

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Mapping Emotion Representations: Evidence for Category–Dimension Relationships and Individual Differences

Permalink

<https://escholarship.org/uc/item/1571b50s>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Wu, Claudia

Yal—âún, ,Äŕzge Nilay

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Mapping Emotion Representations: Evidence for Category–Dimension Relationships and Individual Differences

Claudia Wu (claudia_wu@sfu.ca)¹ & Özge Nilay Yalçın (oyalcin@sfu.ca)²

^{1,2}School of Interactive Arts and Technology, Simon Fraser University, Surrey, BC, Canada

Abstract

A central question in affective computing is how emotions should be represented. Emotion datasets typically adopt either categorical labels (e.g., Ekman categories) or dimensional ratings (e.g., Valence–Arousal–Dominance). Although these frameworks are often treated as comparable descriptions of affect, most work relies on only one representation at a time, limiting cross-dataset comparability and leaving their relationship empirically under-specified. This study explores how VAD ratings map onto categorical emotion judgments, and whether this mapping is consistent across individuals. We fit generalized linear mixed models with participant-level random effects using the MSP-Podcast Corpus to predict categorical labels from dimensional ratings. The results showed a significant effect of Valence, Arousal, and Dominance across several emotion categories (Sadness, Happiness, Anger, Fear, and Disgust) with significant participant-level variability. Overall, our results highlight that categorical and dimensional emotion representations are systematically related but their mapping is shaped by individual differences in affective interpretation. This suggests that affective computing models should explicitly account for rater-dependent variation when using dimensional and categorical emotion representations.

Keywords: affective computing; emotion; dimensional emotion; categorical emotion; emotion representations

Introduction

Affective computing is an interdisciplinary research area that aims to develop intelligent computational systems capable of understanding, interpreting, and responding to human emotions (Picard, 2015). As human-computer interaction applications become more prevalent, integrating emotional capabilities into these systems can facilitate transformative advancements in healthcare, education, and entertainment (Sethu et al., 2019). However, current emotion recognition systems are hindered by their poor performance in emotion recognition and fail to capture the complexity of human emotions (Sethu et al., 2019). The field faces persistent challenges, including accounting for individual and cultural differences and modeling the temporal dynamics of emotion (Barrett et al., 2019). A core theoretical and methodological issue is how emotions should be presented (Calvo & D’Mello, 2010).

Emotion representation in affective computing is grounded in two dominant psychological frameworks: categorical and dimensional. Categorical (or discrete) approaches describe emotions as discrete states, such as happiness, sadness, anger, fear, and disgust, often motivated by basic emotion theory

(Ekman, 1999). Dimensional approaches represent affect in a continuous space—most commonly along Valence (or Pleasure), Arousal, and Dominance axes—capturing graded and mixed emotional experiences (Mehrabian, 1996; Russell, 1980). In principle, these frameworks should be systematically related, since both aim to describe underlying affective phenomena, however, no universally accepted representation exists across disciplines (Sethu et al., 2019; Wang et al., 2022). In practice, emotion datasets and benchmarks typically adopt either categorical or dimensional representations (Ma & Yarosh, 2021), making it difficult to compare results across studies, transfer models between corpora, or interpret model behavior under a shared theoretical lens.

Despite the categorical and dimensional representations being used interchangeably in affective computing, few studies empirically quantify the correspondence between these frameworks under naturalistic conditions, or examine whether such correspondences are stable across raters. This gap is consequential: if the mapping between VAD and emotion categories varies systematically across participants, then treating a single “ground-truth” label as universal may obscure important observer-dependent differences in emotion perception and annotation (Barrett, 2006a, 2006b; Cowen & Keltner, 2017). To address this gap, we focus on the following research questions:

RQ1: Is there a correlational relationship between categorical and dimensional models of emotion?

RQ2: If so, is the relationship consistent between different participants?

To answer these questions, we analyze the MSP-Podcast Corpus (Lotfian & Busso, 2017), a naturalistic speech dataset with dimensional and categorical annotations, using generalized linear mixed-effects models (Bates et al., 2015) to assess the significance of VAD dimensions for multiple basic emotion categories while accounting for participant-level variability. Our results showed that the dimensional ratings of Valence, Arousal, and Dominance were significantly correlated with discrete basic emotion categories. However, ratings varied widely among participants. By empirically characterizing the relationship between categorical and dimensional emotion representations, these results enable principled translation between different emotion models and datasets, supporting greater interoperability, comparability, and generalizability in affective computing research.

Together, these findings make three contributions: (1) an empirical characterization of VAD–category correspondences

in naturalistic speech, (2) evidence that the mapping varies systematically across participants, and (3) implications for interoperability across emotion datasets, evaluation protocols, and modeling approaches in affective computing.

