

Validating a brief measure of four facets of social evaluation

Alex Koch & Austin Smith, University of Chicago, USA

Susan T. Fiske, Princeton University, USA

Andrea E. Abele, University of Erlangen, Germany

Naomi Ellemers, Utrecht University, Netherlands

Vincent Yzerbyt, University of Louvain, Belgium

Authors' note: We pre-registered Studies 4 ([link](#)) and 5a-5d ([link](#), [link](#), [link](#), and [link](#), respectively). In each study, we collected all data before analyzing any of them. All study materials and data, code, and results are available on the Open Science Framework website ([link](#)). Please address correspondence to Alex Koch (alex.koch@chicagobooth.edu).

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

Abstract

Five studies ($N = 7,972$) validated a brief measure and model of four facets of social evaluation (friendliness and morality as horizontal facets; ability and assertiveness as vertical facets). Perceivers expressed their personal impressions or estimated society's impression of different types of targets (i.e., envisioned or encountered groups or individuals) and numbers of targets (i.e., between six and 100) in the separate, items-within-target mode or the joint, targets-within-item mode. Factor analyses confirmed that a two-items-per-facet measure fits the data well and better than a four-items-per-dimension measure that captured the Big Two model (i.e., no facets, just the horizontal and vertical dimensions). As predicted, the correlation between the two horizontal facets and between the two vertical facets was higher than the correlations between any horizontal facet and any vertical facet. Perceivers' evaluations of targets on each facet were predictors of unique and relevant behavior intentions. Perceiving a target as more friendly, moral, able, and assertive increased the likelihood to rely on the target's loyalty, fairness, intellect, and hubris in an economic game, respectively. These results establish the external, internal, convergent, discriminant, and predictive validity of the brief measure and model of four facets of social evaluation.

Keywords: Big Two; four facets; brief measure; validity; adversarial collaboration

Validating a brief measure of four facets of social evaluation

As social beings, people evaluate themselves and others to create opportunities and solve problems. That is, they notice and infer people's attributes, to guide behavior. Ideally, they learn to precisely and efficiently evaluate a few attributes that predict many people's cognition, affect, and behavior across time and many situations. Thus, these summarizing attributes matter. Decades of research have converged on two attribute dimensions. These Big Two have been labelled agency and communion (Abele & Wojciszke, 2014) or competence and warmth (Fiske et al., 2002), for example. They capture evaluations of people's prospect to get ahead in task performance and get along with others, respectively.

Recently, researchers representing five Big Two models set out to compare their findings on social evaluation in the context of an adversarial collaboration (Ellemers et al., 2020). Integrating theoretical predictions and available evidence allowed specifying consensus as well as controversies in need of a resolution (Abele et al., 2021; Koch et al., 2021). A key insight from this adversarial collaboration was that existing research focused on different contexts, modes, and types of social evaluation, and different types and numbers of targets and perceivers. Future research should test whether these differences explain the heterogeneity in previous findings, and thereby resolve controversies and integrate theorizing about social evaluation. This future research requires a validated and agreed-upon measure of the Big Two.

The adversarial collaborators distinguished between two facets of vertical evaluation (ability and assertiveness; see Carrier et al., 2014) and two facets of horizontal evaluation (friendliness and morality; see Brambilla et al., 2012; Leach et al., 2007). This consensual differentiation of each Big Two dimension into two facets (see also Abele et al., 2008; 2016) is worthwhile. People describe groups based on their morality before friendliness, and

based on their ability before assertiveness (Gligorić et al. 2023, Nicolas et al., 2022).

Self-rated friendliness predicts life satisfaction better than self-rated morality, and self-rated assertiveness predicts self-efficacy and self-esteem better than self-rated ability (Abele, 2022; Abele & Hauke, 2020). Then again, impressions of another individual’s ability and morality better predict their esteem / reputation than impressions of their assertiveness and friendliness (Abele & Hauke, 2020).

The present research validated a brief measure of these facets of social evaluation. The measure captures the horizontal friendliness facet with the items “warm” and “friendly,” the horizontal morality facet with the items “honest” and “sincere,” the vertical ability facet with the items “capable” and “skilled,” and the vertical assertiveness facet with the items “confident” and “determined.” The brief measure may prove useful for developing theory about social evaluation and comparing findings across different paradigms and labs.

Building on existing research

Three papers (two published, one under review) made a laudable effort to validate a measure of the facet model (Abele et al., 2016; Abele & Hauke, 2019; Barbedor et al., 2024). Our validation complements this previous and ongoing research in important ways that we discuss below and in no particular order (i.e., they are equally important).

