

Brain neuromarkers predict self- and other-related mentalizing across adult, clinical, and developmental samples

Received: 20 March 2025

Accepted: 20 May 2026

Published online: 06 June 2026

 Check for updates

Dorukhan Açıl ^{1,2,3}, Jessica R. Andrews-Hanna ^{4,5}, Marina López-Solà^{6,7}, Mariët van Buuren⁸, Lydia Krabbendam⁸, Liwen Zhang ⁹, Lisette van der Meer^{10,11}, Paola Fuentes-Claramonte ^{12,13}, Edith Pomarol-Clotet^{12,13}, Raymond Salvador^{12,13}, Martin Debbané ^{14,15}, Pascal Vrticka ¹⁶, Patrik Vuilleumier ^{17,18}, David A. Sbarra⁴, Andrea M. Coppola ⁴, Anita Tusche ^{19,20}, Lars O. White³, Tor D. Wager ²¹ & Leonie Koban ^{22,23} ✉

Human social interactions rely on the ability to reflect on one's own and others' internal states and traits—a process known as mentalizing. Impaired or altered mentalizing is a hallmark of multiple psychiatric and neurodevelopmental conditions. Yet, replicable and easily testable brain markers of mentalizing have so far been lacking. Here, we apply an interpretable machine learning approach to multiple datasets (total $n = 390$) to train and validate fMRI brain signatures that predict i) mentalizing about the self, ii) mentalizing about another person, and iii) both types of mentalizing. Self-mentalizing and other-mentalizing classifiers had positive weights in anterior/medial and posterior/lateral brain areas, respectively, with accuracy rates of 82% and 77% for out-of-sample prediction. The classifier trained across both types of mentalizing showed 98% predictive accuracy and separated (mental) attributional from factual inferences. Classifier patterns revealed better self/other separation in healthy adults compared to individuals with schizophrenia and with increasing age in adolescence. Together, our findings reveal consistent and separable neural patterns subserving trait-based mentalizing about self and others—present at least from the age of adolescence and functionally altered in severe neuropsychiatric disorders. These mentalizing signatures hold promise as candidate neuromarkers of social-cognitive processes in different contexts and clinical conditions.

Mentalizing—representing and inferring the psychological states of oneself and others—is a fundamental process for adaptive navigation through the social world^{1,2}. Delineating the brain systems involved in mentalizing about self and others is important for understanding brain health and dysfunction, as atypical mentalizing is characteristic for many neurodevelopmental and psychiatric conditions^{3–9}. Many studies

have examined the neural correlates of mentalizing, suggesting an interplay of different brain regions, including medial prefrontal cortex (mPFC), temporoparietal junction (TPJ), and precuneus^{10–14}. However, predictive brain measures of mentalizing that can be easily applied to new data to decode the degree of self-related and other-related processing are still lacking. Most cognitive and affective processes cannot

A full list of affiliations appears at the end of the paper. ✉ e-mail: leonie.koban@cnrs.fr

be captured by activity in individual brain regions, as they are reflected in patterns of brain activity distributed across multiple regions and systems, which can be harnessed by pattern-based decoding^{15,16}. For instance, recent work demonstrates that distributed brain activity and connectivity patterns, as indexed by fMRI, enable us to decode the intensity of pain¹⁷, drug and food craving¹⁸, sustained attention¹⁹, depressive rumination²⁰, and clinically relevant behaviors and outcomes^{21–23}. As opposed to the standard brain mapping approach that maps task features onto brain activity, these predictive brain activity patterns or brain signatures are multivariate models that utilize data across the whole brain to make formally testable population-level predictions across subjects and datasets^{15,24}. The predictions address the involvement and/or the intensity of a mental process. Here, we apply this brain signature approach to predict mentalizing about oneself or other people.

Mentalizing, as a hidden state, would arguably benefit from such an approach. Thinking about self and others is inherently multilayered and multidimensional^{12,25,26}. Recent conceptualizations of self-related and social cognition point to multiple dimensions, such as about oneself versus others, affective versus cognitive^{6,12,27}, and multiple layers, such as observing behaviors, inferring internal states, and representing longer-lasting traits²⁵. Moreover, social cognition terms, including but not limited to mentalizing (e.g., empathy, perspective taking), suffer from inconsistent usage in the literature and a lack of consensual definitions²⁸. Developing brain signatures of mentalizing could potentially help account for this heterogeneity, offering the possibility of testing whether the proposed subdimensions of mentalizing are subserved by dissociable neurobiological patterns. In turn, these signatures carry the potential for validating the multiple facets of mentalizing¹⁵. However, it remains unclear—especially considering the complexity of the mentalizing construct—whether distributed neural patterns can reliably predict mentalizing using brain images.

A central question is whether mentalizing about oneself and mentalizing about others (self- and other-mentalizing hereafter) are reliably distinguishable based on brain activity. Recent electrophysiological evidence suggests that self- and other-mentalizing activate overlapping cortical areas following a similar temporal sequence: TPJ, medial temporal gyrus (MTG)/temporal poles (TP), precuneus/posterior cingulate cortex (PCC), mPFC, in a roughly posterior to anterior temporal order^{29,30}. Conversely, differences in brain activity between self- and other-mentalizing have also been observed. In the mPFC, more dorsal areas coincide with other-related processing and ventral areas with self-related processing, with research initially supporting a linear³¹ but more recently a curvilinear³² ventral-to-dorsal gradient for self versus other within the mPFC. Further, self-referential thought typically elicits activation in cortical midline structures, such as anterior cingulate cortex (ACC), subcortical areas, including thalamus, striatum and caudate nucleus, and, to a lesser extent, in insula, temporal poles, and ventrolateral prefrontal cortex^{31–36} (vlPFC). In contrast, other-referential thought typically elicits activation in TPJ, middle and superior temporal gyri extending to the temporal poles^{10,11,32,37–39}, and to a lesser degree in supplementary motor area (SMA), left inferior and medial frontal gyri and medial orbitofrontal cortex^{34,40}. Collectively, the literature remains inconsistent regarding the separation between self- and other-mentalizing, and it is unclear whether generalizable brain models of self- versus other-mentalizing can be identified.

To address these gaps, we leverage fMRI and machine learning to develop four distinct brain signatures, (1) the Mentalizing Signature (MS) for mentalizing overall (i.e., thinking about either the self or another person versus non-mentalizing control conditions), (2) the Self-Referential Signature (Self-RS) to detect specifically self-related thought (here referred to as self-mentalizing), (3) the Other-Referential Signature (Other-RS) to detect other-related thought (here referred to as other-mentalizing), and (4) the Self-vs-Other Referential Signature (SvO-RS) to specifically separate self- and other-related mentalizing.

While the overall Mentalizing Signature aimed at developing a broad marker of social reflection that establishes the benchmark ceiling performance and may generalize across a wide range of social processes, the other three signatures (Self-RS, Other-RS, and SvO-RS) allowed us to test whether dissociable and generalizable neural patterns underlie distinct targets of mentalizing.

To this end, we used data from eleven cohorts and seven independent fMRI studies. The training and testing datasets used variants of a standard trait-evaluation task with three conditions, which involved reflecting on personality traits/statements describing oneself or another person. Conceivably, personality traits represent the enduring mental states that can be inferred from information collected by social perceptual systems⁴¹ across multiple observations^{12,25}. Trait representations are used to predict others' future behavior and become the basis of the mental-self²⁶. Moreover, trait evaluations are one of the few mentalizing tasks used in the literature to include a self-condition, as other commonly used tasks (e.g., mental attribution, false-belief, perspective-taking tasks) only assess mentalizing about others. Thus, trait evaluation is an ideally suited task to study a key mentalizing process, namely representing stable psychological states of self and others.

We first used standard machine learning algorithms—support vector machines—to develop and cross-validate multivariate classifiers (signatures) of mentalizing about traits in a sample of healthy adults. In a second step, we further validated these signatures in seven independent cohorts from different laboratories, countries, and scanners, and with different sample characteristics (healthy, adolescent, and clinical populations), allowing us to test their generalizability and predictive validity in various contexts. Third, we tested the signatures in two additional independent datasets that used different social cognition tasks to see whether mentalizing signatures would generalize to other social contexts in which mentalizing is likely to occur. Finally, we assessed whether local patterns of brain activity in several regions of interest (ROIs) contain sufficient information to predict self- and other-referential mentalizing. Together, the results of these analyses inform us about the functional neural organization of self- versus other-related mentalizing and provide us with distributed brain signatures that can serve as brain targets for measuring and intervening on mentalizing-related brain processes in future studies.

Results

Data overview

The study included a total of 1161 contrast images from 390 participants and 11 cohorts from seven independent studies, including six samples of healthy adults ($n=227$), two samples of healthy adolescents ($n=105$; $M_{\text{age}}=12.9$ and 16), and three samples of adults with a clinical diagnosis of either schizophrenia ($n=40$) or bipolar disorder ($n=18$; see Fig. 1A and Supplementary Table 1). The data from these seven studies were used as the training dataset, testing datasets, or the extension datasets (see Fig. 1A). The training and testing datasets included participants completing variants of a self- and other-referential judgement task. The two extension datasets included participants performing a social feedback task and an inference task. Studies 1–5 and 7 included three conditions, namely a self-condition, an other-condition, and a nonsocial control condition, whereas Study 6 included two conditions in one study (attributional versus factual inferences; Study 6a) and a 2×2 design in the second (Study 6b; attributional versus factual inferences $\times$ social versus nonsocial scenes). Contrast images were computed for each condition (versus implicit baseline) and rescaled using L2-norm to standardize the scale of beta (regression coefficient) weights across participants, studies, and scanners.

Cross-validated training results

The training dataset (Study 1a) consisted of $n=21$ adult participants who completed a trait-evaluation task using a block design (similar to

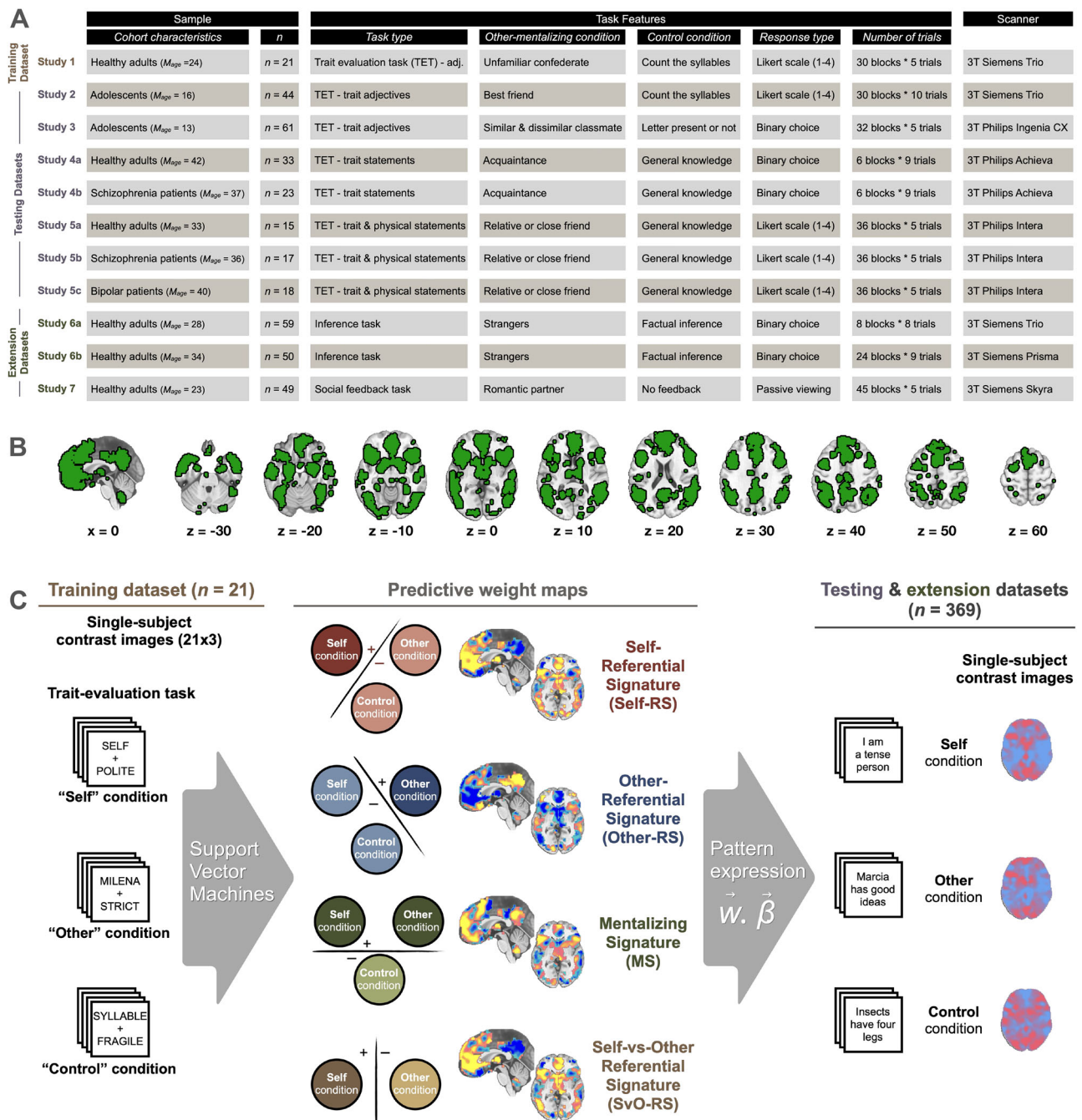


Fig. 1 | Study design and analytic approach. **A** The present study included data from seven independent studies (eleven cohorts, total $n = 390$). Training ($n = 21$) and testing ($n = 211$) datasets included a trait-evaluation task with three conditions: a self-condition, an other-condition, and a nonsocial control condition. Extension datasets ($n = 158$) constituted the novel instances to test signatures in distinct tasks. **B** Display of the inclusive mask of brain areas related to social cognition. The single-subject contrast images included in this study are masked using this mask. For results with a whole-brain grey matter mask, please refer to Supplementary Fig. 1. **C** The training dataset (Study 1) included contrast images of $n = 21$ participants performing a trait-evaluation task with three conditions (Self, Other, Control). We

used a 10-fold cross-validated linear SVM to train four mentalizing signatures. The Self-Referential Signature (Self-RS) was trained to predict the Self condition versus the two remaining conditions. The Other-Referential Signature (Other-RS) was trained to predict the Other condition versus the two remaining conditions. The Mentalizing Signature (MS) was trained to separate the Self and Other conditions from the Control condition. The Self-vs-Other Referential Signature (SvO-RS) was trained to separate the Self and Other conditions. For the validation in independent datasets, we applied the signatures to the single-subject contrast images from Studies 2–7 by computing the dot product between the weight maps and the contrast images.

ref. 42; see Fig. 1C). In each block, participants were presented with several trait adjectives: In the Self condition, they were asked to rate the degree to which each trait adjective described themselves. In the Other condition, they had to rate how much each adjective described another person—a confederate with whom the participant had previously interacted during a decision-making task⁴³. In the non-

mentalizing Control condition, participants indicated the number of syllables in the trait adjectives. To reduce the influence of non-mentalizing-related brain regions (e.g., visual cortex) and the risk of opportunistic classification based on features not related to mentalizing, we used an inclusive mask of brain regions related to social processing (see Fig. 1B).