Related Work

Categorical Representations

The categorical (or discrete) emotion model proposes that there exists a defined subset of emotion labels that are innate, universal, and discrete (Ekman, 1999). Ekman's basic emotion model and its variations are widely accepted within the field, particularly because of its foundation in Darwinian theory (Sethu et al., 2019; Wang et al., 2022). The model posits that basic emotions arise from human instinct, which leads individuals to exhibit similar patterns of expression in similar situations (Wang et al., 2022). Research that adopts this model typically uses variations of Ekman's emotion model, which commonly includes anger, joy, sadness, fear, disgust, surprise, and neutral (Amir et al., 2015). The universality of these basic emotions is supported by a large body of research (Carlson & Hatfield, 1992; Ekman & Cordaro, 2011; Oatley & Jenkins, 1992), although recent work challenges this notion with evidence that emotion recognition differs between cultures, languages, and genders (Gong et al., 2018; Jack et al., 2012; Russell, 1994). However, there is no agreed-upon set of basic emotion categories in affective computing research.

Categorical labels are intuitive, easy to understand, and provide discrete, straightforward classifications that simplify data collection. However, the chosen list of emotion categories can bias participants' responses, as a predefined list may compel participants to select a specific category that may not capture the emotion they actually perceive. Including an "Other" category can alleviate this issue, but it may create complications during data analysis, particularly for small sample sizes (Sethu et al., 2019). Moreover, humans display a wide range of subtle and complex affective states conveyed through numerous possible facial and bodily expressions (Gunes et al., 2011). As such, relying solely on a single label may be inadequate for capturing the complexity of the emotional states exhibited by individuals in natural settings (Gunes et al., 2011).

Dimensional Representations

The dimensional emotion model proposes that affective states are not independent of one another, but are instead systematically related (Gunes et al., 2011). In this framework, emotions are continuous and multi-dimensional, characterized by graded representations. A well-known model is the Valence-Arousal circumplex model, which includes two dimensions: Valence, representing feelings ranging from negative to positive, and Arousal, measuring physiological activity and psychological alertness (Russell, 1980). Another widely accepted model, the Pleasure-Arousal-Dominance (PAD) model (Mehrabian, 1996), adds a third dimension of Dominance, which indicates the extent to which an individual feels restricted or in control of a situation (Bakker et al., 2014; Mehra-

bian & Russell, 1974; Wang et al., 2022). Collectively, these dimensions represent the emotional response that constitutes an individual's state of feeling. As a note, the terminology for these dimensions can vary across different theoretical accounts; for instance, Pleasure and Valence are used interchangeably, as well as Activation and Arousal (Sethu et al., 2019). This study will use the terms Valence and Arousal.

Dimensional ratings provide a graded and nuanced representation that captures emotional complexities such as intensity and variability. However, dimensional ratings are less intuitive, as most people do not naturally think of emotions in these terms. Studies that employ dimensional ratings often need to find methods for defining these terms without influencing participant ratings. Moreover, even extended dimensional models may not fully capture the richness of emotional behaviour and the individual differences, which can complicate computational interpretation (Sethu et al., 2019).

Additional Representations

While the categorical and dimensional emotion models are the most popular models used for emotion classification, other theoretical models of emotion have been proposed (Scherer, 2010; Wang et al., 2022). Appraisal-based theories are another popular model, which denotes that emotions are elicited by cognitive responses to antecedent situations and events (Scherer, 2010). Appraisal is typically perceived as the cause of the behavioral and cognitive changes associated with emotion, such that a specific set of cognitive appraisal allows for the prediction of which emotion type will be experienced (Marsella et al., 2010; Scherer, 2010). Other theories are less popular for computational modelling and often overlap with the aforementioned approaches. For instance, Barrett (2006b) proposes that emotions stem from a core affective state that is informed by conceptual knowledge of the context and semantic meanings. Meanwhile, Gray (1990) postulates that the brain systems mediating emotion and cognition largely overlap, derived from fundamental systems in the brain.

Prior Work on Representation Mapping

Very few studies have examined the relationship between categorical and dimensional models for emotion ratings.

Studies on human perception of emotion, such as Amir et al. (2015) and Hoffmann et al. (2012), observed inconsistent relationships between categorical and dimensional emotion labels. Amir et al. (2015) found that while the Neutral label had a mean Valence close to zero, its mean Arousal was approximately -0.5, based on ratings from 20 female participants. Hoffmann et al. (2012) asked 70 participants to position 22 discrete emotion labels onto a dimensional coordinate space based on their subjective knowledge. The results showed a high level of consistency among participants in positioning categorical emotion labels along the Valence axis, but greater variation along the Arousal and Dominance axes. However, both studies are limited by a small sample size or a reliance on participants' expectations of emotion mapping, rather than

the perception of real-world stimuli. Similarly, Cowen and Keltner (2017) found that categorical labels better captured subjective experience compared to dimensional ratings based on results from self-reported emotion responses elicited by provoking short videos. However, their focus on emotional experience rather than perception limits their applicability to emotion recognition systems.

