

MSP-Podcast

Building Naturalistic Emotionally Balanced Speech Corpus by Retrieving Emotional Speech From Existing Podcast Recordings

Reza Lotfian & Carlos Busso

IEEE Transactions on Affective Computing, 10(4), 471–483 (2019)

Multimodal Signal Processing Laboratory, The University of Texas at Dallas

Corpora II: naturalistic and in-the-wild data

Why this paper matters



It is a methods paper disguised as a corpus paper

403

podcasts downloaded under
permissive Creative Commons licences

84,125

candidate speaking turns
in the unlabelled pool

2,317

turns retrieved and annotated
in this proof of concept

278

crowd workers, at least five
ratings per sentence

● The proposal

Do not record an emotional corpus. Harvest one. Take spontaneous speech that already exists on audio-sharing sites, filter it automatically into clean single-speaker turns, then use emotion models trained on existing corpora to decide which of those turns are worth paying humans to annotate.

● Why the numbers above are the wrong way to judge it

2,317 sentences is small — smaller than IEMOCAP. But the paper is not offering a corpus, it is offering a protocol that scales. By the time it appeared in print the same pipeline had produced over 27 hours; the corpus now stands at 409 hours and is the largest naturalistic emotional speech resource in the field.

1

The problem with naturalistic corpora

Four limitations, and one that nobody had solved



Four limitations of existing emotional corpora



The paper's own diagnosis, and the framework for everything that follows

● Lack of naturalness

Actors reading or improvising produce prototypical displays, not the ambiguous behaviour of daily interaction. Recognisers trained on them degrade in deployment.

● Unbalanced emotional content

In spontaneous recordings the emotion distribution is dictated by the recording protocol. Couples arguing gives you negative valence; colloquial chat gives you positive.

● Limited size

Most emotional corpora run to a few hours. Only IEMOCAP, MSP-IMPROV, TUM-AVIC and FAU-AIBO exceeded nine hours. Deep models need far more.

● Limited speakers

Most corpora have fewer than twenty speakers. Only CREMA-D, FAU-AIBO and Chen Bimodal exceed fifty — and none has enough audio per speaker for speaker-verification work.

● The second one is the interesting problem

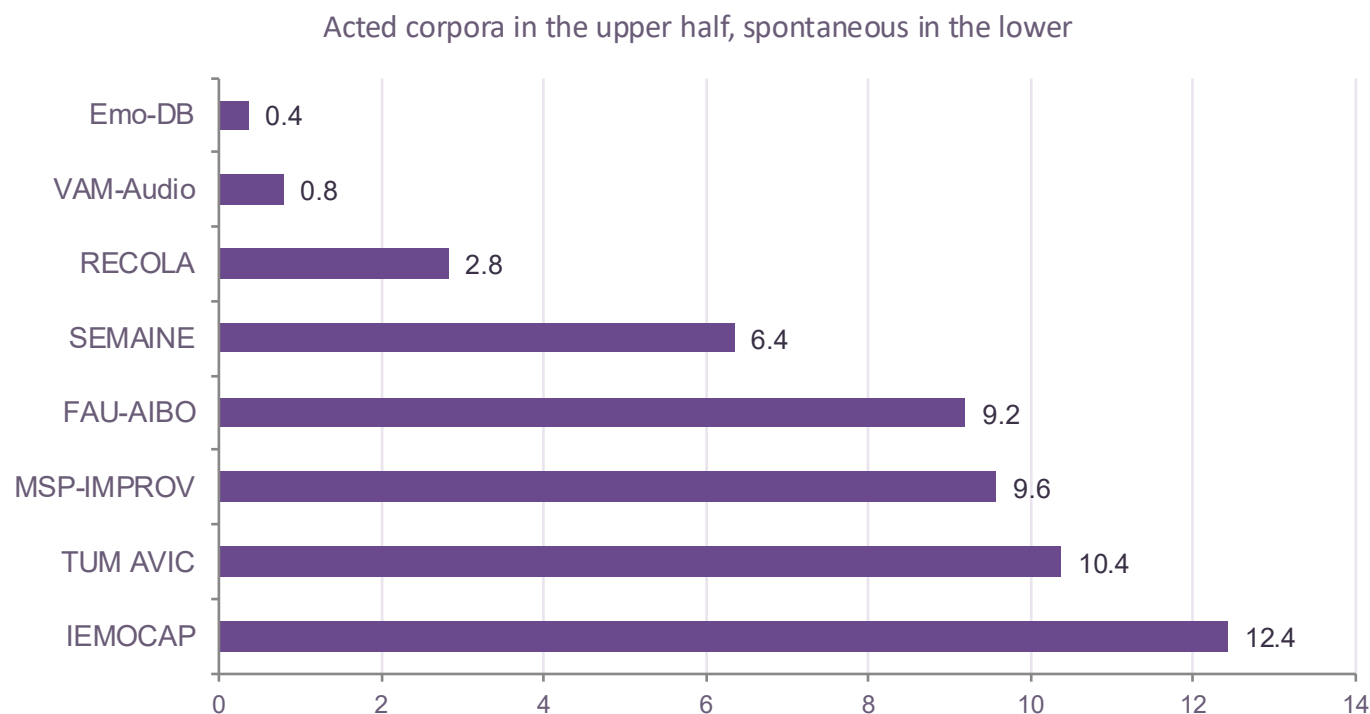
Size and speaker count are engineering problems — collect more. Naturalness is a known tradeoff. But balance is genuinely hard: you cannot ask spontaneous speech to be emotionally diverse, because the very thing that makes it spontaneous is that nobody is controlling what happens.