We used support vector machines (SVM) with default parameters (to avoid overfitting) and 10-fold cross-validation⁴⁴ to train four distinct mentalizing signatures using the masked contrast images from the training dataset. The folds were defined at the subject level, so that all contrast images from a given subject were assigned to the same fold. The Self-Referential Signature (Self-RS) was trained to detect self-reflection (versus the two remaining conditions), the Other-Referential Signature (Other-RS) to detect other-reflection (versus the two remaining conditions), the Mentalizing Signature (MS) to detect both types of mentalizing versus the Control condition, and the Self-vs-Other referential Signature (SvO-RS) to differentiate self- versus other-related processing (see Fig. 1C).

All four signatures showed excellent cross-validated (out-of-sample) prediction accuracy (100% accuracy in two-alternative forced-choice tests, $p < 0.001$, Cohen's d for the Self-RS = 3.18, the Other-RS = 2.45, the MS = 4.92; the SvO-RS = 2.45; AUC = 1.00 for all signatures). To identify which voxels contributed most reliably to the mentalizing signatures, we used bootstrapping (5000 samples) to obtain one p -value per voxel and displayed the thresholded weight maps using false discovery rate (FDR) correction at $q < 0.05$ and cluster extent $k > 10$ voxels. Brain regions with significant positive voxel weights for the Self-RS included the ventromedial (vmPFC) and dorsomedial (dmPFC) prefrontal cortex, frontal eye fields, ventral ACC, frontal operculum, anterior insula, thalamus, and caudate nucleus (see Fig. 2 and Supplementary Table 2). For the Other-RS, significant positive weights were found in left vlPFC, left superior temporal sulcus (STS), bilateral TPJ, and precuneus/PCC (see Fig. 2 and Supplementary Table 3). For the SvO-RS (see Fig. 3 and Supplementary Table 5), brain regions with significant positive voxel weights (thus, associated more with the Self than the Other condition) included vmPFC/ACC, left anterior insula, left inferior frontal sulcus, left caudate nucleus, left middle temporal cortex, and bilateral thalamus. Significant negative clusters in the SvO-RS (thus, associated more with the Other than the Self condition) were found in precuneus, bilateral TPJ/inferior parietal cortex, right PCC, left STS, and left inferior frontal gyrus. The MS had significant positive clusters in mPFC (both dorsal and medial), bilateral vlPFC, dorsal ACC (dACC), frontal operculum, bilateral SMA, bilateral MTG, bilateral TP, left STS, middle cingulate gyrus (MCC), bilateral PCC, precuneus, bilateral angular gyrus/TPJ, and subcortical areas, including right caudate nucleus, left anterior insula, and bilateral thalamus (see Fig. 2 and Supplementary Table 4).

Validation in independent samples

Next, we validated the four brain signatures on several completely independent studies: two samples of healthy adults (Studies 4a and 5a), two adolescent samples (Studies 2 and 3), one cohort of participants with bipolar disorder (Study 5c), and two cohorts of participants with schizophrenia (Studies 4b and 5b). All datasets included comparable trait evaluation tasks and fMRI block designs with three conditions, with some small variations in task designs between studies (see Fig. 1A). To obtain pattern expression values, we computed the matrix dot product between the mentalizing signatures and each subject-level contrast image from these datasets, yielding one scalar value per individual contrast image and per signature (see Fig. 1C). These pattern expression values were then used to test the predictions of the mentalizing signatures. The weighted average prediction accuracy in two-alternative forced-choice (2AFC) tests, across all independent testing datasets, was 81.52% for the Self-RS ($\pm 6.17\%$ average STE; significant in 10 out of 14 validation tests [7 samples $\times$ 2 comparisons]), 77.25% for the Other-RS ($\pm 6.02\%$ average STE; significant in 12 out of 14 tests), 97.87% for the MS ($\pm 1.79\%$ average STE; significant in all 14 tests), and 77.73% for the SvO-RS ($\pm 7.14\%$ average STE; significant in 6 out of 7 tests), suggesting overall high prediction accuracy of the four signatures, even in new samples, although with some variations between datasets (see Figs. 2, 3, and Supplementary Table 6). There were no sex

effects on pattern expression or signature fit—defined as the difference between the pattern expression for the true condition and the false condition in each signature (see “Methods”).

Lower self/other separation in participants with schizophrenia

The results above show that the signatures significantly predicted mentalizing in most samples, including clinical samples. Yet, schizophrenia, in particular, is often associated with impaired social cognition and altered self-perception^{45,46}. Thus, we next tested how well the signatures separated self- from other-related mentalizing in two of the testing datasets (Study 4 and Study 5) that included both participants with schizophrenia (total $n = 40$) and matched healthy participants (total $n = 48$). As expected, the three signatures, the Self-RS ($M_{HC} = 0.32$, $M_{SCZ} = 0.12$; $\beta = 0.21$, STE = 0.07, CI = [0.07, 0.35], $t(86) = 2.99$, $p = 0.004$), the Other-RS ($M_{HC} = 0.43$, $M_{SCZ} = 0.24$; $\beta = 0.19$, STE = 0.07, CI = [0.04, 0.33], $t(86) = 2.57$, $p = 0.012$), and the SvO-RS ($M_{HC} = 0.39$, $M_{SCZ} = 0.19$; $\beta = 0.20$, STE = 0.07, CI = [0.07, 0.34], $t(86) = 2.97$, $p = 0.004$) showed better discrimination between self- and other-referential mentalizing in healthy adults, compared to adults with schizophrenia (i.e., a greater positive difference between the Self-RS responses in the Self compared to the Other condition, and a greater difference between the Other-RS responses for the Other compared to the Self condition). No group differences were found in the correct discrimination of mentalizing versus Control by the MS ($p = 0.29$, see Fig. 4). These results indicate that, compared to healthy adults, participants with schizophrenia have less differentiated brain patterns between self- and other-related thought, indicating the potential clinical utility of the signatures.

For completeness, we also compared the discrimination of self- versus other-related mentalizing activity in participants with bipolar disorder ($n = 18$) versus healthy adults ($n = 15$) using data from Study 5. However, participants with bipolar disorder did not differ significantly from controls for any of the signatures.

Better self/other separation with increasing age in adolescents

Next, we explored potential developmental differences in the classifier performances by testing whether the ability of the classifiers to separate self- from other-related mentalizing depended on the age of the participants in the two adolescent samples ($n = 105$). We combined data from Study 2 (age 12–18 years; $M_{age} = 16$) and Study 3 (age 11–14 years; $M_{age} = 12.9$). We found that (controlling for sex and the source study) older adolescents had better self/other separation for the three signatures trained to distinguish these two conditions: the Self-RS ($\beta = 0.08$, STE = 0.03, CI = [0.018, 0.14], $t(102) = 2.55$, $p = 0.012$), the Other-RS ($\beta = 0.06$, STE = 0.03, CI = [0.001, 0.12], $t(102) = 2.03$, $p = 0.045$), and the SvO-RS ($\beta = 0.07$, STE = 0.03, CI = [0.01, 0.13], $t(102) = 2.31$, $p = 0.023$; see Fig. 5). In contrast, there were no significant associations between age and performance of the MS. These results suggest that, with increasing age, adolescents' brain activity becomes more differentiated for self- versus other-related mentalizing.

Testing the signatures during attributional versus factual inference

So far, we have shown that the signatures performed well in several independent datasets from different labs, but all using a similar explicit mentalizing task. In order to test whether these signatures are implicated specifically in mentalizing or also in broader social or even nonsocial reflections, we tested them in an additional dataset (Study 6a and 6b) of $n = 109$ healthy adults making attributional versus factual inferences about social and nonsocial stimuli (see Fig. 6A for the task). Attributional inferences (e.g., “Is Ada proud of herself?”), especially in a social context, require mentalizing, or in other words, they require representing hidden mental states. Factual inferences (e.g., “Is Ada tall?”), on the other hand, refer to directly observable features and should require less mentalizing. Therefore, we tested whether the

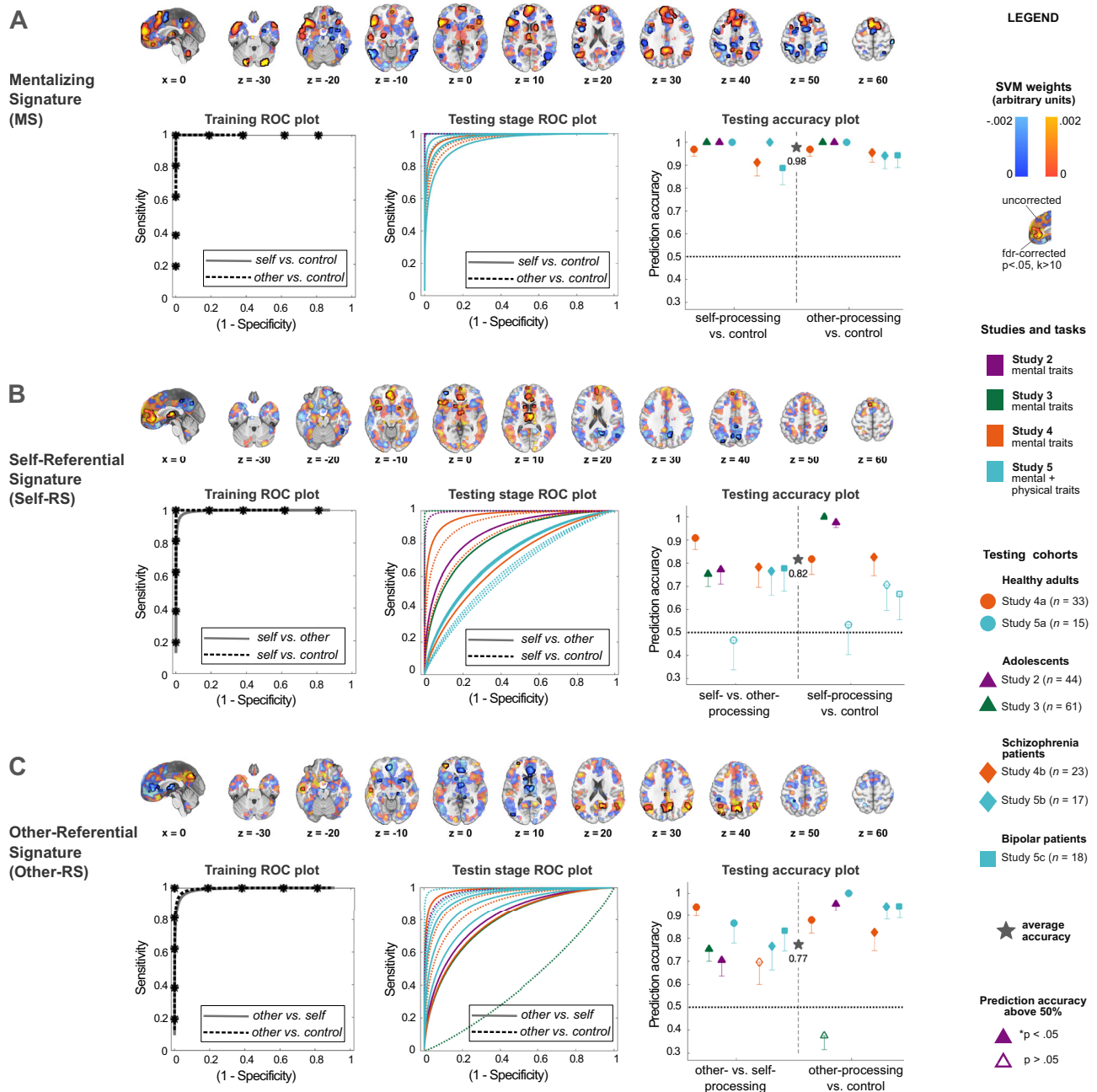


Fig. 2 | Model weights, cross-validated training, and validation results of the mentalizing signatures. SVM classifiers were trained using 5000 bootstrap samples derived from $n = 21$ subjects, yielding statistical maps with two-tailed, uncorrected p -values. On the top row of each box, the weight maps illustrate the uncorrected positive and negative weights for **A** the Mentalizing Signature (MS), **B** the Self-Referential Signature (Self-RS), and **C** the Other-Referential Signature (Other-RS). Voxels significant at an FDR-corrected threshold ($q < 0.05$, minimum cluster size of $k = 10$) are highlighted by black outlines and brighter colors. The left and middle plots in each box depict the receiver operating characteristic (ROC) plots from the cross-validated training and testing datasets. The last column shows the two-alternative forced-choice prediction accuracies of the signatures in each

testing dataset separately (total $n = 211$), with color representing the study and shape representing the cohort type. Filled shapes indicate classification accuracies significantly above chance level (50%) according to two-sided binomial tests at $p < 0.05$; unfilled shapes indicate nonsignificant accuracies. The one-sided error bars below the shapes illustrate the standard error around the accuracy estimates. The star shows the weighted average accuracy across all seven datasets. To ensure that our cross-validated training and validation results were not influenced by the social-processing mask that we used, we repeated our analyses with whole-brain (grey-matter masked) images (see Supplementary Fig. 1). To test whether a different feature selection method would yield similar results, we trained classifiers using recursive feature elimination (RFE), which we report in Supplementary Fig. 2.

Other-RS and the MS could predict attributional against factual inference conditions successfully.

