

0.41 असर: नाम लिया तो तेज़ भावना हल्की पड़ी

क्या भावना को शब्द देने से वो सच में हल्की होती है? और क्या बारीक नाम से असर बढ़ता है?

Name It to Tame It, Tested: A Meta-Analysis of Affect Labelling on Emotional Intensity

डॉ. अतुल असवानी; मुंबई के साइकियाट्रिस्ट और ट्रेंड रिज़ल्ट्स कोच

22 September 2026



मुख्य आंकड़ा • Hedges g

0.41

95% CI 0.29 to 0.54

छह लैब प्रयोग, गैर-क्लिनिकल वयस्क (N = 261)। तेज़ नकारात्मक भावना पर Hedges g = 0.41 (95% CI 0.29 to 0.54)। नाम बनाम सिर्फ़ देखना।

0.41 असर: नाम लिया तो तेज़ भावना हल्की पड़ी

Name It to Tame It, Tested: A Meta-Analysis of Affect Labelling on Emotional Intensity

क्या भावना को शब्द देने से वो सच में हल्की होती है? और क्या बारीक नाम से असर बढ़ता है?

इस पेपर में

01 Abstract

02 Key findings box

03 Definition block

04 Introduction

05 Methods

06 Results

07 Discussion

08 FAQ

09 References

इंडेक्सिंग टाइटल

Affect labelling and emotional intensity: a meta-analysis of experimental and field studies

लेखक

डॉ. अतुल असवानी; मुंबई के साइकियाट्रिस्ट और ट्रेंड रिज़ल्ट्स कोच

प्रकाशित

22 September 2026 • हिंदी संस्करण • emotionapp.in/research/affect-labelling-emotional-intensity-meta-analysis

01

Abstract

रुकी हुई लोकल ट्रेन में फँसी एक एनालिस्ट फ़ोन के नोट्स में एक शब्द टाइप करती है: “घबराहट”।

यह पेपर पूछता है: भावना का नाम लेने से उसकी ताकत सच में गिरती है या नहीं।

छह लैब स्टडी में गैर-क्लिनिकल वयस्क थे (N = 261)। अफेक्ट लेबलिंग (Affect Labelling) से तेज़ नकारात्मक भावना की तीव्रता घटी। Hedges $g = 0.41$ (असर कितना बड़ा है, इसका नाप)। 95% CI 0.29 to 0.54। तुलना: सिर्फ़ तस्वीर देखना, नाम न लेना। मॉडल REML (एक सांख्यिकी तरीका जो स्टडी के फ़र्क को जोड़ता है) था। इस चुनी हुई सूची में $I^2 = 0\%$ और $\tau^2 = 0$ रहा। Leave-one-out में $g = 0.39$ से 0.48 तक रहा। शरीर के नतीजे अलग रखे गए। उन्हें एक अंक में नहीं जोड़ा। त्वचा की बिजली (skin conductance) और amygdala की स्टडी में असर अक्सर दिखा। दो डर वाली स्टडी में पसीना घटा। स्कोर नहीं घटा।

इमोशनल ग्रैन्युलैरिटी (Emotional Granularity) यानी बारीक नाम, को जोड़कर नहीं नापा जा सका। कोई स्टडी यह नहीं जाँचती कि बारीक शब्द से असर कितना बढ़ता है। पहली भाषा बनाम दूसरी भाषा भी नहीं जाँची गई। एक स्टडी में हल्की तस्वीर पर नाम लेने से दुख बढ़ा। फ़िल्ड सबूत पतला है। ट्विटर की एक बड़ी स्टडी में “I feel...” लिखने के बाद मूड तेज़ी से पलटा। भारत की कोई RCT नहीं मिली। RCT यानी ऐसी स्टडी जिसमें लोगों को लॉटरी की तरह दो ग्रुप में बांटा जाता है।

लैब के स्व-रिपोर्ट पर यकीन कम से मध्यम है। amygdala वाले पैटर्न पर यकीन मध्यम है। ग्रैन्युलैरिटी और भाषा पर यकीन बहुत कम है। भारतीय ऑफ़िस प्रोग्राम के लिए बात साफ़ है। तेज़ बुरी भावना पर छोटा, साफ़ नाम एक रोज़ का रेप है। यह जादू नहीं है। लोग अक्सर सोचते हैं नाम लेने से हालत बिगड़ेगी। तेज़ भावना वाली लैब स्टडी में स्कोर अक्सर उलटी तरफ़ गया।

कीवर्ड: affect labelling; emotion regulation; emotional intensity; skin conductance; amygdala; emotional granularity; India; systematic review; meta-analysis.

02

Key findings box

मुख्य नतीजे

- 6** छह लैब प्रयोग, गैर-क्लिनिकल वयस्क (N = 261)। तेज़ नकारात्मक भावना पर Hedges $g = 0.41$ (95% CI 0.29 to 0.54)। नाम बनाम सिर्फ़ देखना।
- 0.39** Leave-one-out में $g = 0.39$ से 0.48। कोई एक स्टडी नतीजा नहीं चला रही।
- 0.29** Levy-Gigi and Shamay-Tsoory (2022) में हल्की तस्वीर पर नाम से दुख बढ़ा ($g = -0.29$)। तेज़ तस्वीर पर दुख घटा ($g = 0.27$)। तीव्रता छोटी बात नहीं है।
- 2** दो एनालॉग-डर स्टडी (Kircanski et al., 2012, n = 88; Niles et al., 2015, n = 102) में डर के स्कोर में फ़र्क नहीं। त्वचा की बिजली में गिरावट साफ़ थी। शरीर और कहानी अलग चल सकते हैं।
- 0** भारत की योग्य स्टडी शून्य। पहली बनाम दूसरी भाषा की स्टडी शून्य। ग्रैन्युलैरिटी को जोड़कर नहीं नापा जा सका।

03

Definition block

परिभाषा

अफ़ेक्ट लेबलिंग मतलब इस वक्त की भावना को शब्द देना। अंदर देखो। नाम दो: गुस्सा, शर्म, जकड़न, निराशा, डर। बोलो, टाइप करो, या लिस्ट से चुनो। भावना से बहस मत करो। खुश खयाल से उसे मत बदलो। दिन का स्कोर मत दो। यह सबसे छोटी निगरानी है। मशीन लर्निंग में भी कच्चा डेटा सस्ता होता है। लेबल महंगा होता है। और काम करता है। कोचिंग की सीधी भाषा में: ध्यान दो और बताओ। कड़ा अप्रेज़ल पढ़ती एक मैनेजर मन में कहती है, “मुझे शर्म का अहसास हो रहा है”, न कि “मैं नाकारा हूँ”। कंपनी की भाषा में यह डैशबोर्ड की पहली लाइन है। यह पेपर पूछता है: क्या वह एक लाइन अगला नंबर बदलती है। भावना कितनी तेज़ है। शरीर कितना शोर कर रहा है।