First, two of the papers (Abele et al., 2016; Barbedor et al., 2024) factor-analyzed the ratings for one type of target by different perceivers. This target-centered analysis asks “if one target is rated as more friendly by one (vs. another) perceiver, is that same target rated as more moral by the first (vs. second) perceiver?”, for example. We factor-analyze the ratings for different targets by one perceiver. This perceiver-centered analysis asks “if one perceiver rates one (vs. another) target as more friendly, does that same perceiver rate the first (vs. second) target as more moral?” Some models of social evaluation prefer the target-centered analysis (e.g., Abele & Wojciszke, 2014; Ellemers, 2017), whereas other models prefer the

perceiver-centered analysis (Fiske, 2018; Koch et al., 2016; Yzerbyt, 2018). It matters to generalize the validity of the facet model from the first to the second approach.

Second, all three papers validated an exhaustive five-items-per-facet measure (Abele et al., 2016; Barbedor et al., 2024) or four-items-per-facet measure (Abele & Hauke, 2019). We validate a more parsimonious two-items-per-facet measure because a five-items-per-facet measure is not always feasible in studies in which perceivers evaluate several or even many targets on all four facets as well as upstream and/or downstream variables. For example, evaluating ten targets on 20 items plus ten items that measure one upstream and one downstream construct makes 300 items and takes roughly 30 minutes (except if the study subsamples items for each participant). 30 minutes is a study duration that arguably erodes data quality and costs the researcher(s) an excessive amount of money if they aim to detect a small effect and thus need to compensate many participants.

Third, the three papers (Abele et al., 2016; Abele & Hauke, 2019; Barbedor et al., 2024) validated the facet model when perceivers rated one target per survey page on all items. Some theorizing focuses on this separate mode of social evaluation (e.g., Abele & Wojciszke, 2014; Nicolas et al., 2022), which prioritizes depth. Other theorizing focuses on a joint mode of social evaluation (i.e., rating all targets on one item per survey page; e.g., Koch, Dorrough, et al., 2020; Imhoff et al., 2018; Judd et al., 2019), which prioritizes breadth. Separate versus joint evaluation can reverse preferences and have other interesting effects (Bazerman et al., 1999; Hsee et al., 1999), but we validate our brief measure of the facet model for both modes.

Fourth, the three papers (Abele et al., 2016; Abele & Hauke, 2019; Barbedor et al., 2024) validated the facet model when perceivers expressed personal evaluations of targets. We generalize the validity of the brief measure and facet model from personal evaluations to personal estimates of cultural evaluations (i.e., the perceivers' estimates of how people in

society evaluate the targets, on average) and cultural evaluations proper (i.e., the perceivers' evaluations of the targets, on average). The size and sign of the statistical relation between two social-evaluative dimensions/facets can vary as a function of (dis)aggregating ratings for targets (Imhoff & Koch, 2017; Koch, Imhoff, et al., 2020; Oliveira et al., 2020; Stoller et al., 2018). Thus, a measure of the facet model needs to be validated for both the personal and cultural type of social evaluations.

Fifth, ideally the facet model applies to perceivers' evaluations of various types of targets. Two of the three papers validated the facet model across evaluations of the self (Abele et al., 2016) and other individuals, including close friends, acquaintances, and celebrities (Abele & Hauke, 2019). We replicate this and generalize the facet model to evaluations of large and society-representative samples of groups (i.e., social categories based on gender, age, race, status, beliefs etc., for a description of how the groups were selected, see Koch et al., 2016). We note that the third paper validated the facet model across evaluations of eight social categories (pre-selected from a study by Koch et al., 2016) plus intimacy and task groups (Barbedor et al., 2024), which differ from social categories (e.g., the entitativity of intimacy/task groups is higher; Lickel et al., 2000). Further, the three papers validated the facet model across perceivers' ratings for familiar and labeled targets (e.g., "think of an acquaintance of yours"). We also validated perceivers' ratings for unknown targets that they encountered by seeing a photo of them.

Sixth, the three papers validated the facet model by predicting perceivers' evaluations of targets on dimensions that are upstream or downstream of the perceivers' evaluations of the targets on the four facets (i.e., life satisfaction, self-efficacy, entitativity, and similarity/identification/likeability; Abele et al., 2016; Abele & Hauke, 2019; Barbedor et al., 2024). In addition, we validate the facet model by predicting perceivers' behavior intentions towards the targets that they evaluated. We note that ongoing research aims to validate the

FOUR FACETS OF SOCIAL EVALUATION

7

Big Two (i.e., not the facet model) by predicting perceivers' incentivized behavior towards targets (Walsh et al., 2023).