Participants in Study 6a ($n = 59$) were asked to make these two types of inferences on social images showing faces or hands with a question prompt (e.g., Is the person proud of themselves?). Attributional (versus factual) inference conditions were predicted with 98% accuracy (2AFC) by both the Other-RS and the MS (see Fig. 6C;

$p < 0.001$, $AUC = 1.00$, sensitivity = 98%; specificity = 98%, Cohen's $d = 2.62$ for the Other-RS and 3.94 for the MS). Study 6b had a 2 (attributional versus factual) $\times$ 2 (social versus nonsocial) design. The Other-RS predicted the attributional inferences against factual inferences with 82% accuracy in the social condition ($p < 0.001$, Cohen's $d = 0.95$, $AUC = 0.85$, sensitivity = 82%, specificity = 82%) and 66% accuracy in the nonsocial condition ($p = 0.03$, Cohen's $d = 0.61$,

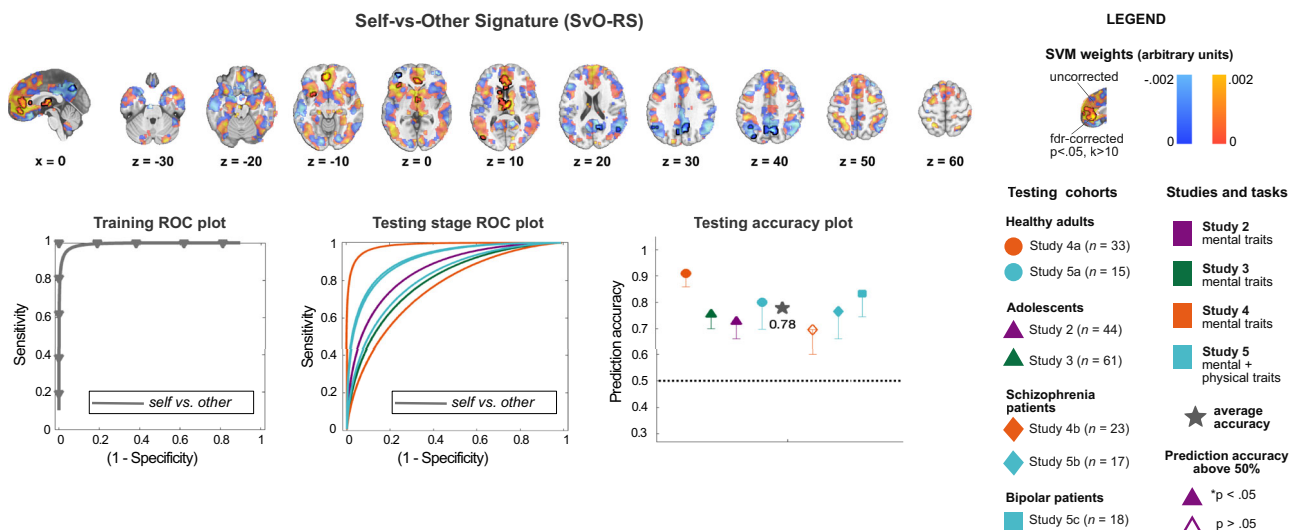


Fig. 3 | Model weights, cross-validated training, and validation results of the Self-vs-Other Signature (SvO-RS). The SVM classifier was trained using 5000 bootstrap samples derived from $n = 21$ subjects, yielding a statistical map with two-tailed, uncorrected p -values. The weight map (first row) illustrates the uncorrected positive and negative weights, as well as the voxels significant at an FDR-corrected threshold ($q < 0.05$, minimum cluster size of $k = 10$). The second row depicts the receiver operating characteristic (ROC) plots from the cross-validated training and testing stages, and the two-alternative forced-choice accuracies of the SvO-RS in

each testing dataset (total $n = 211$). Filled shapes indicate classification accuracies significantly above chance level (50%) according to two-sided binomial tests at $p < 0.05$; unfilled shapes indicate nonsignificant accuracies. The one-sided error bars below the shapes illustrate the standard error around the accuracy estimates. The cohort types are shape-coded, and the studies are color-coded in the validation accuracy plot. The star shows the weighted average accuracy across all seven datasets.

AUC = 0.73, sensitivity = 66%, specificity = 66%; see Fig. 6D). The MS predicted the attributional versus factual inference with 100% accuracy in the social condition ($p < 0.001$, Cohen's $d = 3.56$, AUC = 1.00, sensitivity = 100%, specificity = 100%) and 88% accuracy in the non-social condition ($p < 0.001$, Cohen's $d = 1.21$, AUC = 0.92, sensitivity = 88%, specificity = 88%; see Fig. 6D). Further, we tested how well the signatures separate social from nonsocial inferences (averaged across attributional and factual conditions). While the Other-RS did not significantly separate social from non-social inferences (56% accuracy, $p = 0.48$, Cohen's $d = 0.27$, AUC = 0.62, sensitivity = 56%, specificity = 56%), the MS separated social from nonsocial inferences with almost perfect accuracy (98% correct, $p < 0.001$, Cohen's $d = 2.96$, AUC = 1.00, sensitivity = 98%, specificity = 98%; see Fig. 6B). Together these results suggest that the mentalizing signatures are sensitive to processes that require a reflection about internal or hidden states, and—at least for the MS—especially for social agents.

Testing the signatures in a social feedback task

To examine how the signatures would respond to social interaction tasks in which mentalizing is not directly instructed but likely implicated, we next tested their performance in an fMRI dataset (Study 7) of $n = 49$ romantic partners performing a social feedback task. Participants were led to believe that the task was designed to understand how people rate other people and interpret others' ratings. Therefore, participants first indicated how likable they felt each of 20 individuals (confederates) was, who supposedly rated the participants and their partners in return. To set up a context facilitating spontaneous mentalizing, participants were told that their partners saw the ratings on the screen simultaneously, in a separate room. The task included conditions in which participants received feedback about their own likeability (self-feedback), the likeability of their partner (partner-feedback), as well as a no-feedback control condition. The control condition consisted of trials where no feedback was given, with a prior explanation that such trials signaled that the rater (the individual on the screen) was not shown the participants' picture or did not supply their rating in the time required by the rating task.

Here, we tested whether the signatures' responses paralleled the target of the feedback conditions. For instance, to provide evidence of successful extension, the Self-RS should show higher pattern expression values in the self-feedback condition as opposed to the other two conditions. As expected, the Self-RS responded more strongly to the self-condition, $F(2, 96) = 14.214$, $p < 0.001$, $\eta_p^2 = 0.23$ (Fig. 7). Bonferroni-corrected pairwise comparisons showed that the self-feedback condition ($M = 0.16$) had significantly higher Self-RS expression scores than both the partner-feedback ($M = -0.09$, $p = 0.01$) and control ($M = -0.36$, $p < 0.001$) conditions. Additionally, the Self-RS also produced higher pattern expression values for the partner-feedback ($M = -0.09$) compared to the control condition ($M = -0.36$, $p = 0.04$). Similarly, the MS successfully discriminated both feedback conditions against the control condition, $F(2, 96) = 14.711$, $p < 0.001$, $\eta_p^2 = 0.24$ (Fig. 7). Bonferroni-corrected pairwise comparisons indicated that both self-feedback ($M = 0.17$, $p < 0.001$) and partner-feedback ($M = 0.12$, $p = 0.002$) conditions had significantly higher pattern expression scores than the control condition ($M = -0.17$). However, the Other-RS did not produce any significant differences between task conditions, $F(2, 96) = 2.003$, $p = 0.14$ (Fig. 7). Together this suggests a modulation of the MS and the Self-RS (but not the Other-RS) even in a task where mentalizing was not directly instructed, in line with the idea that feedback to oneself or one's partner should lead to more mentalizing-related brain activity.

Local patterns of self- and other-related mentalizing

We trained local classifiers in regions of interest (ROI) associated with mentalizing to gain further insight into how mentalizing-related information is processed locally and which regions can predict self-versus other-related mentalizing. The whole brain patterns indicate which regions have the most reliable positive and negative contributions to different forms of mentalizing, but they do not necessarily show which regions are not involved and do not inform us whether local patterns alone can predict the target of mentalizing (Self or Other). Training and testing classifiers for brain regions implicated in the literature can inform us in this regard.

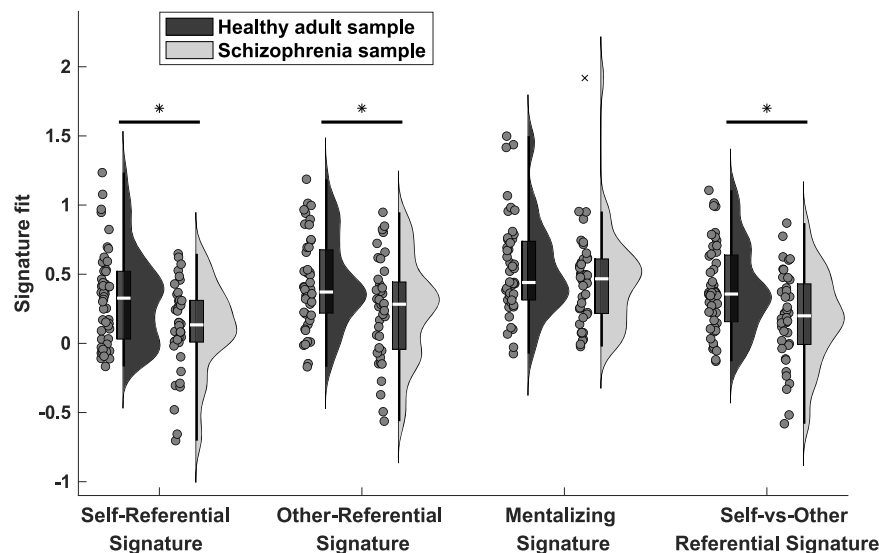


Fig. 4 | Self/other discrimination based on classifier responses in healthy adults compared to individuals with schizophrenia. Signature fits are measured as the pattern expression of true minus false class, separately for each signature in healthy adults ($n = 48$, in dark gray) and adults with schizophrenia ($n = 40$, in light gray). Linear mixed effects models were used to test the group differences (two-sided $p < 0.05$ threshold), controlling for the effects of different cohorts. Self/other discrimination was found to be better in healthy adults compared to patients with

schizophrenia for Self-RS ($\beta = 0.21$, $STE = 0.07$, $CI = [0.07, 0.35]$, $t(86) = 2.99$, $p = 0.004$), Other-RS ($\beta = 0.19$, $STE = 0.07$, $CI = [0.04, 0.33]$, $t(86) = 2.57$, $p = 0.01$), and SvO-RS ($\beta = 0.20$, $STE = 0.07$, $CI = [0.07, 0.34]$, $t(86) = 2.97$, $p = 0.004$). Dots indicate values per person, and boxplots mark the median, lower, and upper quartiles. Outliers (>3 SD from the mean) are marked with an “x.” Whiskers extend from the minimum to the maximum data points, excluding outliers. Asterisks indicate statistical significance: $*p < 0.05$.

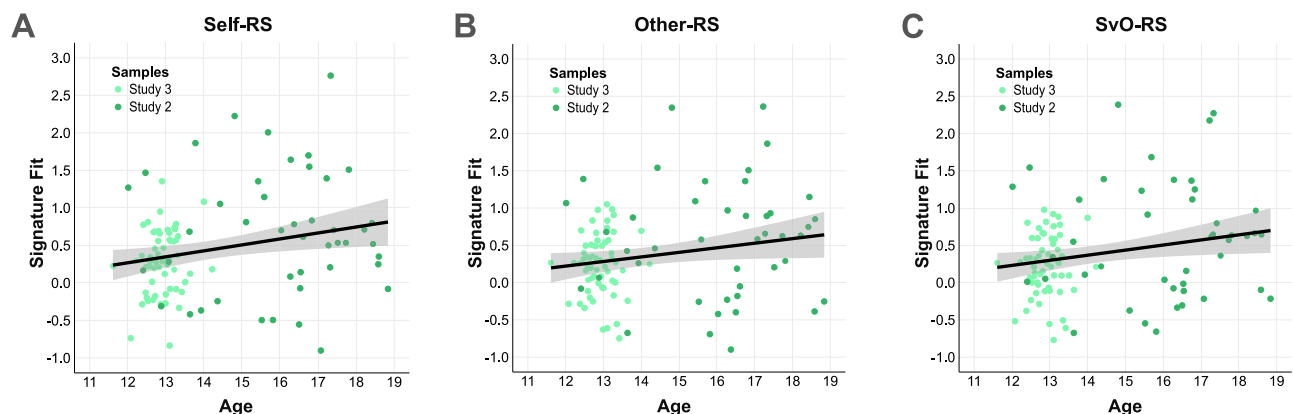


Fig. 5 | Association between self/other discrimination and age of adolescents ($n = 105$). Self/other discrimination (i.e., signature fit) was quantified as the pattern expression difference between the true class (e.g., self-condition) and the false class (e.g., other-condition), separately for each signature. Linear mixed effects models showed that, controlling for sex and cohort, older age was significantly associated with better differentiation of self- and other-related mentalizing for **A** Self-RS

($\beta = 0.08$, $STE = 0.03$, $CI = [0.02, 0.14]$, $t(102) = 2.55$, $p = 0.012$), **B** Other-RS ($\beta = 0.06$, $STE = 0.03$, $CI = [0.00, 0.12]$, $t(102) = 2.03$, $p = 0.045$), and **C** SvO-RS ($\beta = 0.07$, $STE = 0.03$, $CI = [0.01, 0.13]$, $t(102) = 2.31$, $p = 0.023$). Black lines display the fitted linear regression lines. Gray shaded bands indicate the 95% confidence intervals around the fitted regression lines. Dots show values per participant.

To this end, we included ten ROIs (see Fig. 8) previously associated with mentalizing as shown in an automated term-based meta-analysis (NeuroSynth⁴⁷): mPFC, two clusters in the anterior middle temporal gyrus (aMTG) bilaterally, two clusters in TPJ bilaterally, a cluster covering precuneus and posterior cingulate cortex (precuneus/PCC), left SMA, and three clusters in the cerebellum. We trained four types of classifiers in each of these ten ROIs. In other words, we trained each ROI to perform the four following classification tasks separately: (1) self-mentalizing (versus the two remaining conditions); (2) other-mentalizing (versus the two remaining conditions), (3) both mentalizing conditions (Self and Other) against the control condition, and (4) self-mentalizing against other-mentalizing. The ROI classifiers were trained and cross-validated in our training dataset (Study 1a; $n = 21$)

and tested in the combined sample ($n = 211$) of participants from all testing datasets (including all healthy, developmental, and clinical cohorts; for detailed results, see Supplementary Table 7).