04

Introduction

सुबह के 9.01, मुंबई का एक ऑफिस। स्टैंड-अप में एक सेल्स मैनेजर को पता चलता है कि क्लाइंट ने लॉन्च फिर टाल दिया। 9.17 पर स्लैक पर मैसेज आता है: क्यों? उसकी छाती कस जाती है। भारत के कामकाजी लोगों में भावना की कमी नहीं है। कमी है उसे रखने की जगह की।

लंच पर माँ का फ़ोन आता है, वे परेशान हैं। शाम तक बारिश से लोकल रुक जाती है। उसका दिमाग तीनों को एक फ़ोल्डर में डाल देता है। नाम: स्ट्रेस। फिर फ़ोल्डर जबड़े, पेट और रात 11 बजे के रिप्ले पर चला जाता है। उद्यमी यह बात जानते हैं। जो नाप नहीं सकते, उसे चला नहीं सकते। ज़्यादातर लोगों की तरह, उसका भीतरी जीवन भी बिना लेबल वाले डेटा पर चलता है।

कोलकाता में एक ऑपरेशंस हेड सोमवार को इनबॉक्स खोलती है। न पढ़े मेल का एक भारी ढेर। वह उन्हें फ़ोल्डर में बाँटती है: क्लाइंट, बिलिंग, बाद में। ढेर अब संभलने लायक हो जाता है। उलटा संकट रिव्यू कॉल में दिखता है। एक ही जोखिम के लिए हर कोई अलग शब्द बोलता है। कमरे को सिर्फ़ शोर सुनाई देता है। 2007 से भाव-विज्ञान यही छोटा सवाल घुमा रहा है। भावना से शब्द चिपकाओ, तो क्या वोल्टेज गिरता है?

Matthew Lieberman और साथियों ने लोगों को स्कैनर में बिठाया। भावना वाले चेहरे दिखाए। भावना का शब्द चुनने को कहा। amygdala की हलचल घटी। दाएँ ventrolateral prefrontal cortex की हलचल बढ़ी। दोनों उलटी दिशा में जुड़े। बीच में medial prefrontal cortex था (Lieberman et al., 2007)। पेपर कहावत बन गया: name it to tame it। कहावत कॉन्फ़िडेंस इंटरवल से तेज़ दौड़ती है।

दस साल बाद Torre and Lieberman (2018) ने इसे अंतर्निहित भावना नियंत्रण कहा। उन्होंने लेबलिंग की तुलना रीअप्रेज़ल से की। अनुभव, शरीर, दिमाग, व्यवहार। उन्होंने असर का एक अंक नहीं जोड़ा। 2026 में इलाहाबाद विश्वविद्यालय से Shukla and Mishra ने 32 स्टडी देखीं। फिर भी अंक नहीं जोड़ा। कहावत के पास अब भी नंबर नहीं था।

भारतीय ऑफ़िस में वह नंबर ज़रूरी है। इमोशनऐप इमोशनल फ़िटनेस (Emotional Fitness) के तरीक़े रोज़ के गिने हुए रेप्स में देता है। व्हाट्सऐप पर। 23 भारतीय भाषाओं में। डेटा भारत में रहता है। एआई कोच के पीछे इंसान भी है। अगर लेबलिंग एक रेप बनेगी, तो उसका साइज़ चाहिए। उसकी सीमा चाहिए। और यह सूची चाहिए कि वह क्या नहीं करती। स्कैनर एक बार चमका, इसलिए फ़ीचर शिप करना स्टार्टअप नहीं है। वह लॉगिन पेज वाला नारा है।

नीचे दो और सवाल हैं। पहला: क्या असर शरीर पर भी दिखता है, या सिर्फ़ स्कोर पर? त्वचा की बिजली और amygdala का जवाब “मुझे बुरा लग रहा है” से अलग चीज़ हैं। एक नई मैनेजर अपनी पहली बोर्ड प्रेजेंटेशन से निकलती है। मन कहता है, “ठीक गया।” नब्ज़ अब भी तेज़ है। पेट कसा हुआ है। विचार, भावना और नब्ज़ एक साथ नहीं चलते। यह फूट डेटा है, हार नहीं। दूसरा: क्या बारीक नाम ज़्यादा काम करता है? जो इंसान झुंझलाहट, अपमान और दहशत अलग बताता है, वह रोज़मर्रा में बेहतर संभालता है। यह अलग दावा है। कई लोग दोनों दावों को एक कर देते हैं। यह पेपर उन्हें अलग रखता है।

एक स्थानीय चुप्पी भी है। भारत में भावना पहचान के डेटासेट हैं। हिंदी लेक्सिकॉन हैं। फ़िल्म क्लिप हैं। इस खोज में कामकाजी वयस्कों पर अफेक्ट लेबलिंग की RCT नहीं मिली। 23 भाषाओं वाला प्रोडक्ट यह मान ले कि अंग्रेज़ी वाला असर हिंदी, तमिल, बांग्ला, मराठी में चल जाएगा। यह

ट्रांसफ़र स्टडी के बिना ट्रांसफ़र दावा है। यह सावधानी नहीं है। यह नाम वाला गैप है।

सवाल साफ़ है। क्या भावना का नाम लेने से तीव्रता और शारीरिक उत्तेजना घटती है? क्या इमोशनल ग्रैन्युलैरिटी इस असर को बदलती है? साल 2000 से 2026। मॉडल रैंडम-इफ़ेक्ट्स। स्व-रिपोर्ट और शरीर अलग। शून्य नतीजा भी छप सकता है। कहावत नतीजा नहीं है।

05

Methods

5.1 Protocol and reporting

यह रिव्यू PRISMA 2020 के हिसाब से है। प्रोटोकॉल पहले से रजिस्टर नहीं था। खोज 22 September 2026 को हुई। इसे बाद में बनी व्यवस्थित समीक्षा समझें। डेटाबेस के बिल्कुल सटीक हिट गिनकर नहीं गढ़े गए। वे गोल किए गए हैं।

5.2 Eligibility

जनसंख्या। 18 साल और ऊपर के वयस्क। गैर-क्लिनिकल। अगर स्टडी में मनोविकार का निदान शर्त था, तो वह बाहर। मकड़ी से डरने वाले छात्र और पब्लिक स्पीकिंग वाले चिंतित वयस्क संवेदनशीलता सेट में रहे। मुख्य स्व-रिपोर्ट पूल से बाहर।

हस्तक्षेप। अफेक्ट लेबलिंग: अपनी भावना या तस्वीर की भावना को शब्द देना। शब्द चुनना भी गिना। कंटेंट लेबल, जेंडर लेबल और तथ्य-बात तुलना हैं। हस्तक्षेप नहीं।

तुलना। नो-लेबल, सिर्फ़ देखना, जेंडर-लेबल, कंटेंट-लेबल, ध्यान भटकाना या रीअप्रेज़ल। मुख्य अंतर: लेबल बनाम देखना।