Table 0 shows the several ways in which our five studies validated the brief measure and facet model as we had suggested when we had consensually endorsed the facet model (Abele et al., 2021; Ellemers, 2017; Koch et al., 2021).

Table 0

Overview of the validation of the brief measure and facet model

Study	N of Perceivers	Type of Evaluation	N & Type of Targets	N of Items per Facet	Mode of Evaluation	Type of Validity
1a 1b	4,007	Personal & Cultural	Labels of 20 Societal Groups	2	Separate & Joint	Internal, Convergent & Discriminant, External
2	1,502	Personal & Cultural	Labels of Self, Friend, Acquaint., Celebrity	2	Separate & Joint	Internal, Convergent & Discriminant, External
3	1,054	Cultural	Photos of 1k Strangers	2	Joint	Internal, Convergent & Discriminant, External
4	399	Cultural	16 Labels of High-/Low-Scorers on 8 Facet Items	2	Joint	Predictive (Sensitivity & Specificity)
5a-5d	1,225	Cultural	Photos of 1k Strangers	2	Joint	Predictive (Sensitivity & Specificity)

Note. Personal vs. Cultural = impressions by individual vs. aggregated perceivers; Separate vs. Joint = impressions of one target (i.e., focus on depth) vs. many targets (i.e., focus on breadth) per survey page; Internal Validity: The simple-structured model with four correlated facets measured with two items each fits the data well enough and better than alternative models; Convergent & Discriminant Validity: The facets correlate more strongly within (vs. between) the Big Two dimensions; External Validity: The brief measure and facet model generalize across different types and numbers of targets and types of evaluation (i.e., personal vs. cultural) and modes of evaluation (i.e., separate vs. joint); Predictive Validity (Sensitivity & Specificity): Each facet predicts some behavior intentions and one behavior intention better than all other facets, with impressions of a target's friendliness, morality, ability, and assertiveness predicting that the perceiver relies on their loyalty, takes their advice, invests in their performance, and exploits their hubris, respectively.

Overview of the validation

Internal validity. In three studies, we modeled perceivers’ impressions of targets on eight items. We selected the eight items from a list of 20 items, see Supplemental Study 1 (perceivers evaluated groups) and Supplemental Study 2 (perceivers evaluated individuals). Each manifest (i.e., measured) item loaded onto one latent (i.e., estimated) facet. Thus, we tested our assumption that the cause of the variance in each item was variance in one facet and no other facet. We estimated four facets from two measured items each, and we estimated correlations between the facets. This simple-structured model with four facets fit the data well enough, according to standard cutoffs (Hu & Bentler, 1999). Importantly, the model fits the data better than a simple-structured model with two correlated dimensions and four items per dimension (i.e., the Big Two). These eight items were the same as in the better and satisfactory model.

Convergent and discriminant validity. Friendliness and morality correlated more strongly, compared to the correlations between friendliness and the two vertical facets, and compared to the correlations between morality and the two vertical facets. In addition, ability and assertiveness correlated more strongly, compared to the correlations between ability and the two horizontal facets, and compared to the correlations between assertiveness and the two horizontal facets. This pattern of correlations emerged across the three studies and for both the latent facets that we estimated and the manifest scales that measured the facets (e.g., the items “warm” and “friendly” formed the scale that measured the friendliness facet). This pattern corroborated our theorizing about the Big Two such that friendliness and morality are horizontal facets, whereas ability and assertiveness are vertical facets (Abele et al., 2016; 2021; Koch et al., 2021).

For simplicity and to emphasize the facets over the Big Two, the main text reports and interprets the pattern of correlations between the facets rather than a higher-order model

in which the facets (do not correlate but) load on their theoretically designated dimension of the Big Two (which correlate). The supplement reports the higher-order model, which fit as well as the simpler model of our choice (i.e., the model with four correlated facets), see Tables SS1.1, SS2.1, S1.1-S1.3, S2.1, and S3.1. The only research that compared the fit of the two models also found that they fit equally well (Barbedor et al., 2024).

Ecological and external validity. The targets that the perceivers rated were labels of many groups (e.g., “Christians” and “drug addicts”), the self and labels of a few self-selected individuals (i.e., a close friend, acquaintance, and celebrity; see Abele & Hauke, 2019), or pictures of many individuals as they appeared on social media recently (see Connor et al., 2023; Gallardo et al., 2023). These perceivers and targets were approximately representative of today’s U.S. society. Further, the perceivers prioritized depth by rating all the items within the targets (i.e., one target per survey page; the separate mode of social evaluation), or they prioritized breadth by rating all the targets within the items (i.e., one item per page; the joint mode). In addition, they rated their personal impressions of the targets or cultural (i.e., society’s) impressions of the targets (e.g., Fiske et al., 2002), or we computed cultural impressions by averaging the perceivers’ ratings separately for each target and item.

