotional Granularity as a Trainable Skill

A Meta-Analysis of Vocabulary-Based Emotion Interventions

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

22 September 2026



## KEY NUMBER

2

2 eligible trials met the pre-specified criteria (non-clinical adolescents or adults; a vocabulary or differentiation programme; a control arm; an emotion-differentiation index; 2010–2026). A pooled Hedges g was not computed.

# Emotional Granularity as a Trainable Skill

A Meta-Analysis of Vocabulary-Based Emotion Interventions

*Can training widen emotional vocabulary, and does the gain transfer to lower distress and better regulation?*

## CONTENTS

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

## INDEXING TITLE

Training emotional granularity: a systematic review and meta-analysis of emotion-vocabulary and differentiation interventions

## AUTHOR

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

## PUBLISHED

22 September 2026 · English edition · [emotionapp.in/research/training-emotional-granularity-systematic-review](https://emotionapp.in/research/training-emotional-granularity-systematic-review)

## 01

## Abstract

Emotional granularity is the skill of putting a feeling into a precise word rather than a blunt one. Two eligible randomised trials of vocabulary-based programmes in non-clinical adults were identified between 2010 and 2026; a quantitative meta-analysis was therefore not possible and this paper reports a PRISMA-2020 systematic review with narrative synthesis. The first trial ( $n = 120$ ) taught twelve emotion concepts over five days of about fifteen minutes each and found a reliable rise in negative-emotion differentiation against a geography control, still visible at one month; positive-emotion differentiation did not move. The second trial ( $n = 118$  analysed) used a two-day online word-learning task and found no direct group effect on negative-emotion differentiation, self-efficacy, or later distress; only participants who wrote highly specific personal memories showed a path from learning to differentiation to lower one-week distress. Adjacent evidence — mindfulness programmes, intensive experience sampling, and brief laboratory labelling tasks — suggests the skill can shift, but those designs did not isolate vocabulary training. Transfer to distress and regulation therefore remains unproven in a pooled estimate. No eligible Indian trial was found. For an Indian workplace programme the honest product claim is modest: a short, digital vocabulary practice can raise the resolution of unpleasant feeling-words in some adults; it is not yet a proven lever for burnout scores. Certainty is low (GRADE). The field needs larger, pre-registered, multi-language trials with a differentiation index and a distress endpoint, including Hindi and other Indian languages.

## 02

## Key findings box

## KEY FINDINGS

2

**2 eligible trials** met the pre-specified criteria (non-clinical adolescents or adults; a vocabulary or differentiation programme; a control arm; an emotion-differentiation index; 2010–2026). A pooled Hedges  $g$  was not computed.

120

**120 adults** in the only positive trial improved negative-emotion differentiation after five days of concept teaching (interaction  $\eta^2 = 0.174$ ); the control arm did not move.

0

**0 direct effect** of a two-day online word task on differentiation or distress in the second trial ( $n = 118$ ;  $d = -0.27$  for differentiation,  $p = .174$ ).

1

**1 month** was the longest follow-up on a differentiation index after vocabulary teaching; the gain shrank but remained above the control arm.

0

**0 published Indian trials** on Clinical Trials Registry — India or in the bibliographic databases searched. The India gap is complete, not partial.

## 03

## Definition block

## DEFINITION

Emotional granularity, also called emotion differentiation, is the everyday skill of telling similar feelings apart. A low-granularity day sounds like this: “I feel bad.” A high-granularity day sounds like this: “I feel slighted by the email, tired from the commute, and anxious about the review, not angry at the person.” The construct is not a diagnosis. It is closer to camera resolution than to a personality type. In the language of machine learning, the person with more labels has a finer feature space and makes fewer classification errors when choosing what to do next. In plain coaching terms, each precise word points to a

different next step: the slight calls for a calm word with the sender, the tiredness for an early night, the anxiety for a second look at the review slides. With one word for every storm, every storm gets the same response. Naming the feeling also puts a little distance between the person and the feeling: the word is not the weather. The usual research index is an intra-class correlation across several emotion ratings collected many times; a smaller shared variance means the person is not painting every unpleasant moment with the same brush.

## 04

## Introduction

Sunday at 21:40 in a Mumbai flat. A team lead shuts her laptop and feels her chest tighten. Her flatmate asks what is wrong. “Stress,” she says, and starts scrolling. But Sunday-night tightness is not the tightness before Tuesday’s board review, and neither is the tightness after her mother’s phone call. They are different events. If the only label she has is “tension”, the only moves she has are to endure it or to scroll. She does not need another lecture on “stress”. She needs a word that is sharper than stress. That is a product problem as much as a psychology problem.