As expected, all ten regions showed excellent classification accuracy for the general mentalizing classifiers in the training dataset, and all predicted mentalizing successfully in the testing dataset (see blue bars in Fig. 8). However, not all regions could significantly differentiate self-reflection and other-reflection conditions. While mPFC, aMTG, TPJ, and precuneus/PCC were capable of differentiating Self and Other conditions, the cerebellum clusters and left SMA failed at this task (see orange bars in Fig. 8). When trained for self-mentalizing (see red bars in Fig. 8), the right aMTG cluster could not capture any consistent configurations for self-mentalizing. While TPJ subregions

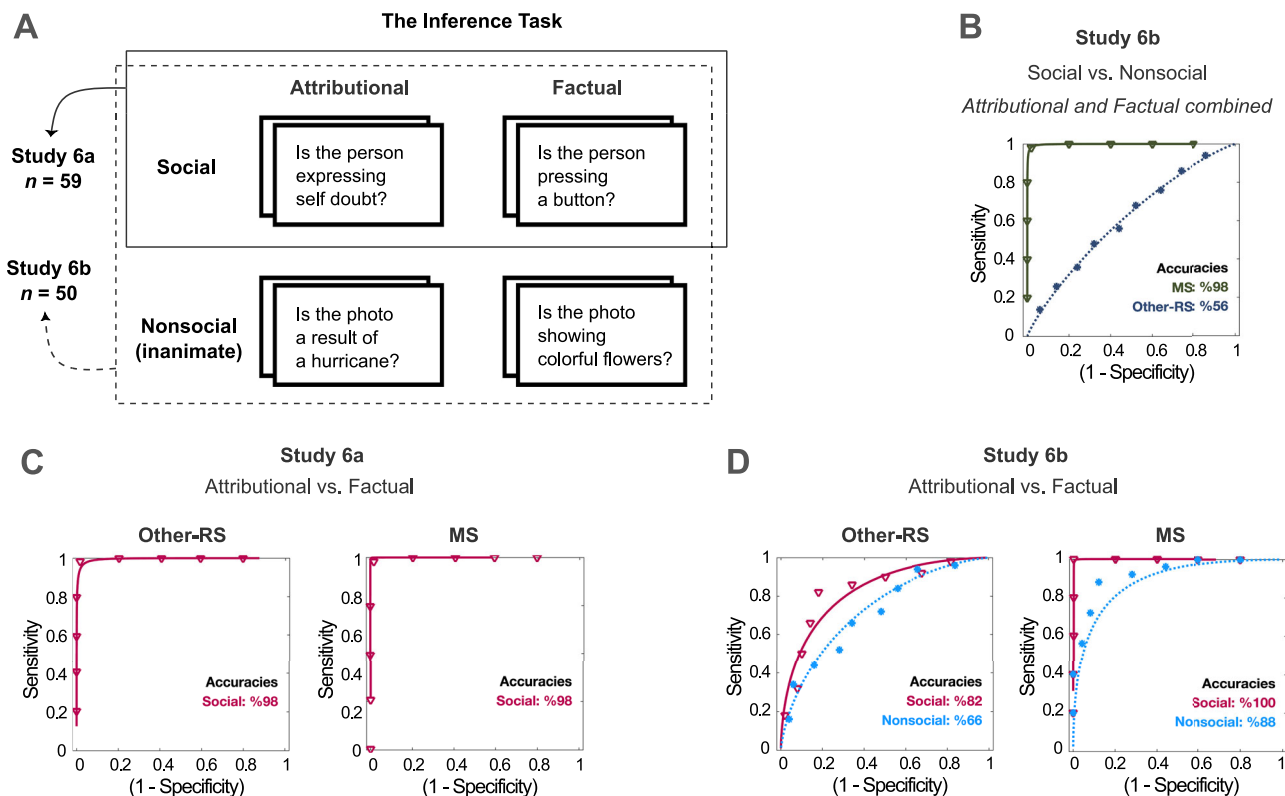


Fig. 6 | Extension of the Other-RS and the MS to an independent inference task (total $n = 109$). **A** Task design. Study 6a had a 1×2 design (attributional versus factual inferences regarding social scenes). Study 6b had a 2×2 design (attributional versus factual inferences $\times$ social versus nonsocial scenes). **B** Binary classification performances of the MS and the Other-RS in social versus nonsocial

attributional (Study 6b). **C** Binary classification results from Study 6a for separating attributional versus factual inferences. **D** Binary classification results from Study 6b, suggesting that the Other-RS and MS signatures are more sensitive to attributional (versus factual) inferences that require reflection about internal states—especially those concerning social agents.

were successful here against the control conditions, they could not predict self-against other-mentalizing, implying a lack of exclusive neural patterns for self-mentalizing within TPJ subregions. Finally, when trained for other-mentalizing, the cerebellum subregions and left SMA were incapable of picking up consistent patterns (see violet bars in Fig. 8).

Together, our analyses suggest that mPFC (including both vmPFC and dmPFC), TPJ, aMTG, and precuneus/PCC clusters encode both self- and other-mentalizing, and that the left SMA and cerebellum clusters most likely are involved in domain-general operations during mentalizing. Besides, our findings suggest that TPJ and aMTG are more strongly associated with other-mentalizing, while mPFC and precuneus/PCC clusters are unique in the sense that they are consistently involved in both self- and other-related mentalizing.

Leave-one-site-out classifiers trained on all datasets

Finally, to ensure that the high accuracies were not an artifact of features unique to the training study (e.g., scanner, sample), we trained new classifiers that combined data across all five training and testing datasets using a leave-one-site-out (LOSiO) cross-validation approach. We combined images from $n = 232$ participants from Studies 1 to 5 that used a trait-evaluation task, pooling data from healthy participants, adolescents, and clinical samples. We used the same parameters to train SVM models corresponding to the four mentalizing signatures.

Cross-validated prediction accuracies of all four classifiers were very good to excellent (see Fig. 9A–D): 83% averaged accuracy for the Self-RS_{LOSiO}, 91% for the Other-RS_{LOSiO}, 96% for MS_{LOSiO}, and 90% for the SvO-RS_{LOSiO} in 2AFC tests, $p < 0.001$ (averaged Cohen's d for the Self-RS_{LOSiO} = 1.10, for the Other-RS_{LOSiO} = 1.36, for the MS_{LOSiO} = 0.64, for the SvO-RS_{LOSiO} = 1.19; averaged AUC for the Self-RS_{LOSiO} = 0.90, for

the Other-RS_{LOSiO} = 0.95, for the MS_{LOSiO} = 0.93, for the SvO-RS_{LOSiO} = 0.93). Thus, these results show that the high prediction accuracies achieved by the mentalizing signatures were not inflated by task-specific idiosyncrasies but rather reflect consistent, robust multivariate patterns that are achieved using different combinations of the training dataset.

Finally, we applied the LOSiO classifiers to the inference task (Study 6) to test how well they classify attributional versus factual and social versus nonsocial conditions (see Fig. 9E). Both Other-RS_{LOSiO} and MS_{LOSiO} predicted social attributional (versus factual) statements in Study 6A with 97% accuracy ($p < 0.0001$; Cohen's d for Other-RS_{LOSiO} = 2.75 and for MS_{LOSiO} = 2.42). In Study 6b, the Other-RS_{LOSiO} separated attributional versus factual statements in the social condition with 84% accuracy ($p < 0.001$, Cohen's $d = 1.32$), whereas this classification was nonsignificant in the nonsocial condition (accuracy = 56%, $p = 0.48$). The MS_{LOSiO} predicted attributional versus factual statements in the social condition with 96% accuracy ($p < 0.001$, Cohen's $d = 2.16$) and in the nonsocial condition with 66% accuracy ($p = 0.033$, Cohen's $d = 0.73$). Finally, MS_{LOSiO} and Other-RS_{LOSiO} separated social from non-social inferences with 96% and 94% accuracy, respectively ($p < 0.001$, Cohen's d for MS_{LOSiO} = 2.12, and for Other-RS_{LOSiO} = 2.5). Together, this illustrates that, similar to the mentalizing signatures, the LOSiO classifiers are most sensitive to attributional processing in the social context, while also capturing nonsocial attributional processing to some degree.

Discussion

Mentalizing—reflecting about others' and one's own internal states—is a fundamental capacity required to function in a social world. Trans-diagnostically, mentalizing deficits typify an array of psychopathological

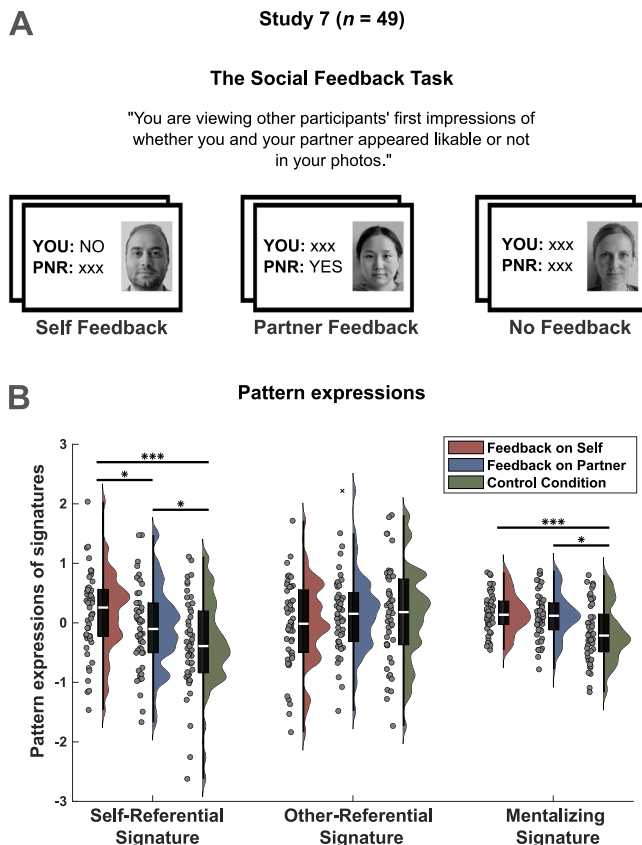


Fig. 7 | Extension of three signatures to an independent social feedback task ($n = 49$). **A** Participants received feedback for themselves, for their partners, or no feedback (control condition). **B** Signature responses for each of the three conditions, demonstrating higher Self-RS responses to feedback concerning oneself than the partner and no feedback, and higher mentalizing-related brain responses (MS) for self- and partner-directed compared to no feedback. Each dot represents one participant ($n = 49$; 3 repeated samples design with 3 conditions). Boxplots mark the median, lower, and upper quartiles. Outliers (>3 SD from the mean) are marked with an "x." Whiskers extend from the minimum to the maximum data points, excluding outliers. Asterisks indicate significant Bonferroni-corrected pairwise comparisons (t-test) following an initial repeated measures ANOVA test within each signature: $*p < 0.05$; $***p < 0.001$. The photographs shown in (A) are illustrative placeholders for display purposes only; the actual face stimuli used in the study were drawn from the Chicago Face Database⁹⁶.

conditions^{3–9}. Here, we developed novel brain signatures of self- and other-related mentalizing (Self-RS, Other-RS, and SvO-RS), as well as mentalizing in general (the Mentalizing Signature or MS), combining data from eleven independent cohorts ($n = 390$). The four signatures showed good to excellent classification accuracies in independent data and provided several insights into the brain circuits of trait-based mentalizing. First, mentalizing about traits appears to coincide with a specific pattern of brain activity, with reliably dissociable activation patterns for self- and other-related mentalizing, based on both whole-brain activity as well as within key nodes of the social brain, especially mPFC, precuneus/PCC, TPJ, and aMTG. Second, the signatures significantly predicted mentalizing not only in healthy adults, but also in adolescents and in individuals with schizophrenia and with bipolar disorder, demonstrating its utility across different (including clinical) samples. Third, our results point to the relevance of these mentalizing signatures for a range of different research questions in social, developmental, and clinical neuroscience, by showing (i) engagement of the signatures in independent social cognition tasks, (ii) that the differentiation of self- and other-related mentalizing increases in transition from adolescence to adulthood, and (iii) that this differentiation is less

pronounced in individuals with schizophrenia compared to healthy controls. Thus, our results provide us with clearly defined *brain models*¹⁵, applicable across contexts to assess the engagement of mentalizing-related brain regions in different experimental conditions and their alteration in clinical and neurodevelopmental conditions.

The Mentalizing Signature (MS) successfully predicted mentalizing across different contexts and extended to tasks without explicit mentalizing instructions (e.g., spontaneous mentalizing). It performed very well in separating higher-level inferences about others and hidden states from factual inferences, with higher classification accuracy in social compared to non-social contexts. The MS (but not the Other-RS) also significantly separated social from nonsocial inferences. Together, these findings demonstrate its sensitivity to mental state compared to external state attributions. The regions that contributed to the MS most positively included mPFC, precuneus, TPJ, STS, and many other regions previously associated with mentalizing^{10,13,29,38,48}, providing convincing face validity of this signature. Its high predictive performance in testing datasets implies that mentalizing recruits consistent and reliable neural configurations in the brain, not only in healthy controls, but also in clinical cohorts and developing adolescents. The MS also extended to other tasks (social-feedback and attributional versus factual inferences), predicting conditions that potentially evoke implicit, spontaneous mentalizing. This suggests that implicit and explicit forms of mentalizing largely share a common neural basis, in keeping with previous research⁴⁹. Thus, subnetworks of mentalizing (implicit-explicit) may be building on a common anatomical architecture that is here covered by the MS.

Recent work proposes that self- and other-mentalizing activate common neural constellations or representations and recruit largely overlapping regions^{29,30,48}. In contrast, our findings suggest that, while they share a large common core (as evident in the MS), they can also be distinguished reliably based on both distributed and local activation patterns. The Self-RS and the SvO-RS had positive weights in anterior mentalizing regions around cortical midline structures, such as mPFC, ACC, thalamus, caudate nucleus, frontal eye fields, and insula, extending to the frontal operculum. This pattern is in line with previous meta-analyses^{31,34,36}. The mPFC has been consistently related to mentalizing, and especially its ventral part, with self-related processing^{33,39,42,50,51}. The ventral-to-dorsal linear³¹ and curvilinear gradient³² hypotheses propose that mPFC is involved in both self- and other-person related processing. Accordingly, our results show that mPFC is involved in both self- and other-related mentalizing with different neural configurations, and that the information encoded in mPFC seems to predict self-mentalizing more consistently. Besides, self-processing recruited reward-processing circuits (e.g., striatum), in line with the idea that introspection about oneself is intrinsically rewarding^{52–54}. The involvement of subcortical areas (e.g., thalamus) and cortical areas associated with interoception (e.g., insula) is in line with previous work suggesting that one has access to affective/physiological information during self-mentalizing, which are less pertinent during mentalizing about others⁵⁵, and that interoceptive information might be an important contribution to one's sense of self^{26,56,57}.

The Other-RS and (negative weights of) the SvO-RS subsumed more posterior and lateral regions, such as TPJ, precuneus/PCC, STS, and vIPFC, resonating with previous findings^{10,11,13,32,37–39}. Interestingly, our results concerning the precuneus/PCC do not align with the known subdivisions of the default-mode⁵⁸ or mentalizing networks³⁰. This region is thought to be part of the midline core of the default-mode network⁵⁸ and of the medial subsystem of the mentalizing network³⁰, with stronger associations with self-processing. However, here in our results, while the precuneus/PCC cluster also contains information about self-mentalizing, it appears more strongly associated with other-related mentalizing, along with the areas that are part of the lateral subdivisions of both networks (e.g., TPJ, temporal lobes). A potential