The internal, convergent, and discriminant validity of the brief measure generalized across these different types and numbers of targets and modes of social evaluation. This established the external and ecological validity of the brief measure and facet model.

Predictive validity (sensitivity and specificity). In two additional studies, we predicted perceivers’ self-rated behavior intentions towards targets in four economic games. Each game captured a different and broadly relevant interpersonal behavior. In each model, the four rivaling predictors were the perceivers’ impressions of the targets on the four facets captured with our brief measure. Each facet was sensitive in the sense that it predicted some type of behavior intention. In addition, each facet was specific in the sense that it predicted

one type of behavior intention better than all three other facets in at least one of the two studies. The four pairs of a facet and the type of behavior intention that it predicted best corroborated our novel theorizing.

Friendliness and loyalty. In the first game, the perceiver decided between relying on unalterable luck (i.e., their fate) and a target who would decide between earning the perceiver a bonus in an act of loyalty or killing the perceiver’s bonus in an act of revenge. Results showed that the perceiver’s impression of a target’s friendliness predicted the perceiver’s reliance on the target (i.e., their loyalty) better than the perceiver’s impressions of the target’s morality, ability, or assertiveness.

Morality and deception. In the second game, the perceiver decided between a fixed bonus and taking the advice of a target who had decided between giving the perceiver honest information in the best interest of the perceiver’s bonus or deceptive information in the best interest of the target’s bonus (Gneezy, 2005). Results showed that the perceiver’s impression of a target’s morality predicted the perceiver’s taking of the target’s advice better than the perceiver’s impressions of the target’s friendliness, ability, or assertiveness.

Ability and investment. In the third game, the perceiver decided between investing a smaller or larger part of their bonus in a target who would earn the perceiver double of what they invested if the target solved an intellectual puzzle correctly. The target would kill the perceiver’s investment if the solution is incorrect. Results showed that the perceiver’s impression of a target’s ability predicted the perceiver’s large investment in the target better than the perceiver’s impressions of the target’s friendliness, morality, or assertiveness.

Assertiveness and hubris. In the fourth game, the perceiver decided whether to make a bold (i.e., higher risk, higher reward/bonus) bet that a target would push their luck very far in a risk-taking game (Lejuez et al., 2002). Results showed that the perceiver’s impression of

a target's assertiveness predicted the perceiver's bold bet on the hubris of the target better than the perceiver's impressions of the target's friendliness, morality, or ability.

Open science and scope of validity

All studies had IRB approval. We pre-registered Studies 4 ([link](#)) and 5a-5d ([link](#), [link](#), [link](#), and [link](#), respectively). In each study, we collected all data before analyzing any of them. All study materials and data, code, and results are available on the Open Science Framework website ([link](#)). We sampled U.S. residents from the online worker platform Prolific Academic whose participation had been approved at a rate of at least 97%. Prolific workers' representativeness of the U.S. population is decent (Douglas et al., 2023; Peer et al., 2017; U.S. Census Bureau, n.d.). We do not generalize our results beyond the U.S. Other research validated the facet model for several European societies as well as China and Australia (Abele et al., 2016; Abele & Hauke, 2019; Barbedor et al., 2024).

Studies 1a and 1b

We aimed to corroborate the internal, convergent, and discriminant validity of a two-items-per-facet measure of social evaluation that we explored in Supplemental Study 1. People rated groups separately or jointly, and they expressed personal evaluations or estimated society's consensual (i.e., cultural) evaluations of the groups. We aimed to generalize the measure across these broadly relevant modes and types of social evaluation, to establish the external validity of the measure and show that it applies across the different social-evaluative contexts that our adversarial models of social evaluation focus on (Abele et al., 2021; Ellemers et al., 2020; Koch et al., 2021).

Method

Participants. Studies 1a and 1b ran at different points in time. People rated the same groups on the same items though, and thus we pooled people's data across Studies 1a and 1b for brevity and conciseness. Across Studies 1a and 1b, we sampled 4,007 people from

Prolific. We excluded 71 people who recommended that we do not analyze their data¹, leaving 3,934 people (49.0% female, 49.0% male, 1.9% other; $M_{age} = 31.62$, $SD = 12.16$).