While the universality of expressions has been a long-standing debate in emotion research, several studies have found differences in emotion perception across different cultures. Ogawa et al. (2005) found that images of facial expressions originally defined as "sadness" or "fatigue" by Russell and Fernández-Dols (1997) were identified as "disgust" among Japanese participants. Gong et al. (2018) had German and Chinese participants rate negative or positive images for emotional Arousal and found differences in ratings across age, gender, and culture.

These studies challenge the assumed relationship between categorical and dimensional emotion ratings and examine their predictive capacity. However, they face limitations in scalability and scope, making it difficult to generalize their results. We address these gaps with a large-scale, data-driven study that quantifies the relationship between these representations using subject-aware statistical analyses, enabling robust assessment of inter-subject variability.

Methodology

Dataset Selection

In our search for relevant datasets, we began by reviewing survey papers that analyzed the current landscape of emotion recognition. These papers provided an overview of the various modalities currently used for gathering emotional stimuli and how stimuli are being annotated (Gunes et al., 2011; Sethu et al., 2019). We noted the datasets mentioned in these articles, used a snowball method from these dataset papers to gather more datasets. We also conducted a separate search for the emotion recognition datasets through Google Scholar and Research networks (eg, ResearchGate). In total, we identified 210 datasets. Since our research focuses on specific rating models, we parsed the results to only include the datasets that contained both categorical and dimensional ratings, ultimately narrowing down our list to 35 datasets.

From this second list of datasets, we further collected information about modality, annotation strategy, and number of raters. The data modality is important for understanding how the affective information was presented to raters, either through internal expressions (biosignals) or external expressions (audio, visual). Meanwhile, the labels and granularity for the categorical and dimensional ratings illustrate the specificity and scope of the annotations. Finally, the datasets were categorized as "perceived" if ratings were based on what emotion the raters perceived in the stimuli, and "feeling" if the ratings were based on what emotions the stimuli invoked in the raters. We then removed the datasets that were based on feelings of participants, and papers that did not provide

rater-specific data.¹

We ultimately chose the MSP-Podcast Corpus, developed by Lotfian and Busso (2017), as it was the most in line with the goals of our study for three main reasons. First, the dataset features natural audio samples collected from podcasts, which increases the ecological validity of our research. Acted speech is a popular stimulus medium, but often mimics more prototypical behaviours, rather than capturing the ambiguous emotional expressions found in everyday interactions (Lotfian & Busso, 2017). Second, the annotation method for the datasets offers a high degree of granularity by including nine primary emotion labels and three common dimensional ratings measured on a 7-point scale, supplemented by self-assessment manikins (Bradley & Lang, 1994). Third, the dataset includes a large number of raters ($n = 13,672$) and ratings (853,302 observations), where each rater annotated multiple stimuli. The number of participants was determined by the number of unique participant IDs in the dataset, which allowed us to take individual differences into account. The inter-evaluator agreement was calculated using Krippendorff's alpha coefficient for Arousal ($\alpha = 0.426$), Valence ($\alpha = 0.459$), and Dominance ($\alpha = 0.402$) for the dimensional annotations, and Fleiss' kappa for the primary emotions ($\kappa = 0.229$) (Lotfian & Busso, 2017). Additionally, the dataset is easily accessible.

Ratings for the MSP-Podcast Corpus (Lotfian & Busso, 2017) were collected through Amazon Mechanical Turk, with no demographic information collected. A task-based survey format was used, with each task group consisting of multiple stimuli to allow the participant to calibrate their assessments across ratings. The dimensional and categorical rating scales were presented in tandem, allowing participants to rate both measures at the same time. To rate the dimensional values, a 7-point Likert scale was used to evaluate Valence (ranging from very negative to very positive), Arousal (from very calm to very active), and Dominance (from very weak to very strong). The scales were supplemented with self-assessment manikins to guide the evaluators' responses. The categorical emotion labels were divided into primary and secondary emotion categories. For the primary emotion, participants were given a list of emotions and asked to select the one that best represented the main emotion conveyed by the audio segment. This list included: Anger, Sadness, Happiness, Surprise, Fear, Disgust, Contempt, Neutral, and Other. If none of the options adequately described the emotional content, participants could choose the Other category and provide a free-form text response. For the secondary emotions, participants could select multiple labels to better describe the emotions they perceived in the audio segment, or "Other" to provide a free-form response. The survey format allowed participants to complete both dimensional and categorical scales simultaneously.