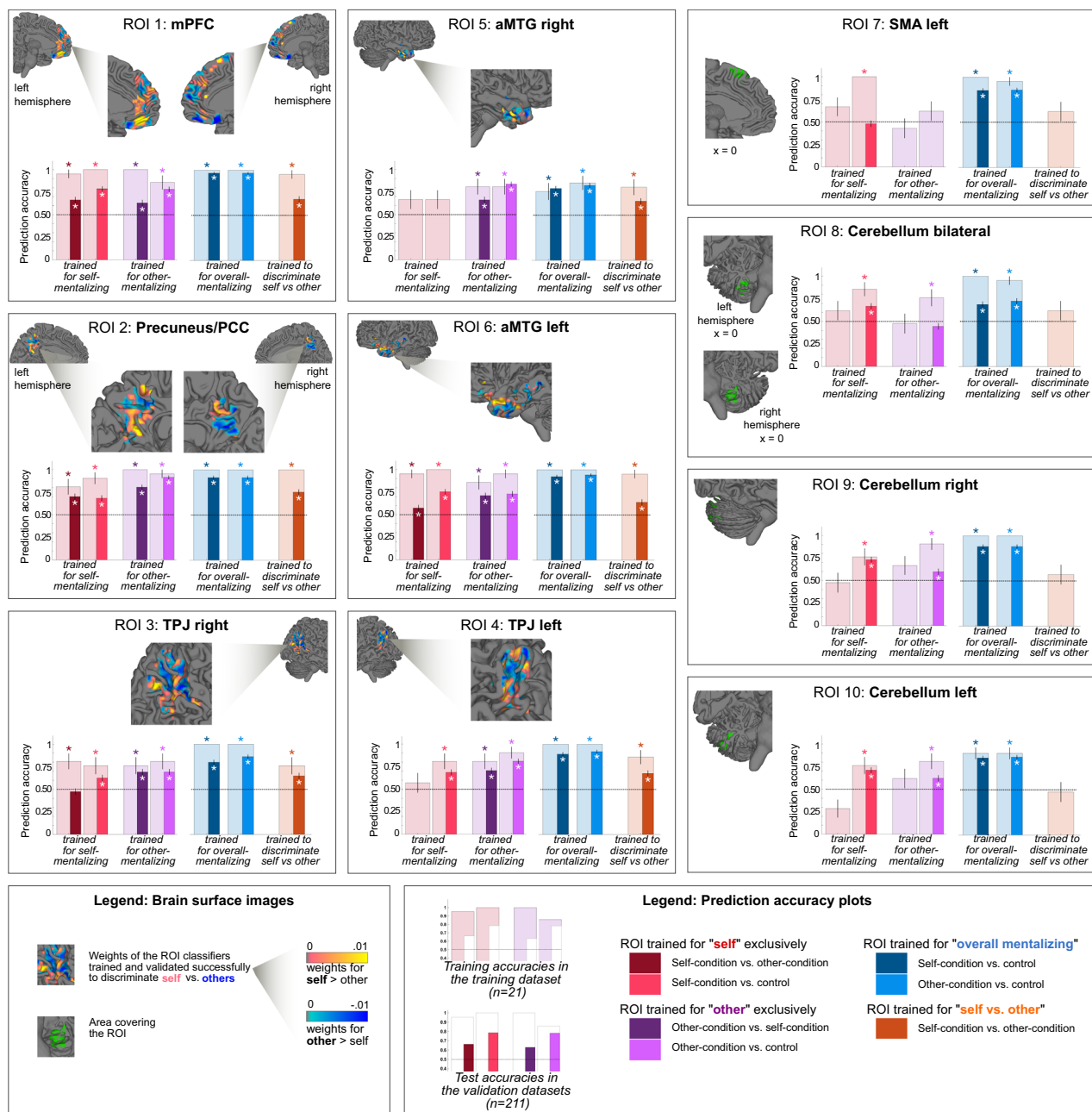


Fig. 8 | Results of the ROI analysis. Cross-validated training ($n = 21$) and validation ($n = 211$) results for ten ROIs derived from a term-based meta-analysis ("mentalizing," NeuroSynth⁴⁷): mPFC, precuneus/posterior cingulate cortex (precuneus/PCC), right temporoparietal junction (TPJ), left TPJ, right anterior middle temporal gyrus (aMTG), left aMTG, left supplementary motor area (SMA), right cerebellum, left cerebellum, and a bilateral cerebellum cluster. The ROIs that were consistently predictive of both self- and other-related mentalizing were: mPFC, precuneus/PCC, TPJ, and aMTG. Four different classifiers were trained for each ROI: (i) to predict Self exclusively, (ii) to predict Other exclusively, (iii) to predict overall mentalizing, and (iv) to discriminate the Self versus Other. If an ROI classifier had significant (above-chance) accuracy in the cross-validated training dataset, it was further tested across

the testing datasets. The cross-validated training accuracies are shown by the wider bars in the background and the validation accuracies by the thinner bars in the foreground. Bar heights illustrate prediction accuracy estimates, and error bars indicate $\pm$ standard error of these accuracy estimates. Classification performances significantly above the chance level (50%) are marked with asterisks (two-tailed $p < 0.05$). The surface images illustrate the weights of the classifiers that were trained to discriminate Self versus Other. Orange-yellow colors show areas with positive weights towards Self, blueish colors show areas associated with positive weights towards Other. ROIs that did not show significant above-chance training and independent validation accuracy for self-versus-other discrimination are shown in green. $*p < 0.05$.

explanation here could be the involvement of the precuneus in mental orientation/imagery^{59,60}.

The mentalizing signatures generalized to data from different types of cohorts, including healthy adults, schizophrenia and bipolar samples, and adolescents. This shows that the overall neurobiological configurations of self- and other-related mentalizing are comparable

across these populations and that the signatures have utility in different types of study populations. However, the Self-RS, the Other-RS, and the SvO-RS also showed sensitivity to clinical status, since self/other differentiation in the pattern responses was significantly less pronounced in participants with schizophrenia compared to matched healthy controls. This finding is noteworthy and in line with clinical

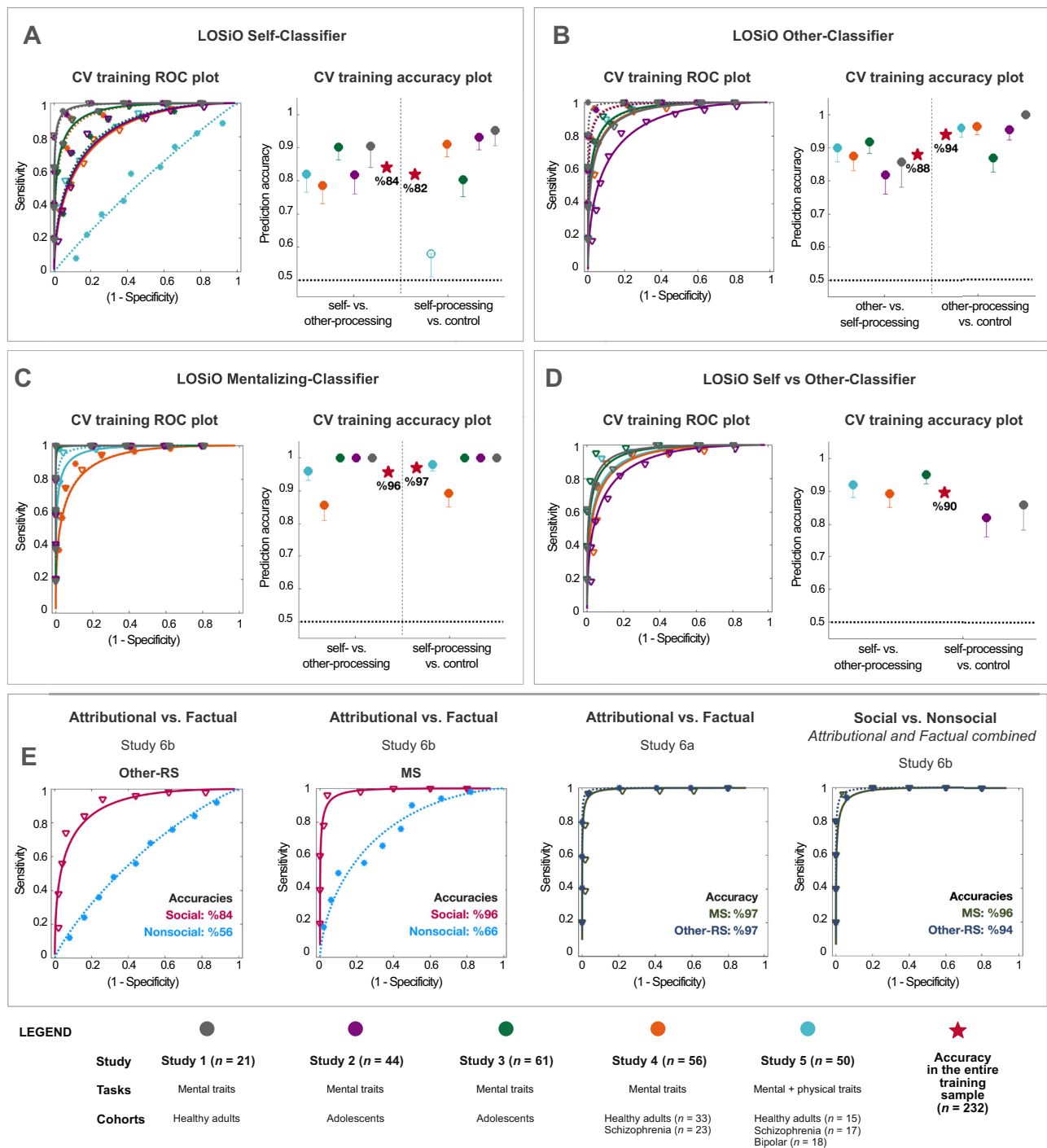


Fig. 9 | Leave-one-site-out (LOSio) classifiers. We trained additional classifiers that combined data across all datasets that used a trait-evaluation task ($n = 232$) to rule out that the results were driven by idiosyncratic features of the original training dataset. The classifiers were trained using the same analytic strategy as the main mentalizing signatures, but using five folds, with each study being held out in one-fold. **A** Self-classifier, **B** Other-classifier, **C** Mentalizing-classifier, **D** Self-vs-Other classifier. The left plot in each panel depicts the cross-validated receiver operating characteristic (ROC) plots. The right plot in each panel shows the cross-validated

classification accuracies in each dataset separately. Red stars show the out-of-sample accuracy across the whole sample ($n = 232$), while the colored circles represent the cross-validated accuracy estimates for each held-out study when classifiers were trained on data from all other studies. The one-sided error bars in accuracy plots illustrate the standard error of these accuracy estimates. **E** ROC plots for classifications in the inference task (Study 6) that involve attributional versus factual statements in social and nonsocial settings.

observations of alterations in self-perception, mentalizing, and the ability to discriminate between self- and other-generated thought^{36,45,46,61}. The responses in the bipolar sample did not differ significantly from those of the healthy adults, but given the limited sample size of the bipolar group in the present analysis, future work is needed to assess more fine-grained (and potentially context-

dependent) differences in mentalizing responses in bipolar disorder, as well as in other psychiatric and neurodevelopmental disorders with potential mentalizing deficits.

In the two adolescent samples (aged 11–18 years), greater age was associated with better self/other differentiation of the Self-RS, the Other-RS, and the SvO-RS. In other words, responses of these brain

signatures to self- and other-related mentalizing were more different from each other in older adolescents. Considering that the signatures were developed using an adult dataset, this reflects an ongoing development of the mentalizing neurobiology^{62,63} to become more adult-like and more differentiated for different targets of mentalization, during the age span covered in our study (11–18 years). Future studies could test the role of pubertal development instead of chronological age, and test the signatures in younger children and those with developmental disorders.

Of note, the mentalizing signatures were also implicated in making inferences about nonsocial targets—albeit to a lesser degree compared to social targets (see Figs. 6D and 9E). An earlier study⁶⁴ that used a similar inference task as Study 6 found that social and nonsocial attributions displayed overlapping activations in dmPFC, STS, TPJ, and lateral orbitofrontal cortex. A common denominator in both conditions may reflect domain-general attributional processing⁶⁵, which at least partially builds on semantic cognition⁶⁶. Likewise, processes supporting mentalizing may also more broadly subserve extraction of conceptual and semantic meaning for non-social stimuli^{67,68}, in line with the default mode network's putative meaning-making function^{51,69,70}. Thus, the mentalizing signatures may involve components of both domain-general representational processing, domain-specific social processing, and target-specific (self and other) processing components.

We note that the signatures' classification performance in the cross-validated training is consistently higher than in the independent testing datasets. This could potentially be due to leakage of information across the different participants in the training dataset, who all performed the same task in the same scanner and experimental settings. However, the additionally developed leave-one-site-out (LOSIO) classifiers that combined data across five studies performed similarly to the brain signatures that were developed using data from only the training dataset (Study 1). This provides evidence that the results of this study are not artifacts of our analytic choices and that the methodological approach fits the aims of the study well, as the results are robust with different permutations. The LOSIO classifiers showed consistently high classifications for the original training dataset (Study 1), which further provided support for the use of this study for training. Yet, because the LOSIO classifiers used information from $n = 232$ participants during model development, it has the potential to generalize very well to independent contexts. A key feature of the brain signature approach¹⁵ employed here is that the brain models are not necessarily final products, but tunable, improvable working models. Therefore, we endorse the use of LOSIO classifiers in future research, along with the main mentalizing signatures, for further evaluation of their utility in diverse settings.

The high separability of self- and other-referential processing in Study 1 might also reflect the fact that the other-condition in the mentalizing task was about a relatively unfamiliar person (a confederate). This may lead to a greater difference in self- versus other-related processing compared to studies in which mentalizing was about a familiar or close person, which is often more closely related to the self and likely also engages self-referential processing. This could potentially also explain why the Other-RS was less successful in separating other-related from self-related feedback in the sample of romantic couples (Study 7), since participants may consider their partners as an extension of themselves. Relatedly, the lower performance of the signatures in Study 5 (compared to the other testing datasets) could be due to the use of physical attributes (in addition to personality traits), which rely less on mentalizing.

Of note, the mentalizing signatures were trained to capture trait-related mentalizing. Here, we provide initial evidence that they extend beyond explicit reflections on trait-level representations. However, they also show above-chance classification in nonsocial inferences (see Figs. 6D and 9E), which might raise concerns about their specificity to

social agents. Future studies are needed to further test how well the mentalizing signatures separate mentalizing about other agents from other higher-order, conceptual, and abstract cognitive processing. Future studies should also test their specificity, as well as their validity in other forms (e.g., affective, spontaneous, or nonverbal) of mentalizing. A natural next step could be to test whether other dimensions of mentalizing^{6,12} (e.g., cognitive versus affective) or different types of mental state content⁷¹ (i.e., stable versus transient or beliefs versus preferences) are also associated with robust, replicable multivariate patterns. Recent work^{72,73} investigated valence and self-relevance as two key components of internal thought. It is an open question how the brain integrates valence with different targets during mentalizing. Finally, future research can also build on this work by testing the signatures in other mentalizing tasks (e.g., false-belief, emotion imagery), on a greater variety of targets⁷⁴, and in other related mental processes (e.g., autobiographical memory retrieval).

In conclusion, we trained and validated four brain signatures that predict self-related, other-related, and both types of mentalizing in multiple independent datasets that used different variants of a standard trait-mentalizing task in diverse populations. Our findings imply that self- and other-mentalizing use largely dissociable neural mechanisms that build on a foundational overall mentalizing capacity. Indeed, the four mentalizing signatures possess potential for use as predictive neuromarkers of mentalizing about self and others, for example, by testing how brain circuits of mentalizing are engaged or modulated by different types of experimental conditions, and how they might be altered in different psychiatric and neurodevelopmental populations.

Methods

This study was conducted in accordance with all relevant ethical regulations. All contributing studies were approved by their respective institutional ethics committees. All participants gave written informed consent and were compensated for their participation via monetary means or gifts. For additional information, please refer to the original studies and Supplementary Table 1.

Participants

The present study pooled data from a total of $n = 390$ participants (174 females and 216 males, $M_{\text{age}} = 26.9$, $SD_{\text{age}} = 11.7$) from eleven cohorts participating in seven independent studies (see Fig. 1A and Supplementary Table 1). These eleven cohorts comprised six samples of healthy (neurotypical) adults, two samples of healthy adolescents, and three samples of adults with a diagnosed psychiatric condition. While most of the datasets have been previously published separately, the analyses reported here are new, and the seven studies have not been previously analyzed in combination.

The training sample (Study 1⁴³) consisted of $n = 21$ healthy adults (10 women and 11 men, $M_{\text{age}} = 23.7$) who were recruited at the University of Geneva, Switzerland. One additional participant with structural abnormalities in the brain was excluded from the original study and the present analysis.