Stimuli. People rated the 20 groups that other people in previous research had listed most frequently when instructed to list the groups that together form today’s U.S. society (see Koch et al., 2016). The 20 groups were defined by their gender, age, race, status, beliefs etc. We list them in alphabetical order: Asian people, Black people, Christians, conservatives, Democrats, elderly people, gay people, Hispanic people, lesbian people, liberals, middle class people, poor people, Republicans, rich people, students, transgender people, upper class people, White people, women, and working class people. These groups are social categories, according to a categorization by Lickel and colleagues (2000; these authors also mention intimacy groups [e.g., friends], task groups [e.g., a committee], and loose associations [e.g., riders on a bus]; research by Barbedor and colleagues [2024] validates the facet model for intimacy and task groups, in addition to a few social categories).

Procedure. People used 7-point scales that ranged from 1 = “not at all” to 7 = “extremely” to rate each of the 20 groups on 10 items. The items “friendly” and “warm” aimed to measure the perceived friendliness of the groups, “honest” and “sincere” aimed to measure the perceived morality of the groups, “capable” and “skilled” aimed to measure the perceived ability of the groups, and “confident” and “determined” aimed to measure the perceived assertiveness of the groups. Finally, “positive” and “good” measure people’s general evaluation of the groups. The analyses omit the latter two items.

People rated their personal evaluations of the groups or their estimates of society’s consensual (i.e., cultural) evaluations of the groups. In addition, people rated one group on all

¹ Across all studies, we asked participants whether they had paid attention throughout the study and if they recommend that we use their data. We consider this an important step to ensure that we are not capturing responses from inattentive participants. Importantly, the exclusion rate never exceeded 7.5% for any study.

ten items (in random order) before rating another group on the next survey page (i.e., separate evaluations), with the order of the groups being random as well. Alternatively, people rated all groups (in random order) on one item before rating them on another item on the next page (i.e., joint evaluations), with the order of the items being random as well. In the separate mode, people read “to what extent do you think of [group; e.g., Asian people] as [items]?” (personal evaluations) or “to what extent do most Americans think of [group] as [items]?” (cultural evaluations). In the joint mode, people read “to what extent do you think of the following groups as [item; e.g., FRIENDLY]?” (personal evaluations) or “to what extent do most Americans think of the following groups as [item]” (cultural evaluations).

Finally, people provided demographic information, including their gender and age.

Results

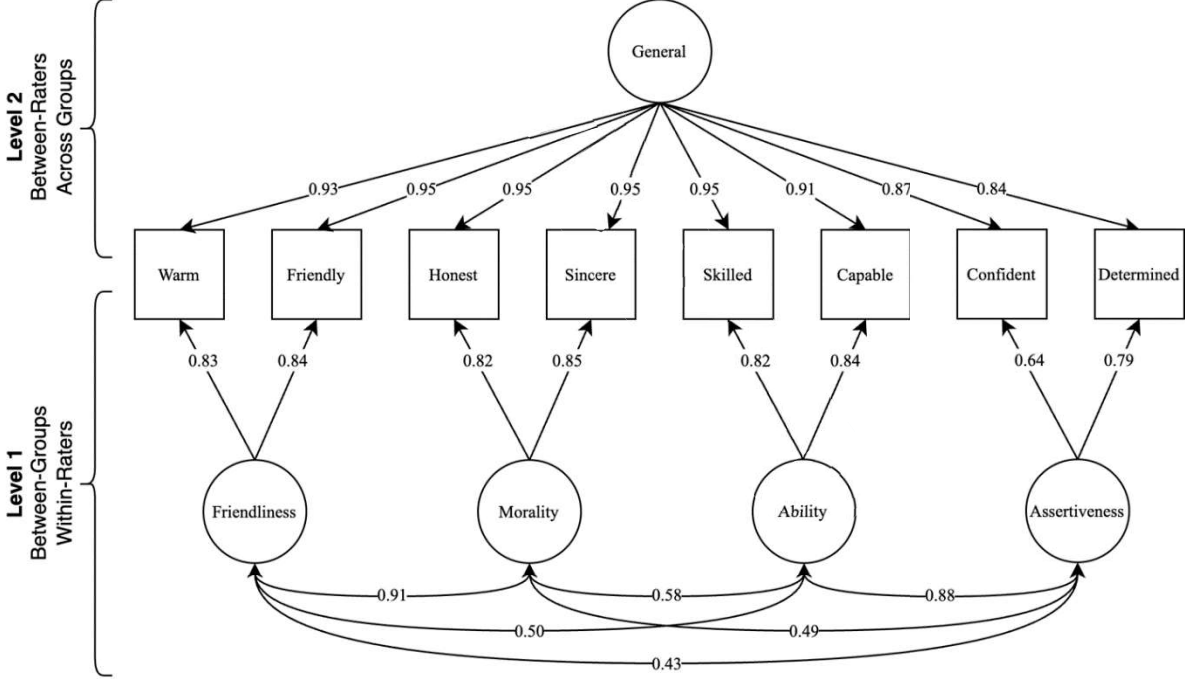
We used the `cfa` function of the R package `lavaan` (Rosseel, 2012) to run four multi-level confirmatory factor analyses (MCFAs; Pornprasertmanit et al., 2014). The four MCFAs modeled all separate evaluations (by 1,955 people), all joint evaluations (by 1,979 people), all personal evaluations (by 1,962 people), or all cultural evaluations (by 1,972 people). In each MCFA, we defined the raters as data clusters (i.e., we treated the groups as nested within the raters), to model the differences between the groups within the raters (vs. modeling the differences between the raters within the groups as in previous work).