This paper's whole contribution is a way to get balance without giving up spontaneity.

The landscape in 2017



Duration of the corpora the paper surveys, in hours



- Nothing in the field exceeded thirteen hours. Compare Librispeech, where 960 labelled hours was the standard supervised setup.
- Speaker counts are similarly small: 10 for IEMOCAP, 12 for MSP-IMPROV, 20 for SEMAINE, 46 for RECOLA.
- The exception is CREMA-D, with 91 speakers — but only 7,442 short clips, and acted.
- So the field had a choice between naturalistic and small, or acted and slightly less small. Neither supports the deep models arriving at the time.

The paper is explicit that this is the bottleneck: the community cannot use powerful learning algorithms because it does not have the data to train them on.

An uncomfortable finding in the survey



The most emotionally balanced corpora are the acted ones

● How the paper measures balance

Plot every speaking turn as a point in the arousal–valence plane, then colour each region by the average distance to its twenty nearest neighbours. Dark regions are densely populated; light regions are gaps.

A thousand turns are sampled from each corpus so the comparison is size-independent.

● What the maps show

IEMOCAP and MSP-IMPROV cover the space most evenly — both acted, both with scenarios chosen to elicit target emotions.

SEMAINE has almost no negative valence. RECOLA is mostly neutral and mildly positive. VAM is concentrated in negative valence, because it comes from a conflict-driven talk show.

● The bind this creates

Acting is what buys you balance. The scenarios in IEMOCAP exist precisely so that anger, sadness, happiness and frustration all appear in reasonable quantity. Remove the actors and you lose the control that produced the coverage.

So the field faced a genuine dilemma, not a tradeoff to be tuned: naturalness and balance appeared to be obtainable only at each other's expense. The retrieval idea in section 3 is an attempt to break that.

2

From podcast to candidate turn

An almost entirely automatic pipeline



Why podcasts



Four properties that make them unusually well suited

● **Redistributable**

Only episodes under CC-BY or CC-BY-SA are downloaded — the least restrictive clauses, which allow the team to modify, redistribute and even commercialise the resulting corpus. Almost no other naturalistic source clears this bar.

● **Music-free by construction**

Producers releasing under these licences generally cannot include broadcast music, because they do not hold the distribution rights. The licence filter is doing acoustic work as a side effect.

● **Topically diverse**

Episodes are chosen manually across science, technology, politics, economics, business, arts, culture, medicine, lifestyle and sport — so the emotional range is not fixed by a single recording scenario.