¹For materials, see: <https://github.com/SFU-HAI-Lab/Mapping-Emotion-Representations.git>

Data Cleaning

All preprocessing, clustering, and predictive modelling were conducted in Python 3.13.3 (Van Rossum & Drake, 2009), using the Natural Language Toolkit (NLTK) (Bird et al., 2009) and scikit-learn libraries (Pedregosa et al., 2011). The original dataset was restructured such that new columns were created for participant ID, primary emotion, secondary emotion, Arousal, Valence, and Dominance. Since the primary focus of this analysis was on the primary emotion categorizations, all columns relating to secondary emotions were removed.

Statistical Analysis

The cleaned MSP-Podcast Corpus was used to investigate the correlational relationship between categorical and dimensional emotion ratings. An analysis of variance was conducted using generalized linear mixed models (GLMM) that employed binomial logistic regression with Arousal, Valence, and Dominance as fixed effects, with participant specified as a random effect including a random intercept and independent random slopes for each predictor. We use the lme4 package (Bates et al., 2015) in RStudio (R Core Team, 2024) for GLMM analysis. The model was run separately for each emotion category including "Other". To examine whether the rating relationship differed among participants, two models were compared: a generalized linear mixed model that included participant ID as a random effect, and a generalized linear model that excluded this effect. The comparisons allowed for an assessment of the extent to which the inclusion of a random effect improved the model's fit.

Results

Data Distribution

Figure 1 shows box plots for the Valence, Arousal and Dominance ratings, respectively, grouped by emotion label. In Figure 1a, the Valence axis represents the 7-point rating scale, with 1.0 representing "very negative", 4.0 as "neutral", and 7.0 as "very positive." The Happy category showed the highest Valence range, with an interquartile range (IQR) of 5-6, while Angry and Sad had the lowest. The absence of a box plot for the Neutral category suggests that at least 50% of the ratings were 4.0.

In Figure 1b, the Arousal axis represents the 7-point rating scale, with the value 1.0 representing "very calm", 4.0 for "neutral", and 7.0 for "very active." Most emotion categories had IQRs between 4.0 and 6.0, indicating that most emotions were generally perceived as moderate to highly arousing.

The Dominance axis in Figure 1c also represents the 7-point rating scale, with 1.0 as "very weak", 4.0 as "neutral", and 7.0 as "very strong." Angry had the highest IQRs, falling between 5.0 and 6.0, while Sad had the lowest range, between 3.0 and 4.0. Otherwise, most emotions had an average rating range that hovered around 5.0, suggesting that most ratings showed a moderate degree of Dominance.

Correlational Analysis

Table 1 summarizes the results of the Type III Wald chi-square test and the corresponding effect sizes (Cohen's *d*) for each dimensional rating across the emotion categories of the generalized linear mixed models. The results showed that for nearly every emotion category, the Arousal, Valence, and Dominance intervals significantly predicted the emotion label at $p < .001$. The only exceptions were the "Contempt" category with a $p < .01$ significance in Arousal and the "Surprise" category, which showed a non-significant effect of Dominance.

Among the three dimensions, Valence appeared to be the most robust predictor with the highest overall χ^2 values and effect sizes. Specifically, the Happy category showed a statistically significant effect of Valence with a medium positive effect size. The Sad, Disgust, Contempt and Angry categories exhibited significant effects with medium negative effect sizes on Valence. The remaining categories showed a very small effect of Valence. Arousal and Dominance were comparatively weaker predictors across the emotion categories. The Neutral category was a notable exception, as it showed significant moderate negative effects for Arousal. Additionally, Anger, Sad and Surprise both exhibited small, yet significant effects for Arousal. For Dominance, most effect sizes were negligible except for small effect sizes for Sad, Neutral and Anger.

Participant Variability

Model comparisons revealed significant improvements in fit when including participant-level random intercepts and independent random slopes for Valence, Arousal, and Dominance across all emotion categories (see Table 2). Both baseline differences between participants and variability in the effects of each dimension was significant ($p < .001$). The magnitude of these effects varied across emotions. Furthermore, the wide whiskers observed in the box plots shown in Figure 1 indicate a broad range of ratings that span the whole rating scale.

Discussion

Our study provides several insights into the relationship between categorical and dimensional emotion representations, showing that while dimensional ratings can predict categorical labels, this relationship is uneven across emotions and varies across individuals.

