

Appendix

Supplementary Table 1. *Overview of action units and emotions.*

Action unit / emotion	Corresponding action unit number
Inner brow raiser	1
Outer brow raiser	2
Brow lowerer	4
Upper lid raiser	5
Cheek raiser	6
Lid tightner	7
Nosewrinkler	9
Upper lid raiser	10
Lip corner puller	12
Dimpler	14
Lip corner depresser	15
Chin raiser	17
Lip stretcher	20
Lip tightner	23
Lips part	25
Jaw drop	26
Blink	45
Happiness	6, 12
Sadness	1, 4, 15
Surprise	1, 2, 5, 26
Fear	1, 2, 4, 5, 7, 20, 26

Anger	4, 5, 7, 23
Disgust	9, 15
Contempt	12 ¹ , 14
Interest	1, 2, 5, 7, r20 ² , r26
Confusion ⁴	4, 7, 15, 17, 23
Valence 1 ⁵	r1, r4, r5, 6, r10, r17, 12, 25, r45
Valence 2 ⁵	6, 12
Arousal	Mean all AU with r45 ⁶
Blinks	45 ⁷

1. For contempt AU12 would only be activated on one side of the face. Dissociating left from right AU12 activations was, however, not possible due to technical limitations.

2. The letter r in front of an AU number indicates reversion (i.e., a higher activity corresponds to weaker corresponding emotion).

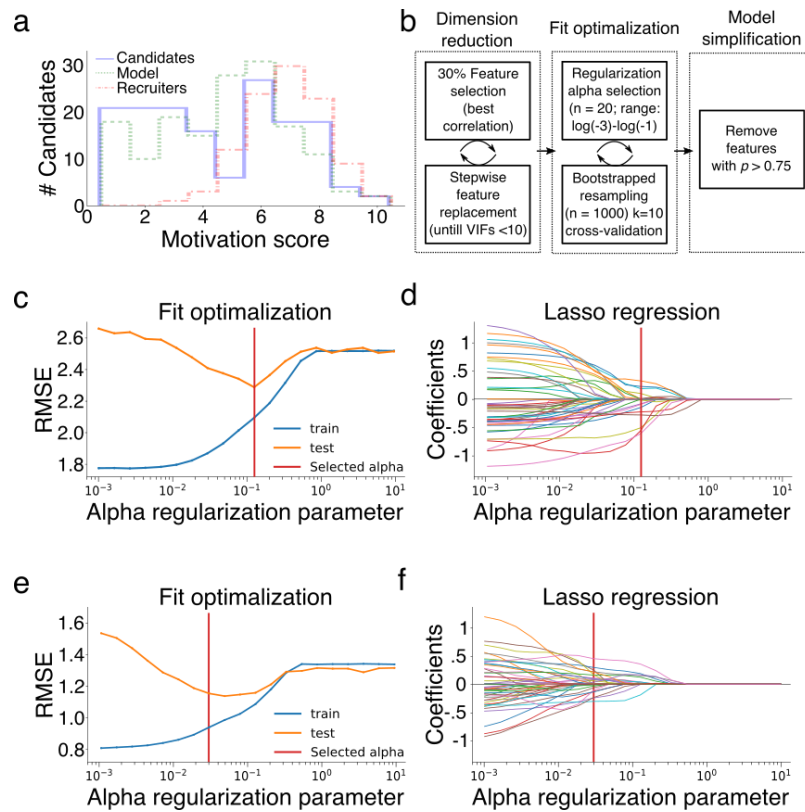
3. Interest is only measured at time points when a threshold activation of ≥ 0.4 for AU1 or AU2 was met.

4. Confusion threshold activation was set at ≥ 0.4 for AU4 and was met.

5. Two versions of valence were computed. The first was based on a pattern developed by McDuff and colleagues⁵¹, and a second by Sayette and colleagues⁵²

6. Action unit 45 was reversed because closing the eyes indicates a state of relaxation rather than arousal.

7. Blink threshold activation was set at ≥ 1 for AU45.



Supplementary Figure S1. Motivation rating distributions, modelling steps, loss function and beta weights across regularization parameters. **a**, Histogram of motivation scores of participants themselves (blue), objective model (green), and raters (red). **b**, Schematic drawing of modelling steps. **c**, Loss function showing the root mean square error (y-axis) for training set and test set as a function of the Lasso's alpha (lambda) regularization parameter (x-axis) for the model predicting candidates' self-reported motivation ratings. **d**, Coefficients (y-axis) of each feature as a function of alpha regularization parameter (x-axis). **e-f**, same as panel **c-d**, but now for the model predicting recruiters' motivation ratings. **d,f**, beta weights are strongly reduced at high alpha's to suppress features that have high collinearity with other features. The vertical red line indicates the selected alpha for the final, well-fitted model; this alpha selection results in a model for which the training and test RMSE were relatively similar, test errors were lowest, and weights of several features were decreased, all to prevent under- and over-fitting.