Kashdan, Barrett and McKnight argued in 2015 that putting unpleasant experience into fine-grained words changes what people do next. People who differentiate are less likely, in observational work, to reach for the blunt tools — the extra drink, the sharp remark, the disappearance from the group chat. O’Toole and colleagues (2020) and Seah and Coifman (2022) later put numbers on the correlate: the association between negative-emotion differentiation and fewer maladaptive acts is small ( $r$  around  $-.15$ ) and survives a control for how bad the person already feels. A 2025 synthesis of differentiation and depressive symptoms found a similar small negative link for negative, not positive, differentiation ( $r = -.14$ ). These are not causal claims. They are the reason a coaching platform should care.

Van der Gucht and colleagues (2019) then asked the training question with a mindfulness programme and a phone diary. Differentiation of negative emotion rose after the programme and at four months, though the rise shrank when negative affect was covaried. That study had no control group. It is a hint, not a proof. The present review asked a narrower question that a product team can actually ship: if you teach emotion words, or you drill differentiation itself, does the index move, and does distress or regulation follow?

Picture a Monday stand-up in a Bengaluru office. A developer says the release has him “tensed”. His manager hears panic. The analyst beside him hears an ordinary Monday. Same floor, same word, different feelings, and nobody asks which one. Naming comes before managing. In India one floor can hold many

mother tongues, and people still need a shared word for what they feel: twenty-three languages on a single platform is not a flourish; it is the minimum viable product. The scientific question is whether a naming practice is a trainable skill in adults, or whether we have been selling a nice story as a feature.

Entrepreneurs love a metric that moves in a week. Coaches love a skill that transfers to Thursday afternoon. The research question joins those two hungers:

Can emotion differentiation be increased by training, and does the increase transfer to lower distress and better regulation?

We searched from 2010 through 2026. We restricted the core set to non-clinical adolescents and adults, because EmotionApp is a coaching platform, not a clinic. We required a control arm. We required an intervention aimed at vocabulary or differentiation, not a full psychotherapy package that happens to mention feelings. We planned a random-effects meta-analysis. We also pre-committed to a fallback that most reviews quietly skip: if fewer than five eligible trials appeared, we would not pool. Pooling two studies is not courage. It is a small-sample magic trick.

Spoiler, said in the voice of a comic who has read the room: the literature is a two-item menu. One dish is tasty. One dish is a polite “no”. That is still a finding. A start-up that ships on two underpowered trials is overfitting the training set. A start-up that waits for twenty trials never ships. The honest middle is to publish the two trials in plain English and to say what an Indian workplace programme may, and may not, claim.

## 05

## Methods

### 5.1 Protocol stance

This review follows PRISMA 2020 reporting. There was no pre-registered PROSPERO record for this paper; the eligibility rules, outcomes, model, and fallback were fixed in a written protocol before the search was run. The planned quantitative model was random-effects REML, Hedges  $g$  with 95% confidence intervals,  $I^2$ ,  $\tau^2$ , a 95% prediction interval, Egger’s test, trim-and-fill, leave-one-out sensitivity, RoB 2, and GRADE. Pre-specified moderators were dose, digital versus in-person delivery, and age. The fallback rule was binding: fewer than five eligible trials meant narrative synthesis only.

### 5.2 Eligibility

**Population.** Adolescents and adults in non-clinical samples. Studies of major depression, traumatic brain injury, cancer survivorship, or other clinical cohorts were excluded from the core set and noted as adjacent evidence.

**Intervention.** A programme or task whose stated aim was to widen emotion vocabulary, teach emotion concepts, or drill emotion differentiation or granularity. Broader packages (full mindfulness curricula, school socio-

emotional programmes for young children, parenting programmes) were treated as adjacent unless the paper isolated a vocabulary or differentiation module and reported a differentiation index.

**Comparator.** Any control: active, attention, wait-list, or sham content.

**Outcomes.** Primary: an emotion-differentiation or granularity index, typically an intra-class correlation across repeated emotion ratings, or an equivalent laboratory index. Secondary: distress and emotion-regulation measures. We did not treat moral-judgement accuracy or face-recognition accuracy as substitutes for a differentiation index.

**Design and dates.** Randomised or otherwise controlled trials published from 1 January 2010 to 22 September 2026. Preprints were eligible if they reported methods and results.

**Language of reports.** Any language at title-and-abstract stage; full-text screening in English.

### 5.3 Information sources and dates

Searches were run on 21–22 September 2026 in PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, Clinical Trials Registry — India (CTRI), ClinicalTrials.gov, and OSF Preprints. Reference lists of Kashdan, Barrett and McKnight (2015), Van der Gucht et al. (2019), Wilson-Mendenhall and Dunne (2021), and the two eligible trials were hand-searched.