First, the correlational analysis (RQ1) confirmed a significant but incomplete relationship between the two models. Valence emerged as the most robust and practically meaningful predictor of categorical labels with the largest effect sizes 1, suggesting that evaluative polarity (positive vs. negative) plays a dominant role in mapping affect to discrete categories. This aligns with prior findings indicating that certain dimensions, particularly Valence, provide the most stable signal across contexts and annotators (Morgan, 2019). In contrast, Arousal and Dominance contributed weaker, category-specific effects, though they remained theoretically meaningful (e.g., Sadness showed negative associations with both,

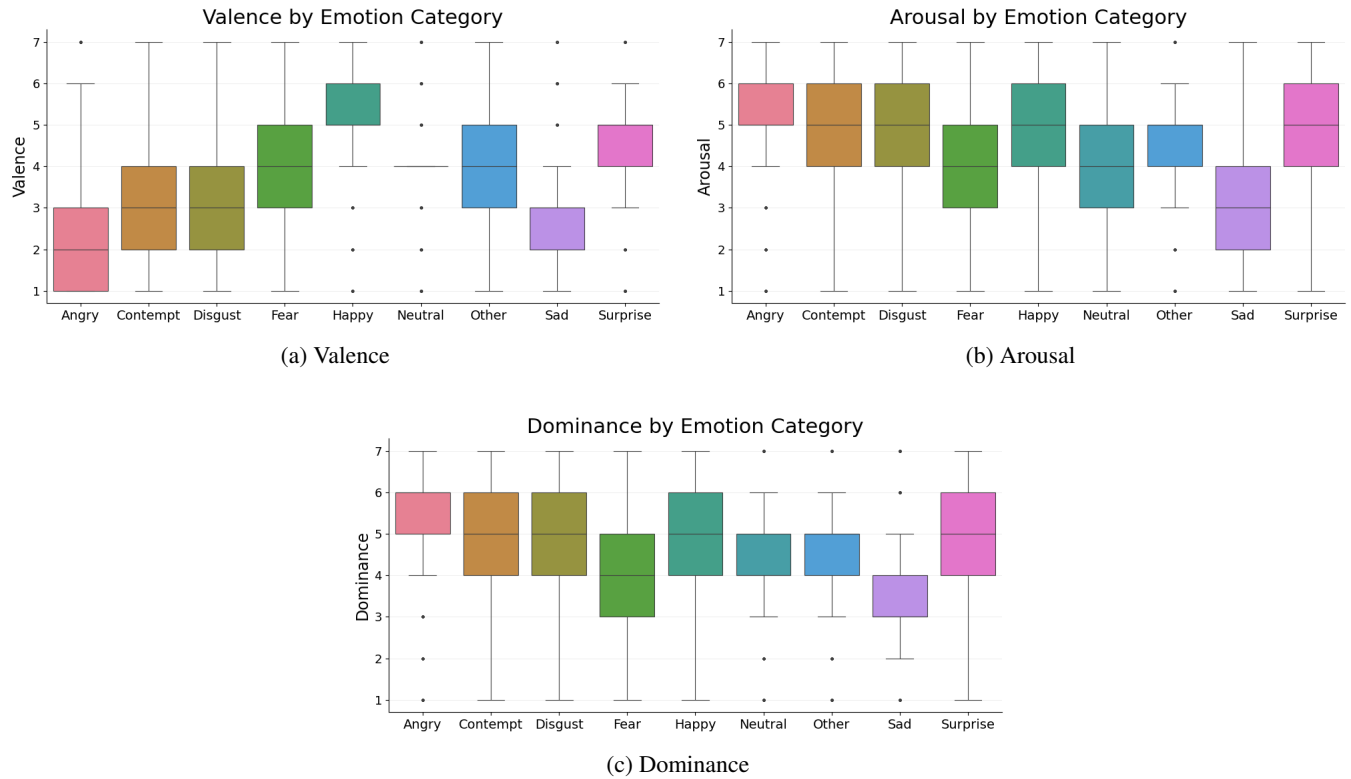


Figure 1: Boxplots for Valence, Arousal and Dominance, per each emotion category.

Table 1: $\chi^2(1, N = 853302)$ and Effect Size for Arousal, Valence, and Dominance Across Emotions

Emotion	Arousal		Valence		Dominance	
	χ^2	d	χ^2	d	χ^2	d
Sad	4357.7**	-0.14	10522.6**	-0.22	5748.4**	-0.16
Neutral	11723.2**	-0.23	2523.8**	0.11	2437.5**	-0.11
Happy	1598.4**	0.09	42184.2**	0.44	75.7**	-0.02
Surprise	2011.1**	0.14	2541.2**	0.14	1.6	0.004
Fear	83.6**	0.02	2275.7**	-0.10	1685.7**	0.09
Disgust	69.7**	0.02	13347.8**	-0.25	276.6**	0.04
Anger	4628.6**	0.15	28404.7**	-0.36	2068.9**	0.10
Contempt	9.4*	0.0	9130.3**	-0.21	1514.1**	0.08

Note. d , Cohen's d ; ** $p < .001$; * $p < .01$

consistent with low activation and control).

These patterns are supported by Figure 1, which shows substantial overlap across categories in VAD space. Valence provides the clearest separation, while Arousal is high for most categories and thus less discriminative. Dominance follows a coherent control-related pattern but with smaller effects. The wide distributions indicate high within-category variability and cross-category overlap, suggesting graded, heterogeneous states rather than discrete clusters.