Study 2⁷⁵ included $n = 44$ healthy adolescents (23 females and 21 males, $M_{\text{age}} = 16$, $SD_{\text{age}} = 1.86$, age range = 12.0–18.8) who were recruited from secondary schools in Geneva, Switzerland. In the original study, one additional participant was excluded due to structural abnormalities in the brain, three due to non-completion of the paradigm, one due to signs of substance use, and five due to excessive movement.

Study 3⁷⁶ included $n = 61$ healthy adolescents (27 females and 34 males, $M_{\text{age}} = 12.9$, $SD_{\text{age}} = 0.43$, age range = 11.6–14.2) who were recruited for a longitudinal project from secondary schools in the Netherlands. An additional 18 participants were excluded from the original study due to excessive movement, incorrect task completion, or measurement errors. Additionally, 5 participants were not included

in this study (original study $n = 66$) because they did not provide explicit consent for data sharing.

Study 4^{77,78} included $n = 33$ healthy adults (14 women and 19 men, $M_{\text{age}} = 41.7$) and $n = 23$ adults with schizophrenia (7 women and 16 men, $M_{\text{age}} = 37.0$). The two cohorts were matched on age, sex, and a measure of general intelligence. The patients were recruited from a local psychiatric hospital in Barcelona, Spain, based on diagnostic interviews. The schizophrenia diagnosis was confirmed using the structured clinical interview for DSM disorders⁷⁹ (SCID).

Study 5⁸⁰ included three cohorts: healthy adults ($n = 15$, 6 women and 9 men, $M_{\text{age}} = 33.3$), adults with schizophrenia (SZ, $n = 17$, 6 women and 11 men, $M_{\text{age}} = 35.5$), and adults with bipolar disorder (BD, $n = 18$, 10 women and 8 men, $M_{\text{age}} = 40.3$). The clinical cohorts were recruited from a local hospital in the north of the Netherlands. The diagnoses of the patients were confirmed using the Mini International Neuropsychiatric Interview-Plus⁸¹ 5.0.0 (MINI-Plus). All BD patients were chosen among those who had a history of at least one psychotic episode. All three cohorts were matched with one another on age, sex, level of education, and a measure of general intelligence. The SZ and BP patients were additionally matched on the level of cognitive and clinical insight as measured by the Schedule of Assessment of Insight⁸²-Expanded version (SAI-E, clinical insight) and the Beck Cognitive Insight Scale⁸³ (BCIS, cognitive insight).

Study 6a⁸⁴ included $n = 59$ healthy adults (26 women and 33 men; $M_{\text{age}} = 28.3$), and Study 6b⁸⁴ included $n = 50$ healthy adults (19 women and 31 men; $M_{\text{age}} = 33.6$) from the Los Angeles metropolitan area; these final sample sizes reflect exclusions described below. One participant from each study was excluded due to poor task performance ($>70\%$ missed trials), and four additional participants from Study 6b were excluded due to outlier behavioral scores (>3 SDs from the mean, $n = 1$) and excessive scanner motion ($n = 3$). All participants were right-handed and had IQs in the normal range (assessed via the WASI-II).

Study 7⁸⁵ included $n = 49$ healthy adults (26 women and 23 men, $M_{\text{age}} = 22.7$) in romantic relationships recruited from the Tucson, Arizona community and surrounding areas. Seven participants were excluded from the original study due to missing or inadequate imaging data. Community members were eligible to participate if they had been in a romantic relationship for at least six months, had no contraindications for MRI scanning, and did not meet criteria for active psychosis or mania at the time of screening. Both members of the couple completed all components of the study, including the social feedback fMRI task.

Tasks

Training dataset. We used a trait evaluation task, adapted from ref. 42 in the training dataset (Study 1). Participants were asked to rate the extent to which certain personality adjectives (e.g., talkative, daring) described themselves (self-condition) or a same-sex confederate (other-condition), with whom they thought they were interacting in a preceding decision-making task⁴³. In the control condition, they were asked to count the syllables for each trait adjective.

Participants met the confederate in person at the beginning of the experimental session, where they were briefed that one of them would be tested in the scanner, the other one in a separate room, and that they would interact via a computer interface. They first participated in a social decision-making task resembling a serial dictator game paradigm, whereby the participant was given the choice to share or keep the resources with the confederate in a series of trials⁴³. After this task, participants were introduced to the trait-evaluation task and were asked to rate their co-player.

The task consisted of 150 trials in total, presented in 30 blocks (ten per condition). Each block contained five trials and lasted for 20 s, with an inter-block interval (fixation cross) of 8 s. Half of the blocks only contained negative and the other half only positive adjectives. Each

word appeared once for each condition (50 adjectives were used in total).

Each trial started with a 3-s cue above the fixation cross, reminding the participant of the condition (Self, Other, or syllables), which was followed by the presentation of the adjective and a Likert scale (1–4) on which the participant was expected to select using a button press device. A word was presented every 4 s. If the participant made a choice sooner, the word would disappear for the remaining time of the 4 s before the next trial. The order of the blocks and words displayed in each block was randomized. The task was programmed in and presented using E-Prime 2.0 (Psychology Software Tools, Pittsburgh, PA).

Testing and extension datasets. All testing datasets included similar mentalizing tasks with the same three conditions (Self, Other, and Control) as the training task and used a block design (for an overview, see Fig. 1A and Supplementary Table 1). The main difference between testing tasks were the stimuli that were presented: Studies 2 and 3 used trait adjectives, Study 4 (with two cohorts) used trait statements, and Study 5 (with three cohorts) used a mix of trait and physical statements. The studies also used a variety of *others* in the other-condition: Best/close friend (Study 2), similar and dissimilar classmate (Study 3), relative or close friend (Study 5), and acquaintance (Study 4). Studies 2 and 5 collected responses using Likert scales, and Studies 3 and 4 asked for binary choices (yes/no). Finally, the control conditions also varied between studies. Studies 4 and 5 presented general knowledge statements; Study 3 asked participants to search for specific letters in each word; Study 2 asked them to count the syllables in the trait words.

The extension datasets (Studies 6 and 7) used different tasks to test how well the mentalizing signatures extend to tasks that involve other forms of social (or nonsocial) processing. Studies 6a and 6b used an inference task based on visual stimuli⁸⁴ (see Fig. 6A). In Study 6a, participants made rapid yes/no judgments in response to social scenes (intentional hand actions and emotional faces) that required either attributional inferences about hidden/internal states (mentalizing) or factual inferences about observable features, resulting in a 1 (stimulus type: social) $\times$ 2 (inference type: attributional, factual) factorial design. Study 6b included an additional nonsocial condition using images of natural scenes (e.g., rain pouring out of a gutter), each paired with attributional or factual inference prompts, resulting in a 2 (stimulus type: social, nonsocial) $\times$ 2 (inference type: attributional, factual) design. Study 7 (see Fig. 7A) investigated spontaneous empathy when negative or positive feedback thought to be from other participants (confederates) was directed to participants' romantic partners (or one's self). Participants received positive and negative feedback either about themselves (self-condition), their romantic partners (other-condition), or no feedback (control-condition). The original study (see ref. 85 for more details) included additional conditions, such as directing feedback to both self and partner, which are not analyzed in the present manuscript.

fMRI data acquisition, preprocessing, and first-level analysis

Training dataset. The training MRI images were acquired on a 3 T Magnetom TIM Trio whole-body scanner (Siemens, Germany) with the product 12-channel head coil. A T1-weighted MPRAGE sequence (TR = 1900 ms, TI = 900 ms, TE = 2.27 ms, voxel size $1 \times 1 \times 1$ mm) was used to acquire structural anatomical images. Functional images were obtained using a standard T2-weighted echo-planar imaging sequence (2D-EP, TR = 2100 ms, TE = 30 ms, flip angle 80° , voxel size $3.2 \times 3.2 \times 3.2$ mm) that scanned the whole brain in 36 sequential slices. An automated shimming procedure was included to minimize magnetic field inhomogeneities.

As it is the standard practice in the development of brain signatures^{15,17,18,21}, we used statistical maps from first-level analysis (e.g., single-subject contrast images for each condition versus implicit

baseline) as input for training and for testing the classifiers. SPM8 (Wellcome Department of Imaging Neuroscience, UCL, London, UK) and Matlab® (2022b; The MathWorks Inc.) were used for image preprocessing and first-level analysis. A standard preprocessing pipeline was performed, which included spatial realignment and reslicing of the functional images, coregistration, unified segmentation and normalization to the Montreal Neurological Institute (MNI) echo planar imaging template (voxel size: 2 mm³) and finally spatial smoothing using an 8 mm full-width at half maximum (FWHM) Gaussian kernel.

During first-level analysis, we included six task (block) regressors that were composed of 3 task conditions by positive and negative valence. The task regressors were convolved with a canonical hemodynamic response function. We also included six additional regressors of no interest to control for motion. A high-pass frequency filter (128 s) and autocorrelation corrections (using restricted maximum likelihood and an autoregressive model) were used in model estimation.

Testing datasets. We conducted a non-systematic literature review to identify recent fMRI studies of comparable mentalizing tasks with at least three conditions (Self, Other, and non-mentalizing Control condition) and in which participants rated trait adjectives or statements. We emailed the authors of seven studies that we identified and received positive responses from four of them. Data from these four studies were included in the current study as testing datasets (Studies 2–5). In addition, to test the signatures in different types of tasks, we included an unpublished dataset⁸⁵ (Study 7) and data from recently published work⁸⁴ (Study 6) as extension datasets. Please refer to the original studies (see Supplementary Table 1) for details of image acquisition, preprocessing, and first-level analysis.

The authors provided the subject-level contrast images of the different experimental conditions (versus the implicit baseline). Voxel weights were normalized within each image using L2-norm, and contrast images were resampled onto the same image space as the training dataset using linear resampling. For Study 3, which included two different Other conditions (similar and dissimilar classmates), we averaged the contrast images across these two conditions, resulting in a single Other condition, as in the training and the other testing datasets (analyzing them separately did not alter the results). Study 6 included two types of social conditions: hands and faces. As we were not interested in the social processing of hand- or face-related information, we averaged the contrast images of these conditions, creating a single social condition image per participant.

Data analysis

Cross-validated training. Using 10-fold cross-validation, we trained three support-vector-machine (SVM) classifiers in Study 1 that discriminate each condition from the other two conditions: the Self-RS was trained to separate the self-condition from the other- and control conditions, the Other-RS to separate the other-condition from the self- and control conditions, and the MS was trained to separate both mentalizing conditions (Self and Other) from the control condition. In addition, we trained one signature to specifically separate the Self from the Other condition (the SvO-Signature). The input for SVMs was subject-level contrast images for each of the three task conditions (versus implicit baseline). The cross-validation folds were organized to ensure that the three contrast images from each participant were kept together within a single fold (nine folds with two participants and one fold with three participants). To reduce the possibility that classifiers opportunistically used non-mentalizing related processes (e.g., visual information), we applied a mask in the training dataset that includes key social-cognition regions (see Fig. 1B). This mask was computed as the union of six term-based meta-analytic maps (association and uniformity maps for “mentalizing,” “self-referential,” and “social,” downloaded from NeuroSynth⁴⁷ on 06/06/2024) to cover a wide range of brain areas typically activated by social and self-related processing.

While the use of a mask was intended to minimize involvement of the confounding voxels that are not directly related to social cognition, this might raise the question of how more extensive, whole-brain classifiers might perform. As shown in Supplementary Fig. 1, applying a whole-brain gray matter mask yielded highly similar results.

SVMs were chosen based on their high performance for binary linear classification problems in high-dimensional and low-sample data^{86–88}. Linear SVM is a commonly used method in cognitive neuroscience decoding problems^{48,89,90}, which is computationally efficient and provides interpretable, robust weight maps. The algorithm fits a hyperplane that classifies true and false classes by assigning weights for each feature (i.e., each voxel). The fitting is performed for each of 10 folds, in which the data of 90% of participants are used for fitting the classifier, and the resulting classifier is tested on the remaining 10% hold-out participants' data, allowing for assessing its classification performance in independent hold-out data. Because using a one versus the rest approach in SVM (e.g., Self versus Other and Control, see Fig. 1C) may add bias into the model by favoring the majority class, we fitted weighted SVM models with a ridge parameter of 0.5 to train the Self-RS, the Other-RS, and the MS (but not the SvO-RS). To avoid overfitting, the SVMs were otherwise trained using default parameters (regularization parameter $C=1$, linear kernel function, number of folds = 10). The cross-validated distance from the hyperplane of hold-out images was used to calculate the receiver operating characteristic (ROC) curves and the accuracy for each classification. Each SVM results in a weight map with one value per voxel. These voxel weights are effectively the predictive weights of the true class, yielding brain signatures of mentalizing that can be applied to other brain images to obtain a single pattern expression score per image. The classification accuracies in the training dataset were tested using two-alternative forced-choice predictions and binomial tests.

We used bootstrapping to train SVM classifiers, producing stable model estimates and voxel-wise p -values. In the first step, cross-validated training generated prediction accuracies and a final model output (i.e., weight map) that is trained on all samples. In the next step, the bootstrapping procedure resampled the training data 5000 times with replacement, yielding statistical maps with two-tailed, uncorrected p -values. These maps were then thresholded using FDR-correction at $q < 0.05$ with a minimum cluster size of 10 voxels. Note that corrected maps were only used for visualization and inference. The unthresholded weight maps constitute the mentalizing signatures and are used for analysis.

Validation in independent datasets. The four mentalizing signatures were applied to the testing datasets by computing the pattern similarity values as the matrix dot product between mentalizing signatures and the subject-level (1st level) contrast images of each study, yielding one scalar value per condition and participant and signature. The predictions followed a binary forced-choice principle using paired observations. The signatures' classification accuracies were assessed using receiver operating characteristic (ROC) analysis and binomial tests using a two-sided significance threshold of $p < 0.05$. The accuracy rate, area under the curve (AUC), and specificity/sensitivity obtained from these analyses provide the benchmark measures of binary classifications in multivariate modeling⁴⁴. We also report Cohen's d in the text as a more standardized measure to quantify prediction effect sizes across studies and samples. Cohen's d is calculated as the mean of differences between condition values (either pattern expression or distance from the hyperplane) divided by the pooled standard deviation of differences (where difference = $\text{input_values}[\text{binary_class}] - \text{input_values}[-\text{binary_class}]$). Because the three signatures (Self-RS, Other-RS, and MS) were each tested in two comparisons (e.g., Self versus Other and Self versus Control), the main text reports the average statistical values across these comparisons, with the complete results provided in the Supplementary Tables. Similarly, when

accuracy statistics from a classifier were combined across conditions or studies, we report the weighted average values in the main text and present the full set of statistics in the Supplementary Material.

To examine the effects of individual differences (e.g., sex and age) on pattern expression values, we used linear mixed-effects models. We pooled data from five testing datasets that used the same task design ($n = 211$; Studies 2–5). The signature fit (true minus false class score) was defined as the difference between the pattern expression for the true condition and that for the false condition. Specifically, we subtracted the false-condition value from the true-condition value for each signature: Self-RS (Self – Other), Other-RS (Other – Self), SvO-RS (Self – Other), and MS ($[\text{Self} + \text{Other}]/2 - \text{Control}$). Age and sex were entered as fixed effects, and study was included as a random intercept in the model: $\text{signature_fit} - \text{age} + \text{sex} + (1 | \text{study})$.