Level 1 – our main interest – had the groups as rows, the items as columns, and individual ratings as the data. The manifest items “friendly” and “warm” loaded on the latent friendliness facet, “honest” and “sincere” loaded on the morality facet, “capable” and “skilled” loaded on the ability facet, and “confident” and “determined” loaded on the assertiveness facet. For the sake of internal validity, we allowed no cross-loadings (i.e., we modeled a simple structure), but we allowed the latent facets to correlate with one another.

Level 2 had the raters as rows, the items as columns, and mean ratings across the groups as

the data. All eight manifest items loaded on a latent general factor because we had no hypothesis for the factor structure of level 2, and thus decided to keep it simple. Level 2’s latent general factor captured that some people gave higher ratings to all groups on all items due to acquiescence, complaisance (i.e., wanting to be liked), philanthropy, or a combination of these (Rau et al., 2021). Figure 1 shows the estimated parameters of the measure of cultural evaluations. We visualize these parameters in the main text because the measure of cultural evaluations fit the data best. However, we also visualize the estimated parameters of the measures of personal, separate, and joint evaluations in Figures S1.1-3 in the supplement.

Figure 1
Studies 1a and 1b: Parameters of a two-items-per-facet measure of cultural evaluations



We report multiple fit indices (χ^2 , CFI, RMSEA, SRMR, and AIC) to account for the limitations inherent in relying on a single index (Kline, 2005). We compare the fit of the measure against the standard cutoffs of ≥ 0.95 for CFI and ≤ 0.06 for RMSEA and SRMR (Hu & Bentler, 1999). Table 1.1 shows that the two-items-per-facet measure fit the data well

FOUR FACETS OF SOCIAL EVALUATION

15

in both social-evaluative modes (i.e., separate vs. joint) and types (i.e., personal vs. cultural), except for the SRMR index in the joint mode and personal type.

Table 1.1 also shows the fit of a measure in which the manifest items “honest,” “sincere,” “friendly,” and “warm” loaded on a latent horizontal dimension, and “capable,” “skilled,” “confident,” and “determined” loaded on a vertical dimension that we allowed to correlate with the horizontal dimension. The fit of the four-items-per-dimension measure (which refrained from differentiating the Big Two into the four facets) was worse than the fit of the two-items-per-facet measure in both modes and both types of social evaluation according to nested chi-square difference tests (all p 's $< .001$). Figures S1.4-7 show the estimated parameters of all four four-items-per-dimension measures.

Table 1.1

Studies 1a and 1b: Satisfactory and superior fit of a two-items-per-facet model

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; $df = 34$						
Separate evaluations	1,649	0.99	0.03	0.06	885,536	
Joint evaluations	2,021	0.99	0.04	0.07	935,804	
Personal evaluations	1,938	0.99	0.04	0.08	876,224	
Cultural evaluations	1,042	0.99	0.03	0.05	944,668	
Four-items-per-dimension; $df = 39$						
Separate evaluations	7,923	0.96	0.07	0.08	891,800	$< .001$
Joint evaluations	3,790	0.97	0.05	0.08	937,562	$< .001$
Personal evaluations	5,547	0.97	0.06	0.09	879,823	$< .001$
Cultural evaluations	4,077	0.98	0.05	0.06	947,692	$< .001$

Note. $p(X^2)$ represents the p -value from a nested chi-square difference test compared to the respective two-items-per-facet model.

Next and separately for each mode of social evaluation, we computed a friendliness scale by averaging the two manifest friendliness items separately for each group within each rater. Likewise, we computed a morality, ability, and assertiveness scale. For each scale, we centered the differences between the groups within each rater as in the MCFA. Table 1.2 shows the correlations between the within-centered facet scales in both modes and both types

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

of social evaluation (see Tables S1.3 and S1.4 for the correlations in Study 1a separately from Study 1b). To conclude sufficient convergent and discriminant validity, the correlations between the friendliness and morality (i.e., horizontal) facets and the correlations between the ability and assertiveness (i.e., vertical) facets had to be positive and larger than all correlations between one horizontal facet and one vertical facet. The data confirmed this, except that in the personal mode of social evaluation, the correlation between the morality and ability facets, and between the friendliness and ability facets, was slightly larger than the correlation between the ability and assertiveness facets. However, the correlations between the latent facets in the two-items-per-facet measure confirmed the sought-after pattern in both modes and both types of social evaluation, see Figure 1 and Figures S1.1-3.