At the category level, several patterns highlight limits of VAD-based translation. Interestingly, Neutral displayed a small positive Valence effect paired with a small negative Dominance effect. This suggests that “Neutral” may not operate purely as a midpoint on all dimensions, but rather as a label

applied to affectively ambiguous or low-intensity content that may still be perceived as slightly pleasant or non-threatening. Such patterns highlight how neutral categories in large corpora may serve as “catch-all” labels rather than representing a psychologically uniform affective state, with implications for dataset construction and evaluation. Contempt and Disgust showed the same pattern in boxplots and low effect sizes in Arousal and Dominance in the mixed-effects model, suggesting that representation of these emotions may not be captured by VAD alone. Surprise showed small positive effects for Valence and Arousal, but no reliable association with Dominance, aligning with the interpretation of surprise as novelty-driven rather than control-related. These mismatches indicate that dimensional–categorical alignment varies systematically

Table 2: Model Fit Improvement from Random Effect Across Emotions $\chi^2(1, N = 853302)$, P annotates participant.

Emotion	P	valence P	arousal P	dom P
Sad	7943**	8304**	1351**	1281**
Neutral	14338**	3202**	6430**	4247**
Happy	26377**	28984**	1383**	1113**
Surprise	3416**	1464**	787**	183**
Fear	861**	2874**	141**	1515**
Disgust	651**	3081**	131**	913**
Anger	4310**	2936**	286**	3191**
Contempt	2119**	6900**	2992**	2079**

Note. ** $p < .001$

across emotions.

Finally, the participant-level analyses (RQ2) revealed substantial inter-individual variability (see Table 2). Models incorporating participant IDs as random effects significantly outperformed those without, and dimensional ratings varied widely across participants for the same categorical labels. This finding emphasizes that although dimensional ratings correlate with emotion categories, these dimensions are not used uniformly across individuals, but are instead interpreted in participant-specific ways. This suggests that variability in emotion perception is not only noise, but reflects systematic differences in how individuals interpret affective dimensions. This variability has direct implications for affective computing: models trained on aggregated labels may capture a population-average mapping that generalizes poorly to new users or contexts, motivating greater attention to individual differences and potentially personalized or adaptive affect modeling. More broadly, the results challenge the common assumption of a shared, universal mapping between dimensional affect and categorical emotion, suggesting that computational models should incorporate mechanisms for adapting to individual differences in affective interpretation.

A key contribution of this work is providing empirical evidence that supports translation between categorical and dimensional emotion representations, but also clarifying when such translation is likely to be reliable. Categories with strong and interpretable VAD profiles (e.g., Happy and Sad) may allow more stable cross-dataset alignment and conversion. In contrast, categories with weak or inconsistent VAD signatures (e.g., Fear and Surprise) require more caution and may benefit from additional representational dimensions, contextual features, or appraisal-based annotations. Furthermore, while the coefficients derived in Table 1 could support the mapping of discrete categories to the VAD space, they do not account for individual differences and therefore are insufficient.

Practically, these findings support a more nuanced approach to interoperability: rather than assuming categorical and dimensional datasets are directly comparable, researchers should consider (1) which categories exhibit strong dimensional structure, (2) which dimensions are most informative for their task, and (3) how individual differences in

ratings may impact generalization. In large-scale affective datasets, reporting effect sizes and participant-level variability—alongside significance tests—can improve interpretability and help prevent overclaiming due to large sample sizes.

Limitations and Future Work

This study has several limitations. First, the observed effects may be influenced by dataset composition, including class imbalance and the prevalence of Neutral and Happy categories. Second, while VAD provides a compact and widely used dimensional representation, it may not capture critical appraisal or contextual variables that structure emotions such as fear and surprise. Third, although random effects analyses demonstrate participant variability, future work should directly examine what drives these differences (e.g., cultural background, emotion vocabulary, rater calibration, or systematic response biases such as scale use).

Future research could extend this work by testing whether the observed mapping replicates across datasets with different elicitation conditions (e.g., acted vs. spontaneous emotion, speech vs. text), by incorporating additional dimensions (e.g., secondary emotion categories, "other" category), or by evaluating whether personalized models improve categorical-dimensional alignment at the individual level. Moreover, studies should consider the integration of formal or declarative emotion models, such as appraisal or logic-based approaches, in modelling emotions in emotion recognition systems. Such efforts would advance both theoretical understanding and the development of more robust, human-aligned affective computing systems. Likewise, repeating these analyses with multiple datasets will help to validate our claims.