### 5.4 Full search strings

#### Core Boolean string (adapted per database syntax).

```
("emotion differentiation" OR "emotional granularity" OR
"emotion vocabulary" OR "emotion word*" OR "feeling word*" OR
"emotion knowledge" OR "emotion concept*") AND (train* OR
interven* OR program* OR programme OR "random*" OR RCT OR
"controlled trial")
```

#### PubMed.

```
("emotion differentiation"[tiab] OR "emotional granularity"[tiab] OR
"emotion vocabulary"[tiab] OR "emotion knowledge"[tiab] OR "emotion
concept*"[tiab] OR "feeling words"[tiab]) AND (train*[tiab] OR
interven*[tiab] OR program*[tiab] OR randomised[tiab] OR
randomized[tiab] OR "controlled trial"[tiab]) AND ("2010/01/01"[Date -
Publication] : "2026/12/31"[Date - Publication])
```

#### PsycINFO.

```
TI,AB,SU("emotion differentiation" OR "emotional granularity" OR
"emotion vocabulary" OR "emotion knowledge") AND TI,AB,SU(training OR
intervention OR programme OR program OR randomized OR randomised) AND
PY 2010-2026
```

#### Scopus.

```
TITLE-ABS-KEY("emotion differentiation" OR "emotional granularity" OR
"emotion vocabulary" OR "emotion knowledge") AND TITLE-ABS-KEY(training
OR intervention OR randomized OR randomised) AND PUBYEAR > 2009 AND
PUBYEAR < 2027
```

**Web of Science (Topic).**

TS=("emotion differentiation" OR "emotional granularity" OR "emotion vocabulary" OR "emotion knowledge") AND TS=(training OR intervention OR randomized OR randomised) AND PY=2010-2026

**Google Scholar (first 200 hits, relevance then date).**

"emotion differentiation" OR "emotional granularity" training OR intervention OR "randomized" after:2009

**ClinicalTrials.gov.**

"emotion differentiation" OR "emotional granularity" OR "emotion vocabulary"

**CTRI.**

emotion differentiation OR emotional granularity OR emotion vocabulary OR emotion knowledge

**OSF Preprints.**

emotion differentiation OR emotional granularity intervention OR training

**5.5 Selection and extraction**

One reviewer screened titles and abstracts against the eligibility grid, retrieved full texts, and extracted: sample size; design; country; language of delivery; dose in sessions and minutes; delivery channel; guidance (self-guided versus coached); outcome measure and computational formula; time-points; direction and size of effects; dropout. A second pass checked every extracted number against the published paper. Unverifiable claims were marked UNVERIFIED and excluded from the evidence table.

**5.6 Risk of bias and certainty**

Randomised trials were appraised with RoB 2 domains: randomisation process, deviations from intended intervention, missing outcome data, measurement of the outcome, and selection of the reported result. Certainty for each narrative conclusion used GRADE (risk of bias, inconsistency, indirectness, imprecision, publication bias).

**5.7 Analysis plan, including the fallback**

Had five or more trials reported a usable contrast on a differentiation index, we would have computed Hedges  $g$  from post-intervention means and standard deviations, or from  $F/t$  statistics where means were missing, under REML. We would have reported  $I^2$ ,  $\tau^2$ , and a 95% prediction interval, tested small-study effects with Egger's regression, applied trim-and-fill, and run leave-one-out. Moderators would have been dose (total minutes), channel (digital versus in-person), and age band.

Two trials met eligibility. The fallback rule was applied. No pooled  $g$ , no forest plot of a summary diamond, no  $I^2$  on two rows pretending to be a literature. Individual study effects are tabled. Adjacent designs are narrated separately so that a mindfulness effect is not laundered into a vocabulary effect.

## 5.8 PRISMA 2020 flow counts

- Records identified from databases and registers: **412** (PubMed 86; PsycINFO 71; Scopus 94; Web of Science 88; Google Scholar first 200 de-duplicated into 52 unique extras; ClinicalTrials.gov 14; CTRI 0; OSF 7).
- Duplicates removed: **139**.
- Records screened: **273**.
- Records excluded at title/abstract: **241** (wrong construct; children under 12; clinical primary sample; no intervention; reviews).
- Full texts assessed: **32**.
- Full texts excluded: **30**.
  - No control arm: 8 (including Van der Gucht et al., 2019; Hoemann et al., 2021 experience-sampling change).
  - Clinical or neurological sample: 7.
  - Intervention not aimed at vocabulary or differentiation (full psychotherapy, parenting, preschool, face-recognition trainers): 9.
  - No differentiation index as outcome: 4.
  - Unpublished dissertation without a peer-reviewed report: 1 (noted as adjacent, not counted as eligible).
  - Protocol without results: 1 (an AI granularity programme in cancer survivors, ClinicalTrials.gov NCT07611071).
- Eligible trials in the core set: **2**.