Group comparisons of pattern expression values. To quantify how well signatures discriminate true versus false conditions and to easily compare self/other separation across groups, we used the signature fits for each signature, the calculation of which is explained in the section above. For statistical comparisons between participants with schizophrenia and healthy controls, we ran linear mixed effects models for each of the signatures, with the signature fits (true minus false class score) as the dependent variable. The group constituted the fixed effect variable (healthy controls [$n = 48$] versus schizophrenia sample [$n = 40$]). Because the data came from two different studies (Study 4 and Study 5), the study was included as a random effect variable. Therefore, the model equation was $\text{signature_fit} - \text{group} + (1 | \text{study})$.

To test whether the bipolar sample ($n = 18$) differed from healthy adults ($n = 15$) in Study 5 regarding their pattern expression values, we performed independent samples t-tests for each of the signatures consecutively. The dependent variable was the “signature fit” as above.

Associations with age. We tested whether the adolescents’ age was associated with self/other discrimination using linear mixed effects models. We included age as the fixed effect, Study (Study 2 and Study 3) as the random effect, and sex as the control variable (fixed effect). The dependent variable was the difference of the true minus false classes for each of the signatures, as detailed above. Therefore, the model formula was $\text{signature_fit} - \text{age} + \text{sex} + (1 | \text{study})$.

Testing the signatures in extension tasks. In order to test how well the signatures generalized to the two extension tasks, we applied them to the subject-level contrast images for each experimental condition by computing their dot products. When a signature is applied to another context (e.g., a different task), one could apply either (i) a binary prediction test or (ii) a linear-effects model to test for condition differences, depending on the study aim. Binary prediction is arguably most appropriate when there is a ground truth or clear directional expectation, whereas linear models are better suited to more distant applications (e.g., novel tasks or populations). Accordingly, we tested the signatures in Study 6 using (i) binary prediction tests, expecting the signatures to selectively predict the attributional statements involving mentalizing, consistent with the original study⁸⁴. We followed the same principles for predictions as outlined above (validation in independent tests), where pattern expression values and predictions are subjected to ROC analyses and binomial tests. In Study 7, we instead used (ii) linear-effects models, submitting pattern expression values to repeated-measures ANOVAs across the three social-feedback conditions, followed by Bonferroni-corrected t-tests for pairwise comparisons.

Region-of-interest (ROI) analyses. To test the ability of local patterns to predict mentalizing and to separate self-related versus other-related

mentalizing, we trained and cross-validated ROI classifiers using a parallel approach to the whole-brain classifiers. We downloaded a term-based meta-analytic map for “mentalizing” from NeuroSynth⁴⁷ on 14/09/2022 that included 151 studies. We selected clusters that contained more than 200 voxels, resulting in the following ten ROIs: mPFC, bilateral TPJ, bilateral anterior MTG, precuneus/PCC, right SMA, and three clusters in the cerebellum. In each ROI, we trained four classifiers in the training dataset, using 10-fold cross-validation, and validated them in the remaining datasets, to (i) predict self-mentalizing (versus the two remaining conditions), (ii) predict other-mentalizing (versus the two remaining conditions), (iii) predict mentalizing (self- and other-mentalizing versus control condition), and (iv) differentiate self-versus other-processing. We tested these ROI classifiers in the testing datasets using the same parameters and analytic approach as outlined above in the main whole-brain analyses.

Leave-one-site-out classifiers trained on all datasets. To fully benchmark how well our approach generalizes across independent datasets, we trained additional classifiers based on the pooled cohorts from all datasets that used a variation of the trait evaluation task (Studies 1–5; $n = 232$). We defined five folds, each representing one study, so that the classifiers were trained on four studies and tested on the remaining study in each fold. The cross-validated training of bootstrapped classifiers followed the same analytic flow as the main classifiers (see cross-validated training above). Besides the number of folds, the only other change was the number of bootstrapped samples, which we limited to 2500 samples for computational efficiency. We additionally tested these classifiers in the inference task from Study 6 using the binary prediction tests as described above for the main signatures.

General statistical approach. All data analysis and data visualizations were performed using Matlab® R2022b software and the Canlab toolbox (<https://github.com/canlab>). Violin plots in Figs. 4 and 7 were generated using the MATLAB function `daviolinplot`⁹¹. Scatter plots in Fig. 5 were generated in R Statistical Software⁹² (v12.1) using `ggplot2`⁹³. Statistical inference used a significance threshold of $p < 0.05$, unless otherwise noted.

Reporting summary

Further information on research design is available in the Nature Portfolio Reporting Summary linked to this article.

Data availability

Source data for all figures and single-subject contrast images of the training dataset (Study 1) are available at⁹⁴ <https://doi.org/10.6084/m9.figshare.29908139>. Contrast images from Study 4a can be accessed via OpenNeuro (<https://openneuro.org/datasets/ds001618/versions/1.0.1>), and data from Study 6 are available through the NIHM Archive (https://nda.nih.gov/edit_collection.html?id=2643). Imaging data from five contributing studies (Studies 2, 3, 4b, 5, and 7) are not shared online for privacy reasons; however, deidentified data from these studies can be made available in line with applicable data protection and privacy regulations upon request to the corresponding authors.

Code availability

Mentalizing signatures for use in future studies, as well as custom MATLAB scripts for analyses are available at⁹⁵: https://github.com/ldmk/2025_MentalizingSignatures.

References

- Moore, C. & Frye, D. The acquisition and utility of theories of mind. in (eds Frye, D. & Moore, C.) *Children’s Theories of Mind: Mental States and Social Understanding*, 1–14 (Lawrence Erlbaum, 1991).

2. Wellman, H. M. *Making Minds: How Theory of Mind Develops* (Oxford University Press, 2014).
3. Brüne, M. & Brüne-Cohrs, U. Theory of mind—evolution, ontogeny, brain mechanisms and psychopathology. *Neurosci. Biobehav. Rev.* **30**, 437–455 (2006).
4. Debbané, M. et al. Attachment, neurobiology, and mentalizing along the psychosis continuum. *Front. Hum. Neurosci.* **10**, 406 (2016).
5. Gray, K., Jenkins, A. C., Heberlein, A. S. & Wegner, D. M. Distortions of mind perception in psychopathology. *Proc. Natl. Acad. Sci. USA* **108**, 477–479 (2011).
6. Luyten, P., Campbell, C., Allison, E. & Fonagy, P. The mentalizing approach to psychopathology: state of the art and future directions. *Annu. Rev. Clin. Psychol.* **16**, 297–325 (2020).
7. Sharp, C. Mentalizing problems in childhood disorders. in *Handbook of Mentalization-Based Treatment*, 101–121 (John Wiley & Sons, Ltd, 2006).
8. Sloover, M., van Est, L. A. C., Janssen, P. G. J., Hilbink, M. & van Ee, E. A meta-analysis of mentalizing in anxiety disorders, obsessive-compulsive and related disorders, and trauma and stressor related disorders. *J. Anxiety Disord.* **92**, 102641 (2022).
9. Johnson, B. N., Kivity, Y., Rosenstein, L. K., LeBreton, J. M. & Levy, K. N. The association between mentalizing and psychopathology: a meta-analysis of the reading the mind in the eyes task across psychiatric disorders. *Clin. Psychol. Sci. Pract.* **29**, 423–439 (2022).
10. Frith, C. D. & Frith, U. The neural basis of mentalizing. *Neuron* **50**, 531–534 (2006).
11. Saxe, R. Uniquely human social cognition. *Curr. Opin. Neurobiol.* **16**, 235–239 (2006).
12. Schurz, M. et al. Toward a hierarchical model of social cognition: a neuroimaging meta-analysis and integrative review of empathy and theory of mind. *Psychol. Bull.* **147**, 293–327 (2021).
13. van Overwalle, F. & Baetens, K. Understanding others' actions and goals by mirror and mentalizing systems: a meta-analysis. *NeuroImage* **48**, 564–584 (2009).
14. Yang, D. Y.-J., Rosenblau, G., Keifer, C. & Pelphrey, K. A. An integrative neural model of social perception, action observation, and theory of mind. *Neurosci. Biobehav. Rev.* **51**, 263–275 (2015).
15. Kragel, P. A., Koban, L., Barrett, L. F. & Wager, T. D. Representation, pattern information, and brain signatures: from neurons to neuroimaging. *Neuron* **99**, 257–273 (2018).
16. Rosenberg, M. D., Casey, B. J. & Holmes, A. J. Prediction complements explanation in understanding the developing brain. *Nat. Commun.* **9**, 589 (2018).
17. Wager, T. D. et al. An fMRI-based neurologic signature of physical pain. *New Engl. J. Med.* **368**, 1388–1397 (2013).
18. Koban, L., Wager, T. D. & Kober, H. A neuromarker for drug and food craving distinguishes drug users from non-users. *Nat. Neurosci.* **26**, 316–325 (2023).
19. Rosenberg, M. D. et al. A neuromarker of sustained attention from whole-brain functional connectivity. *Nat. Neurosci.* **19**, 165–171 (2016).
20. Kim, J. et al. A dorsomedial prefrontal cortex-based dynamic functional connectivity model of rumination. *Nat. Commun.* **14**, 3540 (2023).
21. López-Solà, M. et al. Towards a neurophysiological signature for fibromyalgia. *Pain* **158**, 34–47 (2017).
22. Gabrieli, J. D. E., Ghosh, S. S. & Whitfield-Gabrieli, S. Prediction as a humanitarian and pragmatic contribution from human cognitive neuroscience. *Neuron* **85**, 11–26 (2015).
23. Koban, L. et al. An fMRI-based brain marker of individual differences in delay discounting. *J. Neurosci.* **43**, 1600–1613 (2023).
24. Woo, C. W., Chang, L. J., Lindquist, M. A. & Wager, T. D. Building better biomarkers: brain models in translational neuroimaging. *Nat. Neurosci.* **20**, 365–377 (2017).
25. Tamir, D. I. & Thornton, M. A. Modeling the predictive social mind. *Trends Cogn. Sci.* **22**, 201–212 (2018).
26. Qin, P., Wang, M. & Northoff, G. Linking bodily, environmental and mental states in the self—a three-level model based on a meta-analysis. *Neurosci. Biobehav. Rev.* **115**, 77–95 (2020).
27. Corradi-Dell'Acqua, C., Hofstetter, C. & Vuilleumier, P. Cognitive and affective theory of mind share the same local patterns of activity in posterior temporal but not medial prefrontal cortex. *Soc. Cogn. Affect. Neurosci.* **9**, 1175–1184 (2014).
28. Quesque, F. et al. Defining key concepts for mental state attribution. *Commun. Psychol.* **2**, 1–5 (2024).
29. Tan, K. M. et al. Electroencephalographic evidence of a common neurocognitive sequence for mentalizing about the self and others. *Nat. Commun.* **13**, 1919 (2022).
30. Wang, Y. et al. A large-scale structural and functional connectome of social mentalizing. *NeuroImage* **236**, 118115 (2021).
31. Denny, B. T., Kober, H., Wager, T. D. & Ochsner, K. N. A meta-analysis of functional neuroimaging studies of self- and other judgments reveals a spatial gradient for mentalizing in medial prefrontal cortex. *J. Cogn. Neurosci.* **24**, 1742–1752 (2012).
32. Parelman, J. M. et al. Overlapping functional representations of self- and other-related thought are separable through multivoxel pattern classification. *Cereb. Cortex* **32**, 1131–1141 (2022).
33. Fossati, P. et al. In search of the emotional self: an fMRI study using positive and negative emotional words. *Am. J. Psychiatry* **160**, 1938–1945 (2003).
34. Murray, R. J., Debbané, M., Fox, P. T., Bzdok, D. & Eickhoff, S. B. Functional connectivity mapping of regions associated with self- and other-processing. *Hum. Brain Mapp.* **36**, 1304–1324 (2014).
35. Northoff, G. et al. Self-referential processing in our brain—a meta-analysis of imaging studies on the self. *NeuroImage* **31**, 440–457 (2006).
36. van der Meer, L., Costafreda, S., Aleman, A. & David, A. S. Self-reflection and the brain: a theoretical review and meta-analysis of neuroimaging studies with implications for schizophrenia. *Neurosci. Biobehav. Rev.* **34**, 935–946 (2010).
37. Arioli, M., Cattaneo, Z., Parimbelli, S. & Canessa, N. Relational vs representational social cognitive processing: a coordinate-based meta-analysis of neuroimaging data. *Soc. Cogn. Affect. Neurosci.* **18**, 1–19 (2023).
38. Tamir, D. I., Bricker, A. B., Dodell-Feder, D. & Mitchell, J. P. Reading fiction and reading minds: the role of simulation in the default network. *Soc. Cogn. Affect. Neurosci.* **11**, 215–224 (2016).
39. Wagner, D. D., Chavez, R. S. & Broom, T. W. Decoding the neural representation of self and person knowledge with multivariate pattern analysis and data-driven approaches. *Wiley Interdiscip. Rev. Cogn. Sci.* **10**, e1482 (2019).
40. Arioli, M., Cattaneo, Z., Ricciardi, E. & Canessa, N. Overlapping and specific neural correlates for empathizing, affective mentalizing, and cognitive mentalizing: a coordinate-based meta-analytic study. *Hum. Brain Mapp.* **42**, 4777–4804 (2021).
41. Molapour, T. et al. Seven computations of the social brain. *Soc. Cogn. Affect. Neurosci.* **16**, 745–760 (2021).
42. Kelley, W. M. et al. Finding the self? An event-related fMRI study. *J. Cogn. Neurosci.* **14**, 785–794 (2002).
43. Koban, L., Pichon, S. & Vuilleumier, P. Responses of medial and ventrolateral prefrontal cortex to interpersonal conflict for resources. *Soc. Cogn. Affect. Neurosci.* **9**, 561–569 (2014).
44. Scheinost, D. et al. Ten simple rules for predictive modeling of individual differences in neuroimaging. *NeuroImage* **193**, 35–45 (2019).