● **Unscripted and unprofessional**

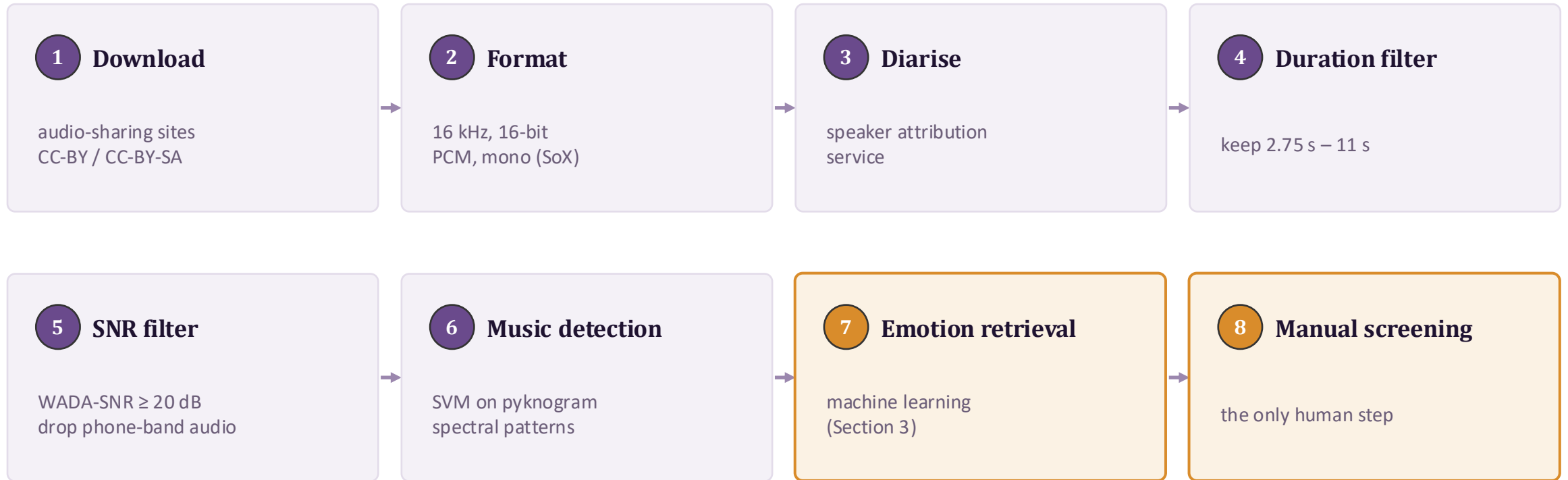
Many podcasts have no fixed transcript and no professional host. The paper argues that ordinary people recording conversations produce more lifelike speech across a broader emotional range than radio presenters would.

Also excluded by design: acted material. Searches target interviews, talk shows, news, discussions, education, storytelling and debates.

The processing pipeline



One manual step, and it is not where you would expect



Every stage except the last is automatic, which is what makes the approach scale. The human reviews only the few hundred segments that survive retrieval — not the tens of thousands in the pool.

What the filters actually reject



And the reasoning behind each threshold

● 2.75 to 11 seconds

A real tradeoff. Longer segments let the emotion drift, so a single label stops being accurate; shorter ones give raters too little to judge and produce unreliable acoustic features.

● Speech activity detection

Diarisation can still return a 2.75-second segment that is mostly silence, so a separate detector removes samples with too little actual speech.

● SNR of at least 20 dB

Estimated with waveform amplitude distribution analysis. Noisy recordings are discarded outright rather than kept and flagged.

● No telephone-band audio

Phone speech is sampled at 8 kHz, so it lacks the spectral content above 4 kHz that emotion features rely on. Segments with no significant energy there are removed.

● No music or background music

A support vector machine over pyknogram spectral patterns, trained on 105 manually segmented podcasts. Of 500

The paper reports that last figure plainly and adds: there is room for improvement. Honest, and worth noticing.

What survives



The pool, and the one place a human still looks

403 → 84,125

podcasts in, candidate speaking
turns out

65.6%

of retrieved segments passed
manual screening

22.3%

rejected because more than
one speaker was present

● Why segments were rejected

Multiple speakers is the dominant failure, at 22.3% — diarisation and the absence of an overlapped-speech detector are the weak links.

Other reasons: residual silence, undetected music, and content problems such as offensive language, explicit sexual references, or speech in languages other than English.

● What the pool is, and is not

84,125 turns is the reservoir, not the corpus. Almost all of it is expected to be emotionally neutral — which is exactly the balance problem from section 1, now sitting in front of them as a concrete pile of audio.

The question the rest of the paper answers: which few hundred of these do you pay people to listen to?

Notice the economics. Annotation is the expensive step, so every design decision upstream exists to make sure the annotation budget lands on segments that are worth annotating.

3

Retrieval

Buying emotional balance with machine learning



The central idea



Reframe corpus construction as an information retrieval problem

● The formulation

Train emotion models on corpora that already have labels. Run them over the unlabelled pool. Use their outputs to rank the pool, and annotate the extremes.