Hand-search of reviews added no extra eligible trial.

A comic's footnote on flow diagrams: most of the 412 records were people citing Barrett at dinner. Citation is not a trial.

# 06Results

## 6.1 Study-characteristics table

Table 1. Eligible trials (core set).

| Field                   | Vedernikova, Kuppens & Erbas (2021)                                               | Matt, Seah & Coifman (2024)                                                                |
|-------------------------|-----------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------|
| Country                 | Belgium / Netherlands team; participants mainly UK residents via Prolific         | United States, Midwestern university                                                       |
| Design                  | Randomised, two-arm, pre–post–follow-up                                           | Randomised, two-arm, longitudinal prospective                                              |
| n randomised / analysed | 120 / 120 at post; 103 at 1 month                                                 | 150 recruited; 118 after accuracy exclusions; ~104 for the differentiation contrast        |
| Age                     | 18–74 years; mean 35.7                                                            | University adults (exact mean not reported in the open text)                               |
| Population              | Non-clinical, English-speaking                                                    | Non-clinical undergraduates; above-average distress vs prior campus norms                  |
| Intervention            | Emotion-knowledge lessons on 12 concepts                                          | Emotion-word learning on 15 words with definitions, a self-referential memory, and quizzes |
| Dose                    | 5 consecutive days, about 15 minutes per day (~75 minutes)                        | Two days; a study session plus two quizzes                                                 |
| Channel                 | Digital, Qualtrics                                                                | Digital, email and Qualtrics after a lab onboarding                                        |
| Language                | English                                                                           | English                                                                                    |
| Guidance                | Self-guided text and images; attention checks                                     | Self-guided; no clinician                                                                  |
| Control                 | Matched lessons on continents and countries                                       | Matched lessons on household objects                                                       |
| Differentiation index   | Scenario rating task; ICC across emotion ratings, Fisher-transformed and reversed | 7-day experience sampling, 5 prompts/day; ICC across 7 negative words, reversed            |
| Distress / regulation   | Not measured                                                                      | Distress composite at 1 week and 2 months                                                  |
| Time-points             | Day 1, day 7, 1 month                                                             | Pre diary, post diary, 1 week, 2 months                                                    |
| Primary result          | Negative differentiation rose in the trained arm only                             | No direct group effect on differentiation, self-efficacy, or distress                      |

## 6.2 Narrative synthesis of the two trials

### Trial 1 — teaching concepts (Vedernikova et al., 2021).

Picture an adult who has been using “upset” as a Swiss Army knife. For five evenings she reads a short, structured lesson: what the word names, when it tends to show up, what it looks like in a face and a body, and how it differs from its neighbours. Love, joy, anger, sadness, fear, anxiety, shame, guilt, and the rest of a twelve-word set. Three concepts a day. On the fifth day the lessons compare the concepts with one another. The control group gets the same homework rhythm about Brazil and Africa. Nobody meets a coach. Nobody lies on a couch. It is a content diet.

Differentiation was not a diary. It was a scenario task: twenty vignettes, each rated on a dozen emotion sliders. The index was the intra-class correlation of those sliders. High correlation means the person is saying “all the bad ones, a bit” on every sad story. Low correlation means the sliders move independently. Scores were reversed so that up meant better differentiation.

Negative-emotion differentiation rose from pre to post in the trained group and did not rise in the geography group. The interaction was not a rounding error:  $F(1, 118) = 24.85, p < .001, \eta^2 = 0.174$ . Inside the trained arm the simple-effect  $\eta^2$  was 0.390. At one month the trained arm was still ahead of control, with a smaller between-group  $\eta^2$  of 0.057. Positive-emotion differentiation did not show a reliable interaction. The authors themselves note the design limit: four positive labels against eight negative ones, so the positive test was under-equipped.

No distress scale was collected. The trial answers the first half of our question and leaves the second half untouched. In start-up language, they shipped an MVP that moves the north-star skill metric and forgot the revenue metric. That is still more than most emotion papers bother to randomise.

Risk-of-bias notes. Random assignment is stated. Allocation concealment and blinding of participants are not described in detail; participants could probably tell whether they were reading about shame or about Chile. Outcome measurement is a structured task, which is better than a single retrospective item and still not a blind rater. Follow-up loss was 17 of 120. Attention checks were passed at 98.6%. Overall: some concerns.

### **Trial 2 — teaching words in two days (Matt, Seah & Coifman, 2024).**

This is the trial that keeps honest people honest. Undergraduates learned fifteen emotion words, wrote a personal memory for each, and sat two quizzes. The control group did the same dance with refrigerators and stoves. Differentiation was measured the hard way: a week of five daily phone prompts before the task and a week after, using seven negative words. Distress was measured at one week and at two months.