Conclusion

This study explored the relationships between categorical and dimensional models of emotions using a large, naturalistic speech dataset. The results demonstrate that Valence, Arousal, and Dominance are significantly associated with discrete emotion categories, but that this relationship is neither uniform across emotions nor consistent across individuals. In particular, we show that participant-level variability is a central feature of emotion representation: individuals systematically differ in how they map dimensional affective signals onto categorical labels. These findings refine the link between dimensional and categorical theories by showing that their correspondence is conditional—strong for some emotions, but weaker or more ambiguous for others—and shaped by individual interpretation. Importantly, this variability reflects meaningful differences in affective perception rather than noise, and is obscured by aggregate analyses. By highlighting the role of individual differences, this work challenges the assumption of a universal mapping between dimensional and categorical emotion representations. It suggests that both theoretical and computational models of emotion should account for variability in how affect is perceived and labeled, contributing to more flexible and ecologically valid representations of emotions.

References

- Amir, N., Rubinstein, R., Shlomov, A., & Diamond, G. (2015). Comparing categorical and dimensional ratings of emotional speech. *2015 11th IEEE International Conference and Workshops on Automatic Face and Gesture Recognition*, 5, 1–5.
- August, B. C., Benally, C. D., & Cadena, D. E. (2007). An example conference paper. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 29, 100–105. <https://escholarship.org/uc/item/k0e1y2>
- Bakker, I., van der Voordt, T., Vink, P., & de Boon, J. (2014). Pleasure, arousal, dominance: Mehrabian and russell revisited. *Current Psychology*, 33, 405–421. <https://doi.org/10.1007/s12144-014-9219-4>
- Barrett, L. F. (2006a). Are emotions natural kinds? *Perspectives on psychological science*, 1(1), 28–58.
- Barrett, L. F. (2006b). Solving the emotion paradox: Categorization and the experience of emotion. *Personality and social psychology review : an official journal of the Society for Personality and Social Psychology*, 20–46. https://doi.org/https://doi.org/10.1207/s15327957pspr1001_2
- Barrett, L. F., Adolphs, R., Marsella, S., Martinez, A. M., & Pollak, S. D. (2019). Emotional expressions reconsidered: Challenges to inferring emotion from human facial movements. *Psychological science in the public interest*, 20(1), 1–68.
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1), 1–48. <https://doi.org/10.18637/jss.v067.i01>
- Bird, S., Klein, E., & Loper, E. (2009). *Natural language processing with python*. O'Reilly Media, Inc.
- Bradley, M. M., & Lang, P. J. (1994). Measuring emotion: The self-assessment manikin and the semantic differential. *Journal of Behavior Therapy and Experimental Psychiatry*, 25(1), 49–59. [https://doi.org/https://doi.org/10.1016/0005-7916\(94\)90063-9](https://doi.org/https://doi.org/10.1016/0005-7916(94)90063-9)
- Calvo, R. A., & D'Mello, S. (2010). Affect detection: An interdisciplinary review of models, methods, and their applications. *IEEE Transactions on affective computing*, 1(1), 18–37.
- Carlson, J. G., & Hatfield, E. (1992). *Psychology of emotion*. Harcourt Brace Jovanovich College Publishers.
- Cowen, A., & Keltner, D. (2017). Self-report captures 27 distinct categories of emotion bridged by continuous gradients. *Proceedings of the National Academy of Sciences*, 114(38), E7900–E7909. <https://doi.org/10.1073/pnas.1702247114>
- Daphne, E. F., & Echo, F. G. (2022). An example journal article. *Journal of Example Studies*, 10(3), 200–215. <https://doi.org/10.1234/eg.article>
- Ekman, P. (1999). Basic emotions. In *Handbook of cognition and emotion* (pp. 45–60). John Wiley Sons, Ltd. <https://doi.org/10.1002/0470013494>
- Ekman, P., & Cordaro, D. (2011). What is meant by calling emotions basic. *Emotion Review*, 3(4), 364–370. <https://doi.org/10.1177/1754073911410740>
- Fitzgerald, G. H., & Galli, H. I. (1985). *An example technical report* (Report No. EX-TR-85-1). Some University, Department of Hypothetical Studies.
- Gong, X., Wong, N., & Wang, D. (2018). Are gender differences in emotion culturally universal? comparison of emotional intensity between chinese and german samples. *Journal of Cross-Cultural Psychology*, 49(6), 993–1005. <https://doi.org/10.1177/0022022118768434>
- Gray, J. A. (1990). Brain systems that mediate both emotion and cognition. *Cognition and Emotion*, 4(3), 269–288. <https://doi.org/10.1080/02699939008410799>
- Gunes, H., Schuller, B., Pantic, M., & Cowie, R. (2011). Emotion representation, analysis and synthesis in continuous space: A survey. *2011 IEEE International Conference on Automatic Face Gesture Recognition (FG)*, 827–834. <https://doi.org/10.1109/FG.2011.5771357>
- Hakuole, I. J. (2001). *An example thesis title* [Doctoral dissertation, Some University]. SU Campus Repository. <https://hdl.handle.net/20.1000/5555>
- Hoffmann, H., Scheck, A., Schuster, T., Walter, S., Limbrecht, K., Traue, H. C., & Kessler, H. (2012). Mapping discrete emotions into the dimensional space: An empirical approach. *2012 IEEE International Conference on Systems, Man, and Cybernetics (SMC)*, 3316–3320. <https://doi.org/10.1109/ICSMC.2012.6378303>
- Issa, J. K. (1963). An example book chapter in an edited volume. In K. L. Juliet & L. M. Kapoor (Eds.), *Example edited volume title* (pp. 300–315). Some Academic Press.
- Jack, R. E., Garrod, O. G. B., Yu, H., Caldara, R., & Schyns, P. G. (2012). Facial expressions of emotion are not culturally universal. *Proceedings of the National Academy of Sciences of the United States of America*, 109, 7241–7244. <https://doi.org/10.1073/pnas.1200155109>
- Lobsang, M. N. (Ed.). (2023). *Example edited volume title* (2nd ed.). Some University Press. <https://doi.org/10.4321/eg.vol2>
- Lotfian, R., & Busso, C. (2017). Building naturalistic emotionally balanced speech corpus by retrieving emotional speech from existing podcast recordings. *IEEE Transactions on Affective Computing*. <https://doi.org/10.1109/TAFFC.2017.2736999>
- Ma, H., & Yarosh, S. (2021). A review of affective computing research based on function-component-representation framework. *IEEE Transactions on Affective Computing*, 14(2), 1655–1674.
- Marsella, S., Gratch, J., & Petta, P. (2010). Computational models of emotion. In K. R. Scherer, T. Banziger, & E. B. Roesch (Eds.), *A blueprint for affective computing: A sourcebook and manual*. Oxford University Press.
- Mehrabian, A. (1996). Pleasure-arousal-dominance: A general framework for describing and measuring individual dif-