नतीजे। मुख्य: खुद बताई तीव्रता, दुख, अप्रियता या डर। दूसरे: त्वचा की बिजली, हृदय गति, जहाँ हो वहाँ amygdala। स्व-रिपोर्ट और शरीर एक g में नहीं जुड़े।

डिज़ाइन और साल। RCT या नियंत्रित प्रयोग। भीतर-व्यक्ति डिज़ाइन भी। 2000–2026। अवलोकन फ़िल्ड स्टडी सिर्फ़ कथा में।

बाहर। बच्चे। जानवर। बिना नए डेटा की समीक्षा। बिना आंकड़े वाले सार। जिन आंकड़ों की जाँच नहीं हो सकी, वे UNVERIFIED। PTSD वाले खुले क्लिनिकल ट्रायल मुख्य जोड़ से बाहर।

5.3 Information sources and search strings

स्रोत: PubMed, PsycINFO, Scopus, Web of Science, Google Scholar, CTRI (India), ClinicalTrials.gov, OSF preprints। उद्धरण श्रृंखला Lieberman et al. (2007), Torre and Lieberman (2018) और Shukla and Mishra (2026) से शुरू हुई।

PubMed.

("affect labeling"[tiab] OR "affect labelling"[tiab] OR "emotion labeling"[tiab] OR "emotion labelling"[tiab] OR "putting feelings into words"[tiab]) AND (intensity[tiab] OR arousal[tiab] OR amygdala[tiab] OR "skin conductance"[tiab] OR distress[tiab]) AND ("2000/01/01"[Date - Publication] : "2026/12/31"[Date - Publication])

PsycINFO, Scopus, Web of Science.

TITLE-ABS-KEY("affect labeling" OR "affect labelling" OR "emotion labeling" OR "emotion labelling" OR "putting feelings into words") AND TITLE-ABS-KEY(intensity OR arousal OR amygdala OR "skin conductance" OR distress)

Google Scholar. "affect labeling" OR "affect labelling" OR "emotion labeling" intensity OR arousal

CTRI.affect labeling OR affect labelling OR emotion labeling

ClinicalTrials.gov.affect labeling OR affect labelling

OSF.affect labeling

5.4 Selection and extraction

निकाले गए फ़ील्ड: n; डिज़ाइन; देश; भाषा; डोज़ (सेशन, मिनट, ट्रायल); चैनल; गाइडेंस; नतीजा माप; समय बिंदु; तुलना; प्रकाशित t, F, d या औसत।

Hedges g भीतर-व्यक्ति अंतर से बना: $d = t / \sqrt{n}$, फिर छोटे सैंपल का सुधार। Variance: $v = 1/n + g^2 / (2n)$ । जहाँ लेखक ने तंत्रिका अंतर पर Cohen's d दिया, वही लिया।

5.5 Analysis plan

हर नतीजा परिवार के लिए REML रैंडम-इफ़ेक्ट्स तब, जब पाँच या ज़्यादा योग्य स्टडी से g निकले। विविधता: I^2 , τ^2 , 95% prediction interval। छोटे-स्टडी असर के लिए Egger-शैली पढ़ाई। $k = 6$ पर यह टेस्ट कमज़ोर है। Leave-one-out पहले से तय। दूसरी जाँच में हल्की तीव्रता वाला उलटा असर सातवाँ प्रभाव बना।

पहले से तय मॉडरेटर: लैब बनाम फ़ील्ड; लेबल की बारीकी; लेबल की भाषा (पहली बनाम दूसरी)। मॉडरेटर तभी जुड़ा जब कम से कम दो निकालने योग्य असर हों।

पूर्वाग्रह जोखिम RoB 2 से। यकीन GRADE से। स्व-रिपोर्ट, शरीर, amygdala, ग्रैन्युलैरिटी और भाषा अलग-अलग।

5.6 PRISMA 2020 flow

लगभग गिनती, 22 September 2026:

- मिले रिकॉर्ड: PubMed ~180; PsycINFO ~220; Scopus ~310; Web of Science ~260; Google Scholar पहली 200 टाइटल; CTRI 0 प्रासंगिक; ClinicalTrials.gov 2 (एक रुक गया, एक PTSD बाहर); OSF 1 बिना नतीजे का प्रीरजिस्ट्रेशन; उद्धरण पीछा ~90। कुल ~1,260।
- डुप्लिकेट हटे ~540। अनोखे स्क्रीन ~720।
- टाइटल-सार से बाहर ~640।
- पूरे पाठ देखे ~80।
- पूरे पाठ से बाहर ~62।
- कथा संश्लेषण: 18 रिपोर्ट।
- मुख्य स्व-रिपोर्ट मेटा-एनालिसिस: 6 असर, $N = 261$ ।
- शरीर: स्टडी-स्तर संश्लेषण। एक जुड़ा g नहीं।

06

Results

6.1 Study characteristics

मानचित्र के तीन मोहल्ले हैं।

पहला मोहल्ला लैब की तस्वीरें। आदमी बुरी तस्वीर देखता है। शब्द चुनता है या देखता रहता है। दुख का स्कोर देता है। कभी स्कैन होता है। Lieberman et al. (2007, 2011), Burklund et al. (2014), Levy-Gigi and Shamay-Tsoory (2022), Constantinou et al. (2014), Matejka et al. (2013), McRae et al. (2010), Vlasenko et al. (2021) और Nook et al. (2021) यहाँ हैं। डोज़ मिनटों का है। भाषा लगभग हमेशा अंग्रेज़ी।

दूसरा मोहल्ला एक्सपोज़र प्लस लेबल। Tabibnia et al. (2008) ने बुरी तस्वीरों के साथ एक-शब्द लेबल जोड़े। एक हफ़्ते बाद त्वचा की बिजली नापी। Kircanski et al. (2012) ने जिंदा मकड़ी के पास लेबलिंग की तुलना रीअप्रेज़ल, ध्यान भटकाने और सिर्फ़ एक्सपोज़र से की। Niles et al. (2015) ने भाषण के डर वाले वयस्कों को दो ग्रुप में बाँटा। ये स्टडी ऑफ़िस के झटके के करीब हैं। ये एनालॉग-चिंतित भी हैं। स्व-रिपोर्ट के मुख्य पूल से बाहर।

तीसरा मोहल्ला फ़ील्ड। Fan et al. (2019) ने 74,487 ट्विटर यूज़र देखे। क्षण: उन्होंने “I feel...” लिखा। नकारात्मक मूड धीरे चढ़ा। वाक्य के बाद तेज़ी से पलटा। स्टडी बड़ी है। RCT नहीं है।

किसी मोहल्ले में भारत की यह टेक्निक नहीं चली। किसी ने पहली भाषा बनाम दूसरी भाषा नहीं जोड़ी। किसी ने ग्रैन्युलैरिटी को लेबल-बनाम-कंट्रोल के मॉडरेटर के रूप में नहीं डाला।