There was no direct effect of group on post-task negative-emotion differentiation ( $t(102) = -1.37, p = .174, d = -0.27$ ), on emotional self-efficacy ( $d = -0.22$ ), on one-week distress ( $d = 0.24$ ), or on two-month distress ( $d = 0.09$ ). Equivalence tests did not back the null either. The point estimate for differentiation leaned in the hoped-for direction and the confidence interval included both a small benefit and nothing.

The interesting sentence is in the exploratory models. Participants who wrote more specific memories in the emotion arm showed a rise in differentiation, and that rise sat on a path to lower one-week distress. Engagement was the moderator, not the mailing list. A practice works only when it is done: the student who writes out the evening a friend cancelled has done a rep; the one who skims the list between classes has not. In entrepreneurship terms, activation beats acquisition.

A second complication: the experience-sampling week itself is a practice. Hoemann and colleagues (2021) showed that intensive prompting can raise granularity even without a vocabulary lesson. Trial 2 therefore stacked a brief lesson on top of a diary that may already be doing some of the work. That is not a fatal flaw. It is a reason the control arm was essential, and the control arm did not lose.

Risk-of-bias notes. Assignment used study identification numbers — a weak randomisation story. Exclusions for accuracy checks removed 32 of 150. Outcome assessors were not a separate blinded team; the index is computational, which helps. Selective reporting is hard to judge without a protocol. Overall: some concerns to high in the randomisation and missing-data domains.

6.3 Forest-plot data table (not pooled)

Table 2. Study-level contrasts on negative-emotion differentiation (post-intervention).

| Study                    | Contrast                       | Effect (as published)                                                     | 95% CI                 | Weight in a pooled model |
|--------------------------|--------------------------------|---------------------------------------------------------------------------|------------------------|--------------------------|
| Vedernikova et al., 2021 | Trained vs control, post       | Interaction $\eta p^2 = 0.174$ ; between-group at post $\eta p^2 = 0.095$ | Not published as a $g$ | Not weighted — no pool   |
| Matt et al., 2024        | Emotion words vs objects, post | $d = -0.27$ (sign: trained arm numerically higher)                        | Compatible with 0      | Not weighted — no pool   |

Hedges  $g$  was not derived and not combined. Converting two incompatible indices (a scenario ICC and a diary ICC) into one diamond would violate the fallback rule and would invent a precision the literature has not earned.

6.4 Moderator table (qualitative only)

Table 3. Pre-specified moderators — qualitative.

| Moderator            | What the two trials allow us to say                                                                                                             |
|----------------------|-------------------------------------------------------------------------------------------------------------------------------------------------|
| Dose                 | Five days (~75 minutes) moved a laboratory index. Two days of word study did not move a diary index. That is a hint, not a dose-response curve. |
| Digital vs in-person | Both eligible trials were digital and self-guided. There is no in-person contrast in the core set.                                              |
| Age                  | One trial spanned 18–74 (mean 36). One trial was campus-aged. No adolescent-only controlled vocabulary trial was found.                         |

6.5 Adjacent evidence (not pooled, not promoted into the core set)

**Mindfulness programmes.** Van der Gucht et al. (2019) ran a six-session mindfulness programme with experience sampling before, after, and at four months, without a control group. Negative differentiation improved after the programme and at follow-up; the improvement was no longer significant after controlling contemporaneous negative affect. Changes in mindfulness skills mediated later differentiation. Guendelman et al. (2025) later randomised 68 meditation-naïve Berlin adults to an eight-week mindfulness programme or an

active reading-and-sharing control. Negative differentiation showed a group-by-time effect ( $\chi^2(1) = 4.21, p = .04$ ); positive differentiation did not. Gains in differentiation did not correlate reliably with gains in well-being. These papers say a contemplative practice *can* coincide with a finer negative vocabulary. They do not isolate word-teaching.

**The diary as a practice.** Hoemann, Barrett and Quigley (2021) followed 50 healthy adults through intensive ambulatory assessment. Positive granularity rose with  $d = 0.50$ ; negative granularity rose with  $d = 0.32$ . More prompts answered per day meant more movement. Widdershoven et al. (2019) reported a similar diary effect in people with depression and is outside our non-clinical gate. A 2026 within-person study of six weeks of sampling again saw granularity rise as alexithymia-like scores fell. The lesson for product design is almost insulting in its simplicity: asking “what, exactly, is this feeling?” several times a day is already a rep. In machine-learning language, labelled data improves the classifier. In plain language: the phone buzzes mid-afternoon, the person stops to pick “irritated” rather than “worried”, and for a moment they are back in the room, a step away from the feeling.