- ferences in temperament. *Current psychology*, 14(4), 261–292.
- Mehrabian, A., & Russell, J. A. (1974). *An approach to environmental psychology*. The MIT Press.
- Mitanni, N. O., & November, O. P. (1972). *Example book title*. Some Publisher.
- Morgan, S. (2019). Categorical and dimensional ratings of emotional speech: Behavioral findings from the morgan emotional speech set. *Journal of Speech, Language, and Hearing Research*, 62, 4015–4029. https://doi.org/10.1044/2019_JSLHR-S-19-0144
- Oatley, K., & Jenkins, J. M. (1992). Human emotions: Function and dysfunction. *Annual review of psychology*, 43(1), 55–85. <https://doi.org/10.1146/annurev.ps.43.020192.000415>
- Ogawa, T., Fujimura, T., & Suzuki, N. (2005). Perception of facial expressions of emotion under brief exposure duration. *Japanese Journal of Research on Emotions*, 12(1), 1–11. <https://doi.org/10.4092/jsre.12.1>
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., & Duchesnay, É. (2011). Scikit-learn: Machine learning in python. *Journal of Machine Learning Research*, 12(85), 2825–2830.
- Picard, R. (2015). The promise of affective computing. In *The oxford handbook of affective computing* (pp. 11–20). Oxford Library of Psychology.
- R Core Team. (2024). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. Vienna, Austria. <https://www.R-project.org/>
- Russell, J. A. (1980). A circumplex model of affect. *Journal of personality and social psychology*, 39, 1161–1178. <https://doi.org/10.1037/h0077714>
- Russell, J. A. (1994). Is there universal recognition of emotion from facial expression? a review of the cross-cultural studies. *Psychological Bulletin*, 115, 102–141. <https://doi.org/10.1037/0033-2909.115.1.102>
- Russell, J. A., & Fernández-Dols, J. M. (1997). *The psychology of facial expression*. Cambridge University Press.
- Scherer, K. R. (2010). Emotion and emotional competence: Conceptual and theoretical issues for modeling. In K. R. Scherer, T. Banziger, & E. B. Roesch (Eds.), *A blueprint for affective computing: A sourcebook and manual*. Oxford University Press.
- Sethu, V., Provost, E. M., Epps, J., Busso, C., Cummins, N., & Narayanan, S. (2019). The ambiguous world of emotion representation. <https://arxiv.org/abs/1909.00360>
- Van Rossum, G., & Drake, F. L. (2009). *Python 3 reference manual* [Version 3.13.3]. CreateSpace.
- Wang, Y., Song, W., Tao, W., Liotta, A., Yang, D., Li, X., Gao, S., Sun, Y., Ge, W., Zhang, W., et al. (2022). A systematic review on affective computing: Emotion models, databases, and recent advances. *Information Fusion*, 83–84. <https://doi.org/10.1016/j.inffus.2022.03.009>