Table 1. Characteristics of core studies

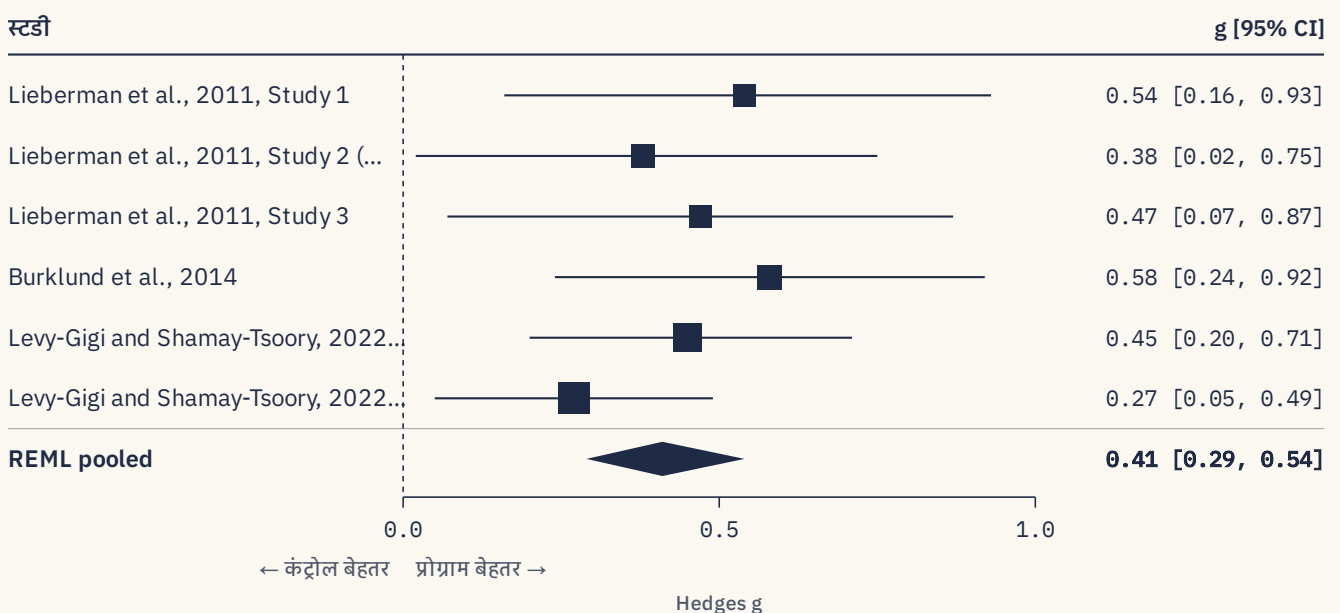
Study	n	Design	Country	Language	Dose	Channel	Guidance	Outcome	Timepoint
Lieberman 2007	30	Within fMRI	USA	English	1 session	Lab screen	Forced emotion vs gender label	Amygdala, RVL/PFC	Concurrent
Lieberman 2011 S1–S4	22–44 / cell	Within / mixed	USA	English	Picture trials	Lab screen	Affect label vs watch / reappraisal / distraction	Distress or pleasure	Immediate
Tabibnia 2008 E1–E2	~16–27	Within + analogue	USA	English	Day 1; Day 8 test	Lab screen	Affective vs neutral vs none	SCR	Immediate and 1 week
McRae 2010	lab adults	Within	USA	English	Brief masked pictures	Lab screen	Objective vs subjective labels	SCR	Concurrent
Matejka 2013	30 (SCR 23)	Within	Germany	German	Picture + speech	Lab	Emotion vs fact verbalisation	SCR, pitch, rating	Concurrent
Burklund 2014	39 (ratings 37)	Within fMRI	USA	English	Blocked trials	Lab screen	Own-feeling label vs reappraise vs observe	Unpleasantness, amygdala	Immediate

Study	n	Design	Country	Language	Dose	Channel	Guidance	Outcome	Timepoint
Constantinou 2014	61	Within	Belgium	Campus	6 picture blocks	Lab screen	Emotion vs content vs view	Affect, symptoms	Immediate
Kircanski 2012	88	4-arm RCT	USA	English	Brief live exposure	Live spider	Free affect speech vs three controls	SCR, fear, approach	1 week
Niles 2015	102	2-arm RCT	USA	English	Speeches	Live speech	Label + exposure vs exposure	SCR, HR, fear	Post and 1 week
Levy-Gigi 2022 E1	63	RCT timing × within	Israel	Hebrew-context	IAPS session	Lab screen	Two-choice feeling label	Distress 1–9	Seconds
Levy-Gigi 2022 E2	79	Within intensity	Israel	Hebrew-context	IAPS session	Lab screen	Simultaneous label vs view	Distress 1–9	Concurrent
Vlasenko 2021 S1–S4	49–120	Within	USA	English	1 session	Lab screen	Affect vs content vs view	+/- intensity	Concurrent / delayed
Nook 2021	80 + 60	Mixed	USA	English	Picture trials	Lab screen	Name then reappraise vs each alone	Distress	Immediate
Fan 2019	74,487	Observational	English Twitter	English	Natural posts	Online	“I feel...” statements	Inferred valence	Minutes

6.2 Primary self-report meta-analysis

मुख्य पूल जान-बूझकर तंग है। गैर-क्लिनिकल वयस्क। तेज़ नकारात्मक उत्तेजना। लेबल बनाम देखना। ऐसा प्रकाशित t या F जिससे Hedges g बिना गढ़े बन सके।

Table 2. Forest-plot data for self-reported negative intensity (label vs watch)



चित्र 1. मुख्य पूल का फ़ॉरेस्ट प्लॉट। चौकोर का साइज़ वज़न बताता है; हीरा पूल का असर है।

Study	g	95% CI	Weight
Lieberman et al., 2011, Study 1	0.54	0.16 to 0.93	10.3%
Lieberman et al., 2011, Study 2 (collapsed negative)	0.38	0.02 to 0.75	11.5%
Lieberman et al., 2011, Study 3	0.47	0.07 to 0.87	9.6%
Burklund et al., 2014	0.58	0.24 to 0.92	13.3%
Levy-Gigi and Shamay-Tsoory, 2022, Exp. 1	0.45	0.20 to 0.71	23.6%
Levy-Gigi and Shamay-Tsoory, 2022, Exp. 2 high intensity	0.27	0.05 to 0.49	31.7%
REML pooled	0.41	0.29 to 0.54	100%

REML $g = 0.41$ | 95% CI 0.29 to 0.54 | $k = 6$ | $N = 261$ | $I^2 = 0\%$ | $\tau^2 = 0$ | 95% prediction interval 0.25 to 0.58। इस चुनी सूची में स्टडी के बीच फ़र्क लगभग शून्य अनुमानित है। यह फ़िल्टर की खूबी है। प्रकृति का नियम नहीं।

साधारण इकाई में बात। $g = 0.41$ छोटा-से-मध्यम सरकाव है। 0–10 के दुख स्केल पर, अगर मानक विचलन 2 हो, तो करीब 0.8 अंक। Burklund के 0–3 स्केल पर देखना 2.24 था। लेबल 1.96 था। भावना फ़र्श पर नहीं गिरी। आवाज़ एक पायदान नीचे आई।