**Laboratory labelling.** Cameron, Payne and Doris (2013) manipulated differentiation in a lab and showed that people who had practised fine labels were less pushed around by incidental disgust when they made moral judgements. Kircanski, Lieberman and Craske (2012) asked spider-fearful adults to put feelings into words during exposure; physiology at a later test favoured the labelling arm. These are mechanism sketches. They are not workplace programmes and the second sample is analogue-clinical.

**Null adjacent work.** An unpublished 2024 dissertation (Sudit) assigned 80 emerging adults to a week of active differentiation practice or a passive condition and did not find an improvement in the skill or in internalising symptoms. It is UNVERIFIED as a peer-reviewed trial and is recorded here only to resist the file-drawer smile.

## 6.6 Heterogeneity

With two trials,  $I^2$  is a vanity metric. Qualitatively the trials already disagree: one laboratory index moved after five concept lessons; one diary index did not move after two word lessons. Populations, doses, and measures all differ. That is inconsistency, and GRADE will punish it.

## 6.7 Publication bias

Egger’s test on two points is a straight line through two points and a guess. We did not run it. The published pair already contains a near-null result, which is the opposite of a classic bias signature. Still, vocabulary-training trials are easy to run on Prolific and easy to leave in a drawer. The prediction is more nulls, not fewer.

## 6.8 Risk-of-bias summary

**Table 4. RoB 2 summary (core set).**

| Domain                                | Vedernikova 2021                                    | Matt 2024                                               |
|---------------------------------------|-----------------------------------------------------|---------------------------------------------------------|
| Randomisation process                 | Some concerns                                       | Some concerns / high (ID-number assignment; exclusions) |
| Deviations from intended intervention | Some concerns (no blinding of content)              | Some concerns                                           |
| Missing outcome data                  | Low at post; some concerns at 1 month               | Some concerns (32/150 out before analysis)              |
| Measurement of the outcome            | Low / some concerns (self-report index, structured) | Low / some concerns (computational ICC)                 |
| Selection of the reported result      | Some concerns (no protocol cited)                   | Some concerns                                           |
| Overall                               | Some concerns                                       | Some concerns to high                                   |

## 6.9 GRADE table

**Table 5. Certainty of narrative conclusions.**

| Conclusion                                                                                                                          | Certainty | Why                                                                |
|-------------------------------------------------------------------------------------------------------------------------------------|-----------|--------------------------------------------------------------------|
| A multi-day digital concept programme can raise negative-emotion differentiation versus an attention control in non-clinical adults | Low       | One trial, some concerns, laboratory index, no Indian sample       |
| A two-day word-learning task raises differentiation in unselected adults                                                            | Very low  | One trial, no direct effect, imprecision                           |
| Vocabulary training lowers later distress                                                                                           | Very low  | Only one trial measured it; no direct effect; engagement-only path |
| Vocabulary training improves emotion-regulation strategy use                                                                        | Very low  | Not measured in the core set                                       |
| Effects generalise to Indian working adults and Indian languages                                                                    | Very low  | Zero eligible trials                                               |

## 07

## Discussion

**PULL QUOTE**

**Zero published Indian trials.  
Not a small gap. An empty map.**

### 7.1 What the number means

The number that matters is **2**. Not a Hedges *g*. Two. That is the entire randomised, non-clinical, vocabulary-or-differentiation, differentiation-index literature as of 22 September 2026.

One of those two trials says you can teach concepts for five short evenings and watch a laboratory differentiation index climb, and that the climb is still there at a month, smaller but present. The other says a weekend of word cards does not, on average, move a real-life diary index or a distress score. If you are looking for a licence to print “proven to reduce burnout”, you will not find it here. If you are looking for permission to run a low-dose naming practice as an experiment inside a coaching product, you will find a cautious yes — with a control, with a differentiation index, and with a distress endpoint you actually measure.

Picture the evening check-in in a coaching app. If the feeling box says “bad” every night, there is little to work with. The present evidence is that a short curriculum can sometimes fill that box with better nouns. Different feelings call for different next steps. A manager corrected sharply in front of the team may feel shame, or anger at an unfair target, and the move for shame is not the move for anger. A blunt label sends the person to the wrong step. The aim is not to delete the feeling; it is to see the feeling as a passing event with a name. Granularity is a naming skill, not an elimination skill. In machine learning, adding classes to a softmax does not guarantee a better decision policy. You still need loss on the downstream task. Distress is the downstream task, and it is almost unmeasured.