45. Bora, E., Yucel, M. & Pantelis, C. Theory of mind impairment in schizophrenia: meta-analysis. *Schizophr. Res.* **109**, 1–9 (2009).
46. Sprong, M., Schothorst, P., Vos, E., Hox, J. & Engeland, H. V. Theory of mind in schizophrenia: meta-analysis. *Br. J. Psychiatry* **191**, 5–13 (2007).
47. Yarkoni, T., Poldrack, R. A., Nichols, T. E., Van Essen, D. C. & Wager, T. D. Large-scale automated synthesis of human functional neuroimaging data. *Nat. Methods* **8**, 665–670 (2011).
48. Oosterwijk, S., Snoek, L., Rottevel, M., Barrett, L. F. & Scholte, H. S. Shared states: using MVPA to test neural overlap between self-focused emotion imagery and other-focused emotion understanding. *Soc. Cogn. Affect. Neurosci.* **12**, 1025–1035 (2017).
49. van Overwalle, F. & Vandekerckhove, M. Implicit and explicit social mentalizing: dual processes driven by a shared neural network. *Front. Hum. Neurosci.* **7**, 560 (2013).
50. Heatherton, T. F. et al. Medial prefrontal activity differentiates self from close others. *Soc. Cogn. Affect. Neurosci.* **1**, 18–25 (2006).
51. Koban, L., Gianaros, P. J., Kober, H. & Wager, T. D. The self in context: brain systems linking mental and physical health. *Nat. Rev. Neurosci.* **22**, 309–322 (2021).
52. Chavez, R. S., Heatherton, T. F. & Wagner, D. D. Neural population decoding reveals the intrinsic positivity of the self. *Cereb. Cortex* **27**, 5222–5229 (2017).
53. Northoff, G. & Hayes, D. J. Is our self nothing but reward? *Biol. Psychiatry* **69**, 1019–1025 (2011).
54. Tamir, D. I. & Mitchell, J. P. Disclosing information about the self is intrinsically rewarding. *Proc. Natl. Acad. Sci. USA* **109**, 8038–8043 (2012).
55. Maresh, E. L. & Andrews-Hanna, J. R. Putting the “Me” in “mentalizing”: multiple constructs describing self versus other during mentalizing and implications for social anxiety disorder. in (eds Gilead, M. & Ochsner, K. N.) *The Neural Basis of Mentalizing*, 629–658 (Springer International Publishing, 2021).
56. Babo-Rebello, M. & Tallon-Baudry, C. Interoceptive signals, brain dynamics, and subjectivity. in (eds Tsakiris, M. & De Preester, H.) *The Interoceptive Mind: From Homeostasis to Awareness*, 46–62 (Oxford University Press, 2018).
57. Garfinkel, S. N., Nagai, Y., Seth, A. K. & Critchley, H. D. Neuroimaging studies of interoception and self-awareness. in (eds Cavanna, A. E., Nani, A., Blumenfeld, H. & Laureys, S.) *Neuroimaging of Consciousness*, 207–224 (Springer, 2013).
58. Andrews-Hanna, J. R., Reidler, J. S., Sepulcre, J., Poulin, R. & Buckner, R. L. Functional-anatomic fractionation of the brain’s default network. *Neuron* **65**, 550–562 (2010).
59. Peer, M., Salomon, R., Goldberg, I., Blanke, O. & Arzy, S. Brain system for mental orientation in space, time, and person. *Proc. Natl. Acad. Sci. USA* **112**, 11072–11077 (2015).
60. Schurz, M., Radua, J., Aichhorn, M., Richlan, F. & Perner, J. Fractionating theory of mind: a meta-analysis of functional brain imaging studies. *Neurosci. Biobehav. Rev.* **42**, 9–34 (2014).
61. Potvin, S., Gamache, L. & Lungu, O. A functional neuroimaging meta-analysis of self-related processing in schizophrenia. *Front. Neurol.* **10**, 990 (2019).
62. Crone, E. A. & Fuligni, A. J. Self and others in adolescence. *Annu. Rev. Psychol.* **71**, 447–469 (2020).
63. Fehlbaum, L. V., Borbás, R., Paul, K., Eickhoff, S. B. & Raschle, N. M. Early and late neural correlates of mentalizing: ALE meta-analyses in adults, children and adolescents. *Soc. Cogn. Affect. Neurosci.* **17**, 351–366 (2022).
64. Spunt, R. P. & Adolphs, R. Folk explanations of behavior: a specialized use of a domain-general mechanism. *Psychol. Sci.* **26**, 724–736 (2015).
65. Spunt, R. P. & Adolphs, R. A new look at domain specificity: insights from social neuroscience. *Nat. Rev. Neurosci.* **18**, 559–567 (2017).
66. Binney, R. J. & Ramsey, R. Social semantics: the role of conceptual knowledge and cognitive control in a neurobiological model of the social brain. *Neurosci. Biobehav. Rev.* **112**, 28–38 (2020).
67. Andrews-Hanna, J. R. & Grilli, M. D. Mapping the imaginative mind: charting new paths forward. *Curr. Dir. Psychol. Sci.* **30**, 82–89 (2021).
68. Baetens, K., Ma, N., Steen, J. & Van Overwalle, F. Involvement of the mentalizing network in social and non-social high construal. *Soc. Cogn. Affect. Neurosci.* **9**, 817–824 (2014).
69. Andrews-Hanna, J. R., Smallwood, J. & Spreng, R. N. The default network and self-generated thought: component processes, dynamic control, and clinical relevance. *Ann. N. Y. Acad. Sci.* **1316**, 29–52 (2014).
70. Yeshurun, Y., Nguyen, M. & Hasson, U. The default mode network: where the idiosyncratic self meets the shared social world. *Nat. Rev. Neurosci.* **22**, 181–192 (2021).
71. Defendini, A. & Jenkins, A. C. Dissociating neural sensitivity to target identity and mental state content type during inferences about other minds. *Soc. Neurosci.* **18**, 103–121 (2023).
72. Kim, H. J., Lux, B. K., Lee, E., Finn, E. S. & Woo, C.-W. Brain decoding of spontaneous thought: predictive modeling of self-relevance and valence using personal narratives. *Proc. Natl. Acad. Sci. USA* **121**, e2401959121 (2024).
73. Kim Lux, B., Andrews-Hanna, J. R., Han, J., Lee, E. & Woo, C.-W. When self comes to a wandering mind: brain representations and dynamics of self-generated concepts in spontaneous thought. *Sci. Adv.* **8**, eabn8616 (2022).
74. Courtney, A. L. & Meyer, M. L. Self-other representation in the social brain reflects social connection. *J. Neurosci.* **40**, 5616–5627 (2020).
75. Debbané, M. et al. Brain activity underlying negative self- and other-perception in adolescents: the role of attachment-derived self-representations. *Cogn. Affect. Behav. Neurosci.* **17**, 554–576 (2017).
76. van Buuren, M. et al. Neural correlates of self- and other-referential processing in young adolescents and the effects of testosterone and peer similarity. *NeuroImage* **219**, 117060 (2020).
77. Fuentes-Claramonte, P. et al. Shared and differential default-mode related patterns of activity in an autobiographical, a self-referential and an attentional task. *PLoS ONE* **14**, e0209376 (2019).
78. Fuentes-Claramonte, P. et al. Brain imaging correlates of self- and other-reflection in schizophrenia. *NeuroImage Clin.* **25**, 102134 (2020).
79. First, M. B. Structured clinical interview for the DSM (SCID). in (eds Cautin, R. L. & Lilienfeld, S. O.) *The Encyclopedia of Clinical Psychology*, 1–6 (John Wiley & Sons, Ltd, 2015).
80. Zhang, L., Opmeer, E. M., Ruhé, H. G., Aleman, A. & van der Meer, L. Brain activation during self- and other-reflection in bipolar disorder with a history of psychosis: comparison to schizophrenia. *NeuroImage Clin.* **8**, 202–209 (2015).
81. Sheehan, D. V. et al. The Mini-International Neuropsychiatric Interview (M.I.N.I.): the development and validation of a structured diagnostic psychiatric interview for DSM-IV and ICD-10. *J. Clin. Psychiatry* **59**, 22–33 (1998).
82. Kemp, R. & David, A. Insight and compliance. in (ed. Blackwell, B.) *Treatment Compliance and the Therapeutic Alliance*, 61–84 (Harwood Academic Publishers, 1997).
83. Beck, A. T., Baruch, E., Balter, J. M., Steer, R. A. & Warman, D. M. A new instrument for measuring insight: the Beck Cognitive Insight Scale. *Schizophr. Res.* **68**, 319–329 (2004).
84. Tusche, A., Spunt, R. P., Paul, L. K., Tyszka, J. M. & Adolphs, R. Neural signatures of social inferences predict the number of real-life social contacts and autism severity. *Nat. Commun.* **14**, 4399 (2023).
85. Ma, S. et al. When empathy gets tough: neural responses to over-coming the self in a novel paradigm predict everyday prosocial behavior, <https://doi.org/10.31234/osf.io/qc45> (2024).

86. Cox, D. D. & Savoy, R. L. Functional magnetic resonance imaging (fMRI) “brain reading”: detecting and classifying distributed patterns of fMRI activity in human visual cortex. *NeuroImage* **19**, 261–270 (2003).
87. Misaki, M., Kim, Y., Bandettini, P. A. & Kriegeskorte, N. Comparison of multivariate classifiers and response normalizations for pattern-information fMRI. *NeuroImage* **53**, 103–118 (2010).
88. Mourão-Miranda, J., Bokde, A. L. W., Born, C., Hampel, H. & Stetter, M. Classifying brain states and determining the discriminating activation patterns: support vector machine on functional MRI data. *NeuroImage* **28**, 980–995 (2005).
89. Woo, C.-W. et al. Separate neural representations for physical pain and social rejection. *Nat. Commun.* **5**, 5380 (2014).
90. Steardo, L. et al. Application of support vector machine on fMRI data as biomarkers in schizophrenia diagnosis: a systematic review. *Front. Psychiatry* **11**, 588 (2020).
91. Karvelis, P. violin and raincloud plots, (<https://github.com/povilaskarvelis/DataViz>), GitHub.
92. R Core Team. *R: A Language and Environment for Statistical Computing* (R Foundation for Statistical Computing, 2024).
93. Wickham, H. *ggplot2: Elegant Graphics for Data Analysis*. (Springer-Verlag, 2016).
94. Açıl, D. et al. Data for: brain neuromarkers predict self- and other-related mentalizing across adult, clinical, and developmental samples. figshare <https://doi.org/10.6084/m9.figshare.29908139.v1> (2026).
95. Açıl, D. et al. Software for: brain neuromarkers predict self- and other-related mentalizing across adult, clinical, and developmental samples. <https://doi.org/10.5281/zenodo.19461625> (2026).
96. Ma, D. S., Correll, J. & Wittenbrink, B. The Chicago face database: a free stimulus set of faces and norming data. *Behav. Res. Methods* **47**, 1122–1135 (2015).

Acknowledgements

This work was funded by a Starting Grant from the European Research Council (ERC, 101041087) to L.Ko.; a German Academic Exchange Service (DAAD) doctoral grant and a Network of European Neuroscience Schools (NENS) exchange fellowship to D.A.; an RO1 grant from the U.S. National Institutes of Mental Health (RO1MH125414-01) to J.A.H. and D.A.S.; a Junior Leader Fellowship from “la Caixa” Foundation (LCF/BQ/PR22/11920017) to P.F.C.; a Consolidator Grant from the European Research Council (ERC, 648082) to L.Kr.; R37, RO1 support from the U.S. National Institutes of Mental Health (R37MH076136 to T.D.W., MH116026 to T.D.W. and L. Chang [PI], RO1EB026549 to T.D.W. and M. Lindquist [MPIs]); an NIMH grant (P50MH094258-01A1) to R. Adolphs; and a CIC Brain and Mental Health Chair from the Neurodis foundation to A.T. L.v.d.M. acknowledges a European Science Foundation EURL grant (044035001) that funded her doctoral studies (PI: A. Aleman). Views and opinions expressed are, however, those of the authors only and do not necessarily reflect those of the European Union, the European Research Council or any other institution. Neither the

European Union nor any granting authority can be held responsible for them. The funders had no role in study design, data analysis, manuscript preparation, or publication decisions.

Author contributions

L.Ko., T.D.W., J.A.H. and M.L.S. conceived the project. D.A. and L.Ko. conceptualized the research goals, built the collaborations, analyzed the data, and prepared the first draft of the manuscript. D.A. created the figures. L.Ko., J.A.H., D.A.S., A.M.C., M.v.B., L.Kr., L.Z., L.v.d.M., P.F.C., E.P.C., R.S., M.D., P.Vr., A.T. and P.Vu. designed the experiments and provided data. L.Ko., D.A. and L.O.W. acquired funding for the project. L.Ko. supervised the project. All authors contributed to the writing and editing of the paper.

Competing interests

The authors declare no competing interests.

Additional information

Supplementary information The online version contains Supplementary material available at <https://doi.org/10.1038/s41467-026-73945-w>.

Correspondence and requests for materials should be addressed to Leonie Koban.

Peer review information *Nature Communications* thanks José Paulo Marques dos Santos, who co-reviewed with Jose Diogo Marques dos Santos, and the other anonymous reviewer(s) for their contribution to the peer review of this work. A peer review file is available.

Reprints and permissions information is available at <http://www.nature.com/reprints>

Publisher’s note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article’s Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article’s Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

© The Author(s) 2026

¹Max Planck Institute for Human Cognitive and Brain Sciences, Leipzig, Germany. ²Department of Child and Adolescent Psychiatry, Psychotherapy, and Psychosomatics, Leipzig University, Leipzig, Germany. ³Department of Clinical Child and Adolescent Psychology and Psychotherapy, University of Bremen, Bremen, Germany. ⁴Department of Psychology, University of Arizona, Tucson, AZ, USA. ⁵Cognitive Science, University of Arizona, Tucson, AZ, USA. ⁶Department of Medicine, School of Medicine and Health Sciences, Institute of Neurosciences, University of Barcelona, Barcelona, Spain. ⁷Institut d’Investigacions Biomèdiques August Pi i Sunyer (IDIBAPS), Barcelona, Spain. ⁸Department of Clinical, Neuro and Developmental Psychology, Faculty of Behavioral and Movement Sciences, Institute for Brain and Behavior Amsterdam, Vrije Universiteit Amsterdam, Amsterdam, The Netherlands. ⁹Institute for Medical Imaging Technology, Ruijin Hospital, Shanghai Jiao Tong University School of Medicine, Shanghai, China. ¹⁰Department of Clinical and Developmental Neuropsychology, University of Groningen, Groningen, The Netherlands. ¹¹Department of Psychiatric Rehabilitation, Lentis Zuidlaren, Zuidlaren, The Netherlands. ¹²FIDMAG Germanes Hospitalàries Research Foundation, Barcelona, Spain. ¹³Centro de Investigación Biomédica en Red de Salud Mental (CIBERSAM) ISCIII, Barcelona, Spain. ¹⁴Developmental Clinical Psychology Research Unit, Faculty of Psychology and Educational Sciences, University of Geneva,

Geneva, Switzerland. ¹⁵Research Department of Clinical, Educational and Health Psychology, University College London, London, UK. ¹⁶Social Neuroscience of Human Attachment (SoNeAt) Lab, Department of Psychology, University of Essex, Colchester, UK. ¹⁷Laboratory of Behavioural Neurology and Imaging of Cognition, Department of Neuroscience, University Medical Center, University of Geneva, Geneva, Switzerland. ¹⁸Swiss Center for Affective Sciences, University of Geneva, Geneva, Switzerland. ¹⁹Department of Psychology, Queen's University, Kingston, ON, Canada. ²⁰Center for Neuroscience Studies, Queen's University, Kingston, ON, Canada. ²¹Department of Psychological and Brain Sciences, Dartmouth College, Hanover, NH, USA. ²²Lyon Neuroscience Research Center (CRNL), CNRS, Inserm, Université Claude Bernard Lyon 1, Bron, France. ²³Le Vinatier Psychiatrie Universitaire Lyon Métropole, Bron, France. ✉e-mail: leonie.koban@cnrs.fr