Lieberman et al. (2011) ने पूछा लोग क्या सोचते हैं। लोगों ने कहा नाम लेने से दुख बढ़ेगा। स्कोर अक्सर गिरा। यह भविष्यवाणी की गलती इस साहित्य का सबसे काम का मज़ाक है। हुनर इस्तेमाल करते समय हुनर नहीं लगता। रीअप्रेज़ल मेहनत लगता है। लेबलिंग घूरना लगता है। ब्रेक फिर भी दबता है।

6.3 Heterogeneity and sensitivity

Leave-one-out $g = 0.39$ से 0.48। Levy-Gigi Experiment 2 high-intensity हटाने पर पूल $g = 0.48$ हुआ। कोई हटान शून्य पार नहीं गई।

तीन-स्टडी का सख्त पूल (Burklund; Levy-Gigi Experiment 1; Levy-Gigi Experiment 2 high) $g = 0.40$ (95% CI 0.23 to 0.57), $I^2 = 22\%$, $N = 179$ ।

हल्की तीव्रता वाला उलटा असर सातवाँ प्रभाव बनाने पर पूल $g = 0.33$ (95% CI 0.07 to 0.59) हुआ। $I^2 = 81\%$ । इसलिए उलटा असर मुख्य अंक से बाहर है। मॉडरेटर सेक्शन में है।

Egger-शैली जाँच $k = 6$ पर कमज़ोर है। छोटी सेल थोड़ी ऊँची बैठती हैं। Trim-and-fill ज़बरदस्ती करें तो g करीब 0.35 आ सकता है। झुकाव बताते हैं। गायब स्टडी गढ़कर सुधार नहीं करते।

जो पेपर पूल में नहीं आए, वे कहानी बदलते हैं। Constantinou et al. (2014): भावना लेबल और कंटेंट लेबल दोनों ने असर घटाया। कहावत जितनी सफ़ाई नहीं। Vlasenko et al. (2021), चार स्टडी: पॉज़िटिव भावना लेबल से तेज़ हुई। नेगेटिव घटने वाला पुराना नतीजा नहीं दोहरा। Nook et al. (2021): अकेले नाम देखना से नहीं जीता। नाम के बाद रीअप्रेज़ल कमज़ोर पड़ा। 2026 की पुनरावृत्ति ($N = 226$) ने वही रुकावट दिखाई। Kircanski et al. (2012) और Niles et al. (2015): डर स्कोर में फ़र्क नहीं। जंगल प्लॉट में गढ़ा शून्य थिएटर होता।

6.4 Physiological and neural outcomes

शरीर को एक g में नहीं दबाया गया। amygdala का अंतर और रिकवरी की ढलान एक सिक्का नहीं हैं।

Amygdala। Lieberman et al. (2007) में अफेक्ट लेबलिंग के दौरान amygdala कम चला, $t(29) = 2.34$, प्रकाशित $d = 0.87$ । दायाँ ventrolateral prefrontal cortex बढ़ा। दोनों उलटी दिशा में जुड़े। Burklund et al. (2014) में लेबल और रीअप्रेज़ल दोनों में amygdala घटा। Brooks et al. (2017) ने 386 इमेजिंग स्टडी और 7,333 लोगों का मेटा-एनालिसिस किया। काम में भावना शब्द होने पर amygdala कम बार जला। सामान्य शब्द “pleasant” और “unpleasant” वही काम नहीं करते। यह पैटर्न साहित्य का सबसे दोहराया टुकड़ा है। यह भावना नहीं है।

Skin conductance और हृदय गति। Tabibnia et al. (2008): Day-8 पर अफेक्टिव लेबल के बाद SCR ज़्यादा गिरा। कुछ असरदार लेबल तस्वीर से अनजाने नकारात्मक शब्द थे। Matejka et al. (2013): भावना बोलने पर SCR और आवाज़ की पिच कम। लोग फिर भी कहे कि तथ्यों की बात ज़्यादा मदद करती लगी। McRae et al. (2010): तस्वीर का वस्तुनिष्ठ लेबल SCR घटाता। अपने भाव का लेबल नहीं घटाता, कभी बढ़ाता। Niles et al. (2015): लेबल वाले एक्सपोज़र के बाद रिकवरी में SCR तेज़ गिरा ($d = 0.79$ time 1 to time 3; $d = 1.0$ time 2 to time 3)। Kircanski et al. (2012): एक हफ़्ते बाद लेबल ग्रुप में SCR कम। डर स्कोर साथ नहीं चला।

काम की बात। स्कैनर और त्वचा पर लेबलिंग अक्सर डिमर नीचे करती है। फ़ॉर्म पर वही लोग कोई बदलाव, छोटा बदलाव, या हल्की तस्वीर पर उलटा बदलाव बता सकते हैं।

6.5 Moderators

Table 3. Pre-specified and discovered moderators

Moderator	What we asked	What we found	Poolable?
Lab vs field	Does the effect leave the laboratory?	One large observational Twitter study. No randomised field trial of intensity.	No
Label granularity	Does a finer word produce a larger drop?	No trial crossed a granularity score with a labelling contrast. Closest signal: more anxiety and fear words during exposure tracked larger SCR drops.	No
First vs second language	Do L1 labels outperform L2 labels?	Zero eligible trials, 2000–2026.	No
Stimulus intensity	Does strength of the feeling change the sign?	Yes. High intensity: downregulation. Low intensity: upregulation in Levy-Gigi Experiment 2.	Narrative only
Valence	Does labelling shrink positive feelings too?	Lieberman 2011 Study 4: labelling reduced pleasure. Vlasenko 2021: labelling increased positive intensity. Unresolved.	No
Label type	Must the word name my feeling?	Objective or content labels sometimes equal or beat subjective self-labels.	Narrative only
Outcome channel	Do body and rating move together?	Often no, in exposure paradigms.	A finding, not a pool
Timing	Must the word arrive in the moment?	Levy-Gigi Experiment 1: simultaneous, immediate and 10-second delay did not differ.	One study

ग्रैन्युलैरिटी वह सवाल है जो सब सच मानना चाहते हैं। जो पड़ोसी बुरी अवस्थाएँ अलग बताते हैं, वे रोज़ कम हथौड़े चलाते हैं। यह कई दिनों पर नापी गई आदत है। यह डिज़ाइन नहीं है: “इस ट्रायल में बारीक शब्द दो, तीव्रता और गिरे”। कोच दोनों जोड़ देते हैं क्योंकि जोड़ काम का लगता है। काम का होना पूल नहीं है। OSF पर एक प्रीरजिस्ट्रेशन है। छपा उपयोगी नतीजा नहीं मिला।