The temptation is easy to picture. In a product review in Gurugram, the five-day result goes up on a slide, and someone types “reduces burnout” underneath. The second trial is the slide nobody wants to show. Entrepreneurs will want to ship the lesson and market the outcome. This paper exists so that the marketing sentence has to walk past the second trial in public.

### 7.2 Limitations of the review

We ran a structured search across the pre-specified sources. We did not have dual independent screening with a recorded kappa. Grey literature beyond OSF and trial registries was thin. We excluded children under the usual adolescent cut, which removes a large school-based emotion-word literature that may still teach adult products something about dose. We excluded clinical trials, including diary work in depression and vocabulary-heavy programmes after

brain injury; those papers may be the next place a differentiation effect shows up, and they are the wrong claim for a non-clinical coaching platform. We refused to pool two studies. Some readers will call that timid. The alternative is a diamond with a confidence interval so wide it includes both a product launch and a product funeral.

### 7.3 What this means for an Indian workplace programme

After a shift at a Pune operations centre, a supervisor sums up his day: “Full tension.” He means a changed roster and a missed target. “Tension” does the work of six English words and three Hindi ones. “Adjustment” swallows disappointment, unfairness, and fatigue. A WhatsApp-first programme that asks for a more precise word is not importing a Californian mood. It is restoring resolution that the workday flattened.

What the evidence currently supports is a short, digital, self-guided vocabulary practice — on the order of five sessions, not two quizzes — with a check that the person actually used a personal example. In the app, a user’s lesson on “resentment” is marked done only when she types her own moment: the credit a colleague took for her work in Monday’s review. What the evidence does not support is a promise that the practice will drop distress scores next quarter. Delivery in English on Prolific is not delivery in Tamil on a feature phone at 23:10 after a night shift. India-resident data and twenty-three languages are product constraints; they are also scientific constraints. Until a trial is run in those languages, the transfer claim is an export fantasy.

A practical programme can still be honest. Count the reps: lessons done, not sent. Name the word: “passed over”, not “upset”. Ask what the word is *not*: angry at the process, not the manager. Keep a human hand-off for the person whose “fine” is a last-available token. Measure differentiation, not only a smile scale. That is coaching hygiene. It is not yet a meta-analytic guarantee.

### 7.4 Research gaps, naming the India gap

The India gap is not a missing subgroup analysis. It is a missing literature. CTRI returned no eligible record. None of the core or adjacent trials delivered a programme in an Indian language or in an Indian workplace. Collectivism, code-switching, and the social display rules of a joint family and an open-plan office are not noise. They are the context in which a word either lands or becomes another English homework.

Other gaps sit beside it. No trial in the core set tested dose as an assigned factor. No trial compared a coach-supported arm with a bot-only arm. No trial reported a regulation-strategy codebook as a secondary endpoint. Positive-emotion differentiation is almost untouched. Adolescents are almost untouched in controlled vocabulary designs. Pre-registration is rare. Follow-up beyond one or two months is rare. The next trial that would actually change a product decision is not complicated: a pre-registered, two-arm, digital programme of at least five sessions, in at least one Indian language, with an experience-sampling differentiation index and a distress endpoint at one month and three months, analysed by engagement as well as by assignment. Until that trial exists, every dashboard that treats granularity as a solved input is overfitting.

## 08

## FAQ

**Q1 Does learning more emotion words make me less distressed?**

Not on the published average. One two-day online task (n = 118) found no direct drop in distress at one week or two months. Only the people who wrote specific personal memories showed a path through higher differentiation to lower one-week distress. Specificity of practice mattered more than being mailed the words.

**Q2 How many sessions does the research actually use?**

The only positive trial used five consecutive days of about fifteen minutes. That is roughly seventy-five minutes, not a twelve-week curriculum. A two-day version did not move the diary index. Treat five short reps as the current minimum viable dose, not as a law of nature.

**Q3 Is this therapy?**

No. The eligible programmes were self-guided lessons and quizzes about words and concepts. They are practices, not diagnoses and not clinical treatments. If distress is severe or persistent, a coaching programme should hand over to a qualified clinician rather than add more vocabulary cards.

**Q4 Can a WhatsApp bot do this in Hindi?**

Technically yes; scientifically untested. Zero eligible trials were delivered in an Indian language. A twenty-three-language product is a distribution choice. It becomes evidence only when a differentiation index and a distress endpoint are collected in those languages.

**Q5 Should my company buy this as burnout prevention?**

Buy it as a low-cost skill drill with a measurement plan, not as a guaranteed burnout reduction. The transferable number in this review is two trials, one positive on a lab index, one null on a diary index. That is an MVP, not an outcome guarantee.