Four retrieval targets, run separately: high arousal, low arousal, high valence, low valence. Attacking the four corners is what produces coverage of the space.

● Why attributes rather than categories

Two reasons the paper gives. First, emotional category sets differ wildly between corpora, so training a category model across them is awkward. Arousal and valence annotations exist almost everywhere.

Second, dominance correlates strongly with arousal, so it adds little; only two dimensions are used.

● What is really being proposed here

Cross-corpus emotion recognition is usually treated as a benchmark — you train on one corpus, test on another, and report how much accuracy you lose. This paper treats it as a tool instead. The models do not need to be accurate enough to deploy; they only need to be better than chance at sorting a list.

That is a much lower bar, and it turns a known weakness of the field into a way of building the data that would fix it.

Five methods across three families



Each turns the retrieval problem into a different learning problem

● Classification

Train a binary classifier for high versus low, then rank by distance from the decision boundary.

SVM with a linear kernel

Bayesian-optimal classifier for dichotomised labels (BOC-DL), which uses the original continuous scores rather than throwing them away when binarising

● Preference learning

Train on pairs — 'A has higher arousal than B' — and rank the pool by a learned scoring function.

Rank-SVM

Gaussian-process ranking, a probabilistic kernel method for preference learning

● Regression

Predict the attribute value directly and take the top and bottom of the sorted list.

Support vector regression

● Why preference learning is a natural fit

Emotional attribute scores are not really absolute quantities — a rating of 4 from one annotator is not the same as a 4 from another. Relative judgements survive that; absolute ones do not. Ranking sidesteps the calibration problem entirely.

How the retrieval models were trained



Two acted corpora, used as a springboard for a natural one

13,222

sentences from IEMOCAP
and MSP-IMPROV combined

50 / 50

speaker-independent train
and test partitions

eGeMAPS

the Geneva minimalistic
acoustic parameter set

P@K

precision at K — the
retrieval evaluation metric

● Why eGeMAPS

A small, theoretically motivated set of prosodic, spectral and voice-quality parameters, chosen for their standing as markers of affective change.

The paper's stated reason for using it: the set is small enough that no feature selection is needed, which makes the results reproducible by other groups.

● How precision at K is defined here

Split the test set at the median for each attribute. A retrieved sample counts as correct if it falls on the intended side.

P@10 of 87% for arousal means: retrieve the top 10% and bottom 10% of the ranked list, and 87% of those land on the correct side. P@50 covers everything, so it equals binary accuracy.

Note the irony sitting inside the method: the models that find naturalistic emotional speech are trained entirely on acted speech, from twenty-two actors between the two corpora. We will come back to what that means.

Which methods work



Retrieval performance on the held-out acted data

● Arousal is easy

Support vector regression, Rank-SVM and Gaussian-process ranking are all competitive, with precision above 80% when 20% of the list is retrieved.

Arousal is carried by loudness, pitch range and speaking rate — properties that acoustic features capture directly.

● Valence is hard

Precision rates are markedly lower across every method. The two ranking approaches do best, with BOC-DL also competitive.

This is a standing result in the field: few acoustic parameters discriminate valence, which is why text and facial cues usually carry it.

● The three methods taken forward

BOC-DL for classification, Gaussian-process ranking for preference learning, and support vector regression for regression — the best in each family. Using three rather than one hedges against any single model's blind spots being baked into the corpus.

The arousal–valence asymmetry is the single most consequential result in the paper. It shapes what the corpus ends up containing, and it recurs in the final results. Watch for it.

What was actually retrieved



Twelve targeted conditions plus a control

	GP-rank	SVR	BOC-DL
High arousal	200	200	200
Low arousal	200	200	200
High valence	200	200	200
Low valence	200	200	200
Random control	100 sentences drawn at random from the pool		

2,400 + 100

sentences requested — but the three methods often selected the same segments

2,317

distinct sentences after removing the overlap, all annotated

Rejected segments were replaced by the next entries in each sorted list until 200 were reached per condition.

4

Annotation

Crowdsourcing that watches its workers in real time



The questionnaire



Three attributes, one primary emotion, and an open secondary set

● Attributes

Valence, arousal and dominance, each on a seven-point Likert scale.

Self-assessment manikins guide the ratings visually — the same lexicon-free instrument used in IEMOCAP, chosen because raters do not have to share a definition of any emotion word.

● Primary emotion

One choice only, from: anger, sadness, happiness, surprise, fear, disgust, contempt, neutral, and other with a free-text box.