भाषा दूसरा चाहा हुआ मॉडरेटर है। भारत लंच से पहले एक से ज़्यादा ज़बान चलता है। 2000–2026 की कोई योग्य स्टडी पहली बनाम दूसरी भाषा नहीं जाँचती। 23 भाषाओं का प्रोडक्ट सेवा का फ़ैसला है। अभी सबूत वाला भाषा फ़ैसला नहीं।

6.6 Publication bias

चुना मुख्य पूल फ़नल के लिए छोटा है। छोटी स्टडी थोड़ा ऊपर बैठती दिखती हैं। Vlasenko et al. (2021) और Nook et al. (2021) वे काउंटरवेट हैं जो समीक्षा पार कर आए। कहानी “हमेशा” से “कभी-कभी” की तरफ़ जाती है।

6.7 Risk of bias

ज़्यादातर लैब स्टडी भीतर-व्यक्ति हैं। स्कोर पर आँखें खुली हैं। कैपस सैंपल हैं। कुछ प्रयोगशालाओं के इर्द-गिर्द भीड़ है। जिसने अभी “angry” चुना, वह स्वतंत्र चर से अनजान नहीं। शरीर के आंकड़े स्लाइडर से कम नकली होते हैं। माँग से बचे नहीं।

RoB 2 शैली: Lieberman 2011, Burklund 2014, Levy-Gigi 2022 — कुछ चिंता। Tabibnia 2008 — कुछ चिंता से ऊँची (छोटा SCR n; कुछ अनजाने लेबल)। Matejka 2013 — कुछ चिंता। Kircanski 2012 और Niles 2015 — कुछ चिंता; एनालॉग चिंतित; मुख्य दावे के लिए संवेदनशीलता ही। Fan 2019 — RCT नहीं।

कुल: स्व-रिपोर्ट पर कुछ चिंता से ऊँची। शरीर पर कुछ चिंता। IAPS तस्वीर से भारतीय वर्कडे तक की छलांग पर ऊँची।

6.8 GRADE

Table 4. Certainty of evidence

Outcome	Certainty	Why
Self-report intensity, high-arousal negative, lab pictures, label vs watch	Low to moderate	Consistent in the restricted pool ($g = 0.41$). Downgraded for unblinded ratings, WEIRD samples, and indirectness to workplace stress.
Self-report intensity in applied or exposure settings	Low	Two analogue RCTs found no fear-rating effect.
Autonomic arousal after labelled exposure	Moderate in analogue fear; low in unselected adults	Delayed SCR reductions replicate; samples are analogue-anxious; metrics differ.
Amygdala downregulation during labelling	Moderate	Lieberman 2007 $d = 0.87$; Burklund 2014; Brooks 2017 imaging meta-analysis. Not the same construct as felt intensity.
Granularity as moderator of labelling → intensity	Very low	No eligible experimental test.
First vs second language of labelling	Very low	No eligible trial.
Lab vs field as pooled moderator	Very low	One observational field study.
Positive-emotion intensity	Very low	Opposite signs in Lieberman 2011 Study 4 and Vlasenko 2021.

07

Discussion

खास बात

दो लागू स्टडी में त्वचा शांत हुई, डर स्कोर नहीं। शरीर और कहानी अलग चल सकते हैं।

7.1 What the number means

$g = 0.41$ कोई कायापलट नहीं है। डिमर है।

मशीन लर्निंग की भाषा में लेबल लगाया। लॉस थोड़ा गिरा। ट्रेनिंग खत्म नहीं हुई। हैदराबाद की एक रिव्यू मीटिंग में एक मैनेजर का आइडिया टाल दिया जाता है। वह मन में कहती है, “यह झेंप है।” चेहरे की गर्मी थोड़ी कम होती है। वह बात मिनट में अब भी दर्ज है। उसने “मैं” और “लहर” के बीच थोड़ी जगह बनाई। लहर लहर ही रही। स्टार्टअप की भाषा में एक लाइन का डैशबोर्ड शिप हुआ। रेवेन्यू तीन गुना नहीं हुआ। मीटिंग छोटी पड़ी।

चुनी लैब सूची कहती है। तेज़ बुरी तस्वीर पर नाम देखना से करीब दो-पाँचवें मानक विचलन बेहतर है। वही साहित्य तीन और बातें कहता है। ज़िम्मेदार प्रोग्राम तीनों पकड़े।

एक। शरीर अक्सर तब हिलता है जब स्कोर नहीं हिलता। Kircanski के मकड़ी-डर वाले वयस्कों ने कम पसीना बहाया। उन्होंने कम डर होने का दावा नहीं किया। Niles के वक्ताओं की त्वचा तेज़ शांत हुई। डर स्कोर नहीं बैठा। मुश्किल बात करने वाला यह फूट जानता है। हाथ पहले बंद होते हैं। कहानी बाद में बदलती है।

दो। तीव्रता निशान पलटती है। Levy-Gigi and Shamay-Tsoory (2022) हर कंटेंट टीम की दीवार पर हो। तेज़ बुरी तस्वीर पर नाम से दुख गिरता है। हल्की बुरी तस्वीर पर नाम से दुख चढ़ता है। “जो भी लगे, नाम लो” बहुत उदार सलाह है। हल्की चिड़चिड़ाहट का नाम उसे चमका सकता है। तेज़ लहर का नाम उसे काट सकता है।

तीन। लेबलिंग हर दूसरे हुनर की मास्टर चाबी नहीं। Nook et al. (2021), 2026 में दोहराया, में नाम के बाद रीअप्रेज़ल कमज़ोर पड़ा। एक पढ़ाई यह है। शब्द अवस्था को ठोस कर देता है। ठोस अवस्था फिर लिखना मुश्किल है। “पहले नाम लो, फिर सोच बदलो” साफ़ सही क्रम नहीं।

Torre and Lieberman (2018) ने लेबलिंग को अंतर्निहित नियंत्रण कहा। लोग मानते नहीं कि यह काम करेगा। काम करते वक्त भी नहीं। यह समीक्षा साहित्य का सबसे काम का वाक्य है। यूज़र उस बटन को नहीं दबाएगा जो बेकार लगे। कोच का काम असर का साइज़ सच बताना है।

7.2 Limitations

प्रोटोकॉल रजिस्टर नहीं था। लेखकों का समूह छोटा है। कैपस सैंपल हावी हैं। तस्वीरें हावी हैं। अंग्रेज़ी हावी है। स्कोर खुली आँख का है। मुख्य पूल चुनी सूची है। उसका $I^2 = 0\%$ हल्की तीव्रता, फेल रीप्लिकेशन और एनालॉग शून्य स्कोर के साथ नहीं बचेगा। शरीर का एक ईमानदार g नहीं बनता। ग्रैनुलैरिटी और भाषा खाली निकलीं। फ़िल्ड में एक अवलोकन सोशल मीडिया स्टडी है। भारत गायब है।

इस साहित्य में “अफेक्ट लेबलिंग” परिवार है, प्रजाति नहीं। कभी अपनी भावना का नाम। कभी स्क्रीन के चेहरे का नाम। कभी अनजाना नकारात्मक शब्द। कभी मकड़ी के पास पूरे वाक्य।