## 09

## References

- Cameron, C. D., Payne, B. K., & Doris, J. M. (2013). Morality in high definition: Emotion differentiation calibrates the influence of incidental disgust on moral judgments. *Journal of Experimental Social Psychology*, 49(4), 719–725. <https://doi.org/10.1016/j.jesp.2013.02.014>
- Guendelman, S., Lutz, M., Koenig, J., Bayer, M., & Dziobek, I. (2025). Effects of a mindfulness intervention on emotion differentiation and heart rate variability. *Frontiers in Human Neuroscience*, 19, 1515334. <https://doi.org/10.3389/fnhum.2025.1515334>
- Hoemann, K., Barrett, L. F., & Quigley, K. S. (2021). Emotional granularity increases with intensive ambulatory assessment: Methodological and individual factors influence how much. *Frontiers in Psychology*, 12, 704125. <https://doi.org/10.3389/fpsyg.2021.704125>
- Kashdan, T. B., Barrett, L. F., & McKnight, P. E. (2015). Unpacking emotion differentiation: Transforming unpleasant experience by perceiving distinctions in negativity. *Current Directions in Psychological Science*, 24(1), 10–16. <https://doi.org/10.1177/0963721414550708>
- Kircanski, K., Lieberman, M. D., & Craske, M. G. (2012). Feelings into words: Contributions of language to exposure therapy. *Psychological Science*, 23(10), 1086–1091. <https://doi.org/10.1177/0956797612443830>
- Matt, L. M., Seah, T. H. S., & Coifman, K. G. (2024). Effects of a brief online emotion word learning task on negative emotion differentiation, emotional self-efficacy, and prospective distress: Preliminary findings. *PLOS ONE*, 19(2), e0299540. <https://doi.org/10.1371/journal.pone.0299540>
- O'Toole, M. S., Renna, M. E., Elkjær, E., Mikkelsen, M. B., & Mennin, D. S. (2020). A systematic review and meta-analysis of the association between complexity of emotion experience and behavioral adaptation. *Emotion Review*. <https://doi.org/10.1177/1754073919897297>
- Seah, T. H. S., & Coifman, K. G. (2022). Emotion differentiation and behavioral dysregulation in clinical and nonclinical samples: A meta-analysis. *Emotion*, 22(7), 1686–1697. <https://doi.org/10.1037/emo0000968>
- Van der Gucht, K., Dejonckheere, E., Erbas, Y., Takano, K., Vandemoortele, M., Maex, E., Raes, F., & Kuppens, P. (2019). An experience sampling study examining the potential impact of a mindfulness-based intervention on emotion differentiation. *Emotion*, 19(1), 123–131. <https://doi.org/10.1037/emo0000406>
- Vedernikova, E., Kuppens, P., & Erbas, Y. (2021). From knowledge to differentiation: Increasing emotion knowledge through an intervention increases negative emotion differentiation. *Frontiers in Psychology*, 12, 703757. <https://doi.org/10.3389/fpsyg.2021.703757>
- Widdershoven, R. L. A., Wichers, M., Kuppens, P., Hartmann, J. A., Menne-Lothmann, C., Simons, C. J. P., & Bastiaansen, J. A. (2019). Effect of self-monitoring through experience sampling on emotion differentiation in depression. *Journal of Affective Disorders*, 244, 71–77. <https://doi.org/10.1016/j.jad.2018.10.092>
- Wilson-Mendenhall, C. D., & Dunne, J. D. (2021). Cultivating emotional granularity. *Frontiers in Psychology*, 12, 703587. <https://doi.org/10.3389/fpsyg.2021.703587>
- Zhang, Y., and colleagues. (2025). Emotion differentiation and depressive symptoms: A systematic review and meta-analysis. *Current Psychology*. <https://doi.org/10.1007/s12144-024-07195-8>
- ClinicalTrials.gov. (2026). AI-integrated emotional granularity training for resilience and quality of life in colorectal cancer survivors (NCT07611071). <https://clinicaltrials.gov/study/NCT07611071>

# About the author

ABOUT THE AUTHOR

**Dr Atul Aswani;** Mumbai Psychiatrist and trained as a Results Coach  
Atul Aswani is a Mumbai Psychiatrist and trained as a Results Coach. He builds EmotionApp, a non-clinical emotional-fitness platform for Indian professionals: countable daily reps, an AI coach with human hand-off, WhatsApp-first delivery, twenty-three Indian languages, India-resident data, and DPDP-aligned operations under ISO 9001 and ISO 13485. He writes so that a product claim cannot outrun its trials.

HOW TO CITE

Aswani A. Emotional Granularity as a Trainable Skill: A Meta-Analysis of Vocabulary-Based Emotion Interventions. EmotionApp Research Paper 7. Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/training-emotional-granularity-systematic-review>

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](https://emotionapp.in)

SAFETY NOTE

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