Nine options, against the five that IEMOCAP originally targeted and the six in CREMA-D.

● Secondary emotions

Select all that apply, from an extended list adding amused, frustrated, depressed, concerned, disappointed, excited, confused and annoyed.

Similar classes are grouped in the interface to reduce cognitive load.

● Why the secondary layer exists

Naturalistic speech routinely conveys blends that no single label describes — sadness with frustration, amusement with contempt. Forcing one choice discards that; recording both a forced primary and an unforced multi-label secondary keeps it.

Compare the two corpora from the first seminar: IEMOCAP allowed multiple labels but shipped a majority vote; CREMA-D forced six-way choice with no 'other'. This design keeps both readings.

Real-time quality control



The most transferable piece of craft in the paper

● The mechanism

Reference sentences whose emotional content is already known are interleaved with new ones. The system tracks each worker's agreement with those references as they go, and stops the task when performance drops.

Workers do at least 12 sentences and may do up to 100 if quality holds.

● Three changes on the previous protocol

One reference every four new sentences instead of five in twenty — finer temporal resolution, and overhead falls from 28% to 20%.

The reference position within each block is randomised, so workers cannot tell when they are being tested.

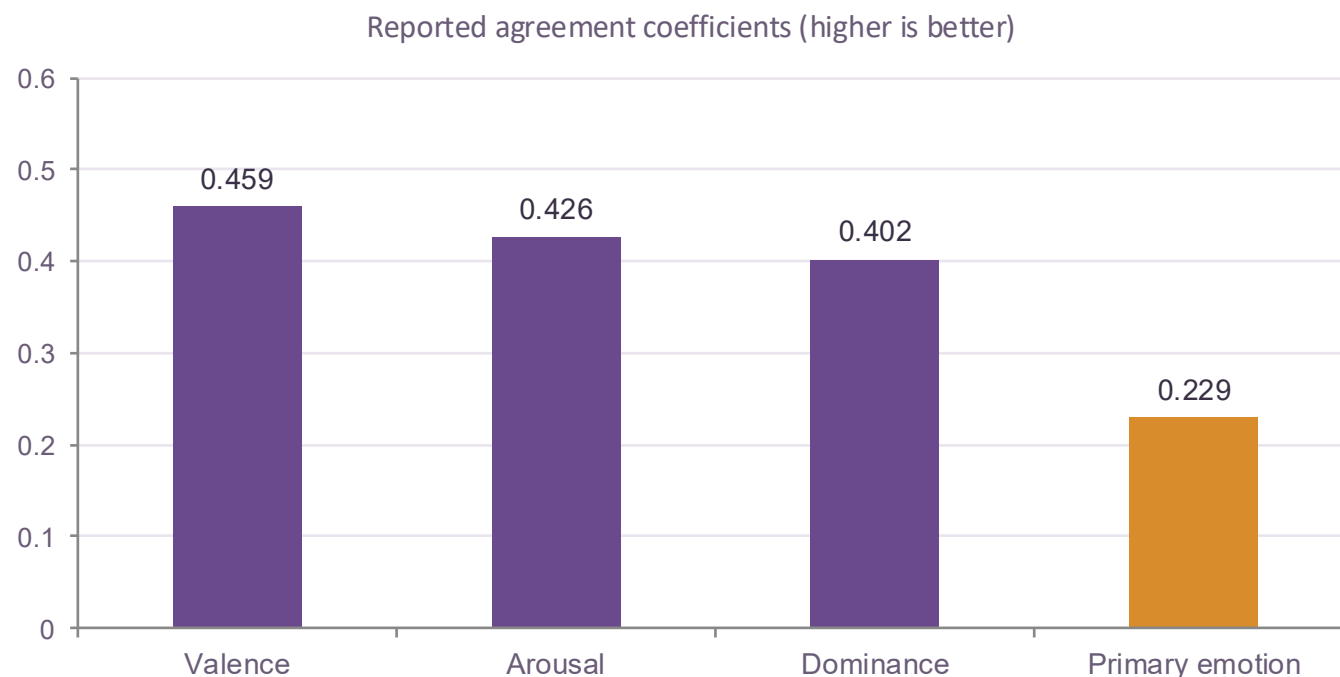
The stopping rule now watches arousal and valence as well as primary emotion.