FOUR FACETS OF SOCIAL EVALUATION

17

Table 1.2*Studies 1a and 1b: Correlations between the four facet scales centered within the raters*

Separate					Joint				
Friendliness -	1.00	0.79	0.51	0.35	1.00	0.74	0.44	0.24	
Morality -	0.79	1.00	0.55	0.36	0.74	1.00	0.50	0.27	
Ability -	0.51	0.55	1.00	0.63	0.44	0.50	1.00	0.56	
Assertiveness -	0.35	0.36	0.63	1.00	0.24	0.27	0.56	1.00	
Personal					Cultural				
Friendliness -	1.00	0.78	0.57	0.28	1.00	0.75	0.41	0.31	
Morality -	0.78	1.00	0.60	0.27	0.75	1.00	0.48	0.36	
Ability -	0.57	0.60	1.00	0.51	0.41	0.48	1.00	0.66	
Assertiveness -	0.28	0.27	0.51	1.00	0.31	0.36	0.66	1.00	
	Friendliness	Morality	Ability	Assertiveness	Friendliness	Morality	Ability	Assertiveness	

Note. More intense hues indicate larger correlations.

Figures S1.8-23 and Tables S1.2-5 show the results of the two studies when analyzing their people/data separately. The results fully replicate the below pattern of results.

Discussion

Based on goodness of absolute and relative model fit as well as correlations between both manifest scales and latent factors, Studies 1a and 1b largely confirmed the validity of our efficient two-items-per-facet measure of four correlated facets of social evaluation (ability, assertiveness, morality, and friendliness; Abele et al., 2021; Koch et al., 2021), and demonstrated its robustness across two modes and two types of evaluating many groups (i.e., separate or joint evaluations, and personal or cultural evaluations).

Study 2

We aimed to generalize our two-items-per-facet measure across separate, joint, personal, and cultural evaluations of several individuals per rater, to further corroborate the external validity of the measure. Put differently, Study 2 aimed to replicate the results of Studies 1a and 1b, except that people evaluated several individuals as in previous work (Abele & Hauke, 2019; see also Supplemental Study 2), instead of many groups.

Method

Participants. We sampled 1,502 people from Prolific. We excluded 54 people who recommended that we do not analyze their data, leaving 1,448 people (49.2% female, 49.0% male, 1.7% other; $M_{age} = 31.46$, $SD = 11.55$).

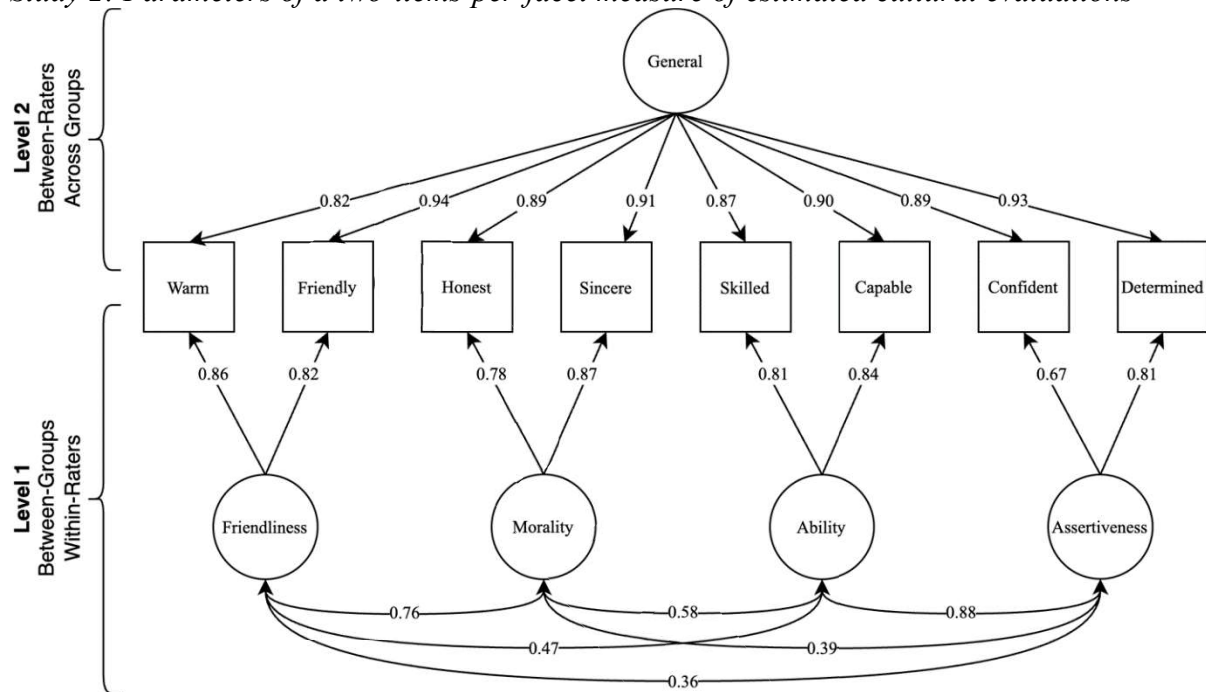
Stimuli and procedure. People began by typing into text boxes the names of a same-sex close friend, acquaintance, and celebrity. These targets of social evaluation differ in terms of their closeness to the perceiver and positivity in the eyes of the perceiver, two continua that characterize the focus of many, if not all, models of social evaluation (Abele et al., 2021; Koch et al., 2021). As in Studies 1a and 1b, they separately or jointly evaluated these individuals, themselves, and these groups on the same 10 items as in Studies 1a and 1b. In addition, people rated their personal evaluations of the social entities or their estimates of society’s consensual (i.e., cultural) evaluations of the social entities.