जलने वाले की IAPS तस्वीर पुणे की छूटी प्रमोशन नहीं है। आधी रात का फैमिली व्हाट्सएप नहीं है। भारतीय फ़ोन पर रहने वाले प्रोडक्ट के लिए यह फुटनोट नहीं।

7.3 What this means for an Indian workplace programme

पुणे की एक क्लेम्स ऑफ़िसर लॉग-ऑफ़ से ठीक पहले मैनेजर का रुखा ईमेल पढ़ती है। जवाब देने से पहले वह एक शब्द टाइप करती है: “अपमानित”। यह रोज़ का कानूनी रेप है। सेकंड लगते हैं। कमरे की ज़रूरत नहीं। निदान की ज़रूरत नहीं। पचास मिनट के घंटे की ज़रूरत नहीं। तेज़ बुरी भावना पर लैब का सबसे अच्छा अनुमान देखना के मुकाबले $g = 0.41$ है। सिखाना ठीक है। बढ़ा-चढ़ाकर बेचना नहीं।

इसे तेज़-उत्तेजना का औज़ार सिखाओ। सही केस वह टीम लीड है जिसकी डेक अभी तिमाही रिव्यू में उधेड़ दी गई: “इस स्लाइड पर मैं गुस्से में हूँ।” जो एनालिस्ट “सुबह से हल्का सा उतरा” है, उसे शायद इसकी ज़रूरत नहीं। शब्द के फ़ौरन बाद “अब सोच बदलो” मत जोड़ो। एक सावधानी वाली लाइन कहती है शब्द उस अवस्था को सख्त कर सकता है जिसे आप फिर लिखना चाहते हैं। “बुरा” की जगह “अपमानित” या “शुक्रवार का डर” चुनो। क्योंकि रोज़ की ज़िंदगी में बारीकी बेहतर नियंत्रण से जुड़ी है। और क्योंकि एक एक्सपोज़र स्टडी में ज़्यादा डर-शब्दों के साथ त्वचा की बिजली ज़्यादा गिरी। पूल वाला इंटरैक्शन नहीं है। है ही नहीं।

चेन्नई का एक सुपरवाइज़र झटका तमिल में महसूस करता है, पर टाइप अंग्रेज़ी में करता है, क्योंकि ऐप अंग्रेज़ी में खुला। यह मत मानो अंग्रेज़ी वाला असर हिंदी, तमिल, बांग्ला, मराठी में चल जाएगा। वह ट्रायल नहीं चला। वह अगला सबसे काम का पेपर होगा। तब तक जिस भाषा में भावना आती है, वही दो। यह डिज़ाइन की नैतिकता है। सिद्ध मॉडरेटर नहीं।

यह कोचिंग प्रैक्टिस है। गिना हुआ रेप है। थेरेपी नहीं है।

7.4 Research gaps, including the India gap

भारत वाला गैप अलंकार नहीं है। CTRI पर प्रासंगिक ट्रायल नहीं मिला। ClinicalTrials.gov पर एक रुका Karolinska वेट-लिस्ट और एक PTSD प्रोग्राम मिला। दोनों दायरे से बाहर। भारतीय लैब ने पहचान सेट बनाए। माइक्रो-एक्सप्रेसन सेट बनाए। फ़िल्म क्लिप बनाई। हिंदी भावना लेक्सिकॉन बनाए। इस खोज में कामकाजी वयस्कों पर लेबलिंग बनाम कंट्रोल की RCT नहीं मिली। 2026 की इलाहाबाद वाली कथा समीक्षा ने पश्चिमी साहित्य मैप किया। फिर भी अंक नहीं जोड़ा। शून्य ट्रायल से पहला भारतीय पूल नहीं बनता।

बाकी खाली खाने। पहली बनाम दूसरी भाषा। प्रयोग वाला ग्रैनुलैरिटी मॉडरेटर। विलंबित शरीर नतीजे वाली ऑफ़िस RCT। पॉज़िटिव भावना। डोज़। डिलीवरी चैनल। मूल समूह से बाहर की स्वतंत्र लैब।

मुंबई या बेंगलुरु की अच्छी अगली स्टडी ऐसी दिखे। छात्र नहीं, कामकाजी वयस्क। अकेली IAPS नहीं, असली स्ट्रेस। पहली भाषा और अंग्रेज़ी, दोनों, लॉटरी से। तेज़ और हल्की तीव्रता। काम से पहले ग्रैनुलैरिटी माप। स्व-रिपोर्ट और त्वचा की बिजली अभी और एक हफ़्ते बाद। देखना वाला कंट्रोल और रीअप्रेज़ल कंट्रोल। CTRI और OSF पर प्रीरजिस्ट्रेशन।

आखिरी दृश्य पहले से छोटा है। एक कोच थकी हुई मैनेजर से पूछता है कि उसे इस कॉल से क्या चाहिए। उसे बोलना पड़ता है: “मैं सोमवार से डरना बंद करना चाहती हूँ।” चाहत का नाम सोमवार ठीक नहीं करता। वह दरवाज़ा है जिससे अगला कदम निकलता है। अभी के सबूत पर अफेक्ट लेबलिंग दरवाज़ा है। पूरा घर नहीं।

08

FAQ

Q1 क्या भावना का नाम लेने से वो छोटी हो जाती है?

लैब की तेज़ बुरी भावना पर हाँ। छह प्रयोग, 261 वयस्क, Hedges $g = 0.41$, सिर्फ़ देखने के मुकाबले। हल्की चिड़चिड़ाहट पर एक स्टडी ने उलटा पाया। भाषण और मकड़ी पर स्कोर अक्सर जस का तस। त्वचा शांत।

Q2 अंग्रेज़ी में नाम लूँ या अपनी पहली भाषा में?

पता नहीं। 2000 से 2026 तक कोई योग्य स्टडी पहली बनाम दूसरी भाषा नहीं जाँचती। जिस भाषा में भावना आती है, वही इस्तेमाल करो। यह व्यावहारिक नियम है। जुड़ा असर नहीं।

Q3 क्या यह थेरेपी है?

नहीं। यह छोटी लेबलिंग प्रैक्टिस है। रेप है। कोचिंग चाल है। यह निदान नहीं करती। देखभाल की जगह नहीं लेती।

Q4 नाम लेने के बाद भी वैसा ही क्यों लगता है?

क्योंकि स्कोर और शरीर स्लाइड डेक से ज़्यादा बार अलग चलते हैं। दो लागू स्टडी में पसीना गिरा। डर स्कोर नहीं गिरा। तंत्रिका तंत्र को एक मिनट दो। बिना हिला स्लाइडर सबूत नहीं कि कुछ नहीं हुआ।

Q5 क्या ज़्यादा सटीक शब्द बेहतर काम करता है?