- Two thresholds per metric. Average: the new labels are as good as those in the reference set. Low: the new labels beat only 10% of reference annotations.
- Evaluation stops if any metric falls below the low threshold, or if two metrics fall below the average threshold, or if the worker quits after the first 12 sentences.
- Payment is designed to match: a fixed rate for the first 12 sentences, then a bonus per sentence at twice that rate — so staying and staying accurate pays better than starting a fresh task.
- Quality is assessed over the three most recent reference sentences, so a worker who tires is caught within minutes rather than after twenty.

How reliable are the resulting labels?



Agreement across at least five raters per sentence



- Valence is the most reliable attribute, as it was in IEMOCAP — judging pleasantness is easier than judging arousal.
- Primary emotion is far weaker, and the paper explains why: the category set grew from five options to nine, and every added option costs agreement.
- A wording note worth catching: the table labels all four coefficients α , while the text says primary emotion was measured with Fleiss' kappa. Quote the number, and say which you mean.
- These values are in the same fair-to-moderate band we found for IEMOCAP and CREMA-D. Emotion annotation lives here.

Read this alongside the previous slide. All that quality-control machinery buys agreement that is still only moderate — which tells you how much of the disagreement is irreducible rather than sloppy.

5

Results and appraisal

Did retrieval deliver balance?



Did retrieval work?



Judged by the human ratings the retrieved segments actually received

● Arousal: clearly yes

The histograms of high-arousal and low-arousal retrievals separate cleanly for all three methods. The models found what they were asked to find.

● Valence: partially

There is a real difference between the high and low sets, but the separation is much weaker. Gaussian-process ranking gives the best valence separation of the three.

● The control: exactly as predicted

The 100 randomly drawn sentences average close to zero on both arousal and valence.

Most podcast speech is emotionally neutral. Random sampling would have wasted the annotation budget.

● A leakage the authors flag themselves

Segments retrieved for high or low valence tend to have high arousal scores as well. The arousal rankers, by contrast, return samples that are symmetric in valence.

So the coverage is better than any comparable natural corpus, but it is not uniform: the low-arousal corners of the space stay thin.

● The headline claim

Redraw the dispersion map from section 1 with the retrieved MSP-Podcast segments and it has smaller light areas than SEMAINE, RECOLA or VAM. Naturalistic recordings with coverage approaching that of designed acted corpora.

That is the result the paper set out to demonstrate, and it does.

What the corpus is not balanced for



Attribute balance does not deliver category balance

- The retrieval targeted arousal and valence. Nothing in the pipeline targeted emotional categories.
- The resulting distribution of primary emotions is lopsided: neutral, happiness and anger dominate; fear and sadness are scarce.
- A substantial share of segments reach no majority agreement at all and carry no categorical label — the same silent filtering step we found in IEMOCAP.
- So a corpus that is well spread across the arousal–valence plane can still be unusable for six-way categorical classification.

● Why this happens

Categories and attributes are not interchangeable. Fear and sadness occupy regions that arousal and valence rankers do not specifically target, so nothing pulls them into the sample.

● The authors' response

They treat it as the next research problem rather than a defect: build category-specific rankers — a 'sad ranker', an 'angry ranker' — and add podcasts likely to contain the missing classes.

A useful general lesson: 'balanced' is never a property of a dataset on its own. It is a property of a dataset with respect to a particular label space, and balancing one can leave another as skewed as before.

Critical appraisal



What the design buys, and what it costs

● Strengths

Genuinely spontaneous speech, from many speakers, in many channel conditions — none of it acted.

A licence position almost nobody else has: the corpus can be redistributed and used commercially.

The pipeline is automatic apart from one screening step, so it scales. The corpus has grown roughly two hundred-fold since this paper.

Attribute and categorical descriptors together, with primary and secondary categorical labels, and at least five raters per segment.

A random control that makes the central claim testable rather than merely plausible.

The annotation quality-control protocol is reusable by anyone, in any domain.

● Limitations

Retrieval bias: the models were trained on twenty-two actors, so the corpus over-represents podcast speech that sounds like acted emotional speech. Emotional behaviour those models cannot detect will never be annotated.

Filtered naturalness: clean, single-speaker, wideband, music-free, 2.75 to 11 seconds. Real conversation is noisier and overlaps constantly.

Category imbalance persists, and fear and sadness remain scarce.

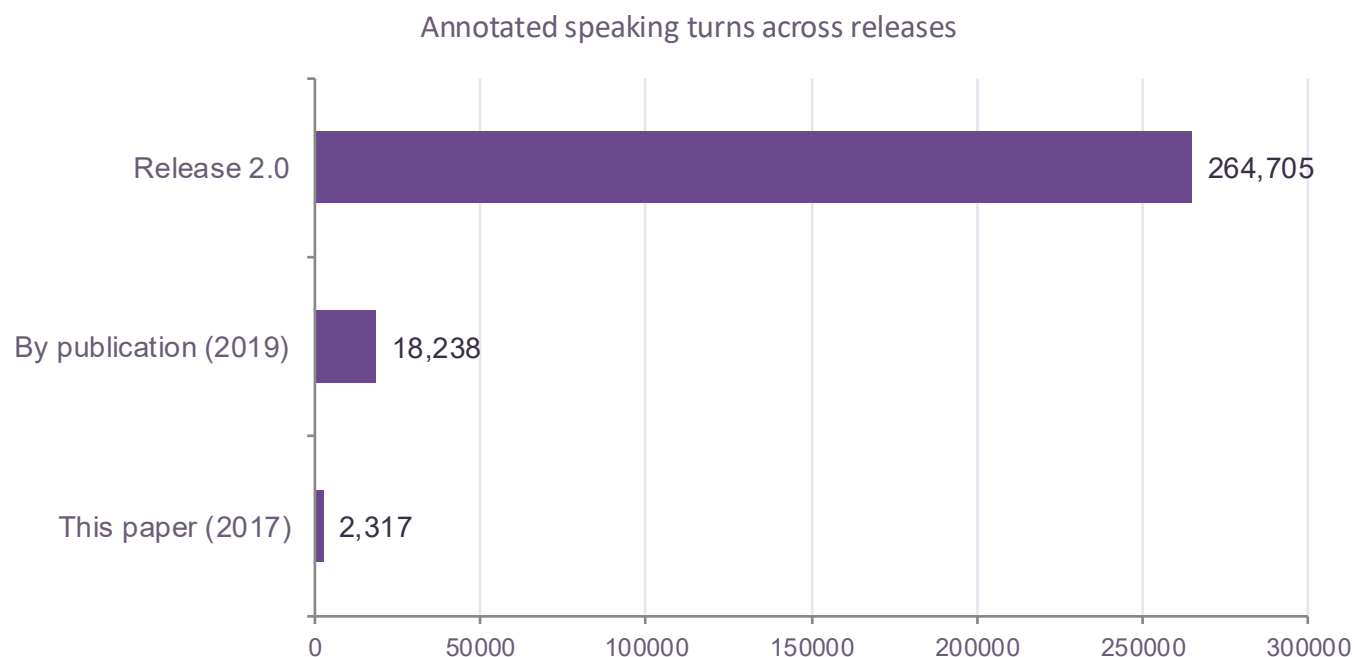
Podcasts are a genre, not the world. English-only, self-selected speakers, broadly performative public speech.

Speaker identities were unavailable at collection and had to be annotated afterwards, so speaker-independent splits were not possible from the start.

What it became



The protocol validated by what it produced



- By the time this paper appeared: 920 podcasts, 18,238 sentences, 27 hours 42 minutes, speaker identities annotated for 9,670 sentences across 151 speakers.
- Release 2.0: 264,705 turns and 409 hours, with defined speaker-independent train, development and three test partitions.
- One test set is drawn at random rather than by retrieval, so the retrieval bias can be measured directly.
- A third test partition withholds its labels entirely and is scored through a web interface — a corpus that ships with a leaderboard.

Note what the later releases fixed: defined speaker-independent splits, a non-retrieved control partition, and held-out labels. Every one of those is a direct answer to a problem we identified in IEMOCAP and CREMA-D — this is the field learning from its own corpus papers.

Questions for discussion



- 1 The retrieval models were trained on acted speech from twenty-two actors. How much does that shape what ends up in a corpus advertised as naturalistic — and how would you measure it without simply annotating everything?
- 2 Attribute balance was achieved; category balance was not. Is 'emotionally balanced' a meaningful claim at all without naming the label space it refers to?
- 3 The quality-control protocol prevents bad annotations; CREMA-D's redundancy averages them out. On a fixed budget, which do you buy?
- 4 All that monitoring still yields only moderate agreement. At what point do we stop treating rater disagreement as noise to be reduced and start treating it as the signal?
- 5 Podcasts are public, performative, self-selected speech from people who chose to publish. What kinds of emotional behaviour can this source never contain?