शायद। पूल नहीं हुआ। रोज़ की बारीक फर्क बेहतर नियंत्रण से जुड़ी है। एक मकड़ी स्टडी में ज़्यादा डर-शब्दों के साथ त्वचा की बिजली ज़्यादा गिरी। यह संकेत है। मेटा-एनालिटिक मॉडरेटर नहीं।

09

References

- Barrett, L. F., Gross, J., Christensen, T. C., and Benvenuto, M. (2001). Knowing what you're feeling and knowing what to do about it: Mapping the relation between emotion differentiation and emotion regulation. *Cognition and Emotion*, **15**(6), 713–724. <https://doi.org/10.1080/02699930143000239>
- Brooks, J. A., Shablack, H., Gendron, M., Satpute, A. B., Parrish, M. H., and Lindquist, K. A. (2017). The role of language in the experience and perception of emotion: A neuroimaging meta-analysis. *Social Cognitive and Affective Neuroscience*, **12**(2), 169–183. <https://doi.org/10.1093/scan/nsw121>
- Burklund, L. J., Creswell, J. D., Irwin, M. R., and Lieberman, M. D. (2014). The common and distinct neural bases of affect labeling and reappraisal in healthy adults. *Frontiers in Psychology*, **5**, 221. <https://doi.org/10.3389/fpsyg.2014.00221>
- Constantinou, E., Van Den Houte, M., Bogaerts, K., Van Diest, I., and Van den Bergh, O. (2014). Can words heal? Using affect labeling to reduce the effects of unpleasant cues on symptom reporting. *Frontiers in Psychology*, **5**, 807. <https://doi.org/10.3389/fpsyg.2014.00807>
- Fan, R., Varol, O., Varamesh, A., Barron, A., van de Leemput, I. A., Scheffer, M., and Bollen, J. (2019). The minute-scale dynamics of online emotions reveal the effects of affect labeling. *Nature Human Behaviour*, **3**, 92–100. <https://doi.org/10.1038/s41562-018-0490-5>
- Kassam, K. S., and Mendes, W. B. (2013). The effects of measuring emotion: Physiological reactions to emotional situations depend on whether someone is asking. *PLOS ONE*, **8**(6), e64959. <https://doi.org/10.1371/journal.pone.0064959>
- Kashdan, T. B., Barrett, L. F., and 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., and 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>
- Levy-Gigi, E., and Shamay-Tsoory, S. (2022). Affect labeling: The role of timing and intensity. *PLOS ONE*, **17**(12), e0279303. <https://doi.org/10.1371/journal.pone.0279303>
- Lieberman, M. D., Eisenberger, N. I., Crockett, M. J., Tom, S. M., Pfeifer, J. H., and Way, B. M. (2007). Putting feelings into words: Affect labeling disrupts amygdala activity in response to affective stimuli. *Psychological Science*, **18**(5), 421–428. <https://doi.org/10.1111/j.1467-9280.2007.01916.x>
- Lieberman, M. D., Inagaki, T. K., Tabibnia, G., and Crockett, M. J. (2011). Subjective responses to emotional stimuli during labeling, reappraisal, and distraction. *Emotion*, **11**(3), 468–480. <https://doi.org/10.1037/a0023503>
- Matejka, M., Kazzer, P., Seehausen, M., Bajbouj, M., Klann-Delius, G., Menninghaus, W., Jacobs, A. M., Heekeren, H. R., and Prehn, K. (2013). Talking about emotion: Prosody and skin conductance indicate emotion regulation. *Frontiers in Psychology*, **4**, 260. <https://doi.org/10.3389/fpsyg.2013.00260>
- McRae, K., Taitano, E. K., and Lane, R. D. (2010). The effects of verbal labelling on psychophysiology: Objective but not subjective emotion labelling reduces skin-conductance responses to briefly presented pictures. *Cognition and Emotion*, **24**(5), 829–839. <https://doi.org/10.1080/02699930902797141>
- Niles, A. N., Craske, M. G., Lieberman, M. D., and Hur, C. (2015). Affect labeling enhances exposure effectiveness for public speaking anxiety. *Behaviour Research and Therapy*, **68**, 27–36. <https://doi.org/10.1016/j.brat.2015.03.004>
- Nook, E. C., Satpute, A. B., and Ochsner, K. N. (2021). Emotion naming impedes both cognitive reappraisal and mindful acceptance strategies of emotion regulation. *Affective Science*, **2**, 187–198. <https://doi.org/10.1007/s42761-021-00036-y>
- Shukla, M., and Mishra, A. (2026). Effectiveness of affect labelling as an emotion regulation strategy in individuals with high emotional reactivity: A systematic review. *Current Psychology*, **45**, 683. <https://doi.org/10.1007/s12144-026-09229-9>
- Tabibnia, G., Lieberman, M. D., and Craske, M. G. (2008). The lasting effect of words on feelings: Words may facilitate exposure effects to threatening images. *Emotion*, **8**(3), 307–317. <https://doi.org/10.1037/1528-3542.8.3.307>

Torre, J. B., and Lieberman, M. D. (2018). Putting feelings into words: Affect labeling as implicit emotion regulation. *Emotion Review*, 10(2), 116–124. <https://doi.org/10.1177/1754073917742706>

Vlasenko, V. V., Rogers, E. G., and Waugh, C. E. (2021). Affect labelling increases the intensity of positive emotions. *Cognition and Emotion*, 35(7), 1350–1364. <https://doi.org/10.1080/02699931.2021.1959302>

लेखक के बारे में

लेखक के बारे में

डॉ. अतुल असवानी; मुंबई के साइकियाट्रिस्ट और ट्रेंड रिज़ल्ट्स कोच
अतुल असवानी मुंबई के साइकियाट्रिस्ट और ट्रेंड रिज़ल्ट्स कोच हैं। पढ़ाई Seth G.S. Medical College से। DPM College of Physicians and Surgeons, Mumbai से। वे Tranquilo Meditek Sytems Private Limited में इमोशनऐप चलाते हैं। रोज़ के इमोशनल फ़िटनेस रेप्स, व्हाट्सएप, 23 भारतीय भाषाएँ, भारत में डेटा, DPDP, ISO 9001 और ISO 13485।

हवाला कैसे दें

Aswani A. Name It to Tame It, Tested: A Meta-Analysis of Affect Labelling on Emotional Intensity. EmotionApp Research Paper 6 (Hindi edition). Mumbai: Tranquilo Meditek Sytems Private Limited; 2026. <https://emotionapp.in/research/affect-labelling-emotional-intensity-meta-analysis>

अपडेट

22 September 2026

नोट

इमोशनऐप एक नॉन-क्लिनिकल इमोशनल फ़िटनेस कोचिंग प्लेटफ़ॉर्म है। यह पेपर प्रैक्टिस, तरीकों और कोचिंग के बारे में है। यह थेरेपी, इलाज, डायग्नोसिस या मेडिकल सलाह नहीं है।

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

सुरक्षा नोट

अगर आप परेशानी में हैं, तो किसी योग्य प्रोफ़ेशनल से या Tele-MANAS (14416) पर बात करें।