At the end of the study, people provided demographic information, including their gender and age.

Results

Figure 2

Study 2: Parameters of a two-items-per-facet measure of estimated cultural evaluations



We used the `cfa` function of the R package `lavaan` (Rosseel, 2012) to run four multi-level confirmatory factor analyses (MCFAs; Pornprasertmanit et al., 2014) on people's evaluations of the four individuals. The four MCFAs modeled all separate evaluations (by 729 people), all joint evaluations (by 719 people), all personal evaluations (by 723 people), or all cultural evaluations (by 725 people). In each MCFA, we defined the raters as data clusters (i.e., we treated the individuals as nested within the raters), to model the differences between the individuals within the raters. We specified Levels 1 and 2 in the same way as for the two-items-per-facet measures in Studies 1a and 1b. Figure 2 shows the estimated parameters

of the measure of cultural evaluations. Figures S2.1-3 show the estimated parameters of the other three measures (i.e., separate, joint, and personal evaluations).

We report multiple fit indices, and we compare the fit of the measure against the standard cutoffs of ≥ 0.95 for CFI and ≤ 0.06 for RMSEA and SRMR (Hu & Bentler, 1999). Table 2.1 shows that the two-items-per-facet measure fit the data well in both modes of social evaluation (i.e., separate and joint) and both types of social evaluation (i.e., personal and cultural), except for the SRMR index.

Table 2.1
Study 2: Satisfactory and superior fit of a two-items-per-facet model

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; df = 34						
Separate evaluations	216	0.99	0.04	0.08	65,349	
Joint evaluations	204	0.99	0.04	0.09	65,000	
Personal evaluations	248	0.98	0.05	0.09	65,150	
Cultural evaluations	164	0.99	0.04	0.07	65,306	
Four-items-per-dimension; df = 39						
Separate evaluations	1,005	0.92	0.09	0.12	66,127	<.001
Joint evaluations	779	0.93	0.08	0.12	65,565	<.001
Personal evaluations	967	0.92	0.09	0.11	65,858	<.001
Cultural evaluations	834	0.93	0.08	0.11	65,967	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

Table 2.1 also shows that a four-items-per-dimension measure (that we specified as in Studies 1a and 1b) fits the data worse than the hypothesized two-items-per-facet measure in both modes and both types of social evaluation according to nested chi-square difference tests (all p's <.001). Figures S2.4-7 show the estimated parameters for the four worse-fitting four-items-per-dimension measures.

FOUR FACETS OF SOCIAL EVALUATION

21

Table 2.2*Study 2: Correlations between the four facet scales centered within the raters*

	Separate				Joint			
Friendliness -	1.00	0.62	0.34	0.21	1.00	0.64	0.32	0.20
Morality -	0.62	1.00	0.42	0.19	0.64	1.00	0.37	0.14
Ability -	0.34	0.42	1.00	0.62	0.32	0.37	1.00	0.69
Assertiveness -	0.21	0.19	0.62	1.00	0.20	0.14	0.69	1.00
	Personal				Cultural			
Friendliness -	1.00	0.63	0.26	0.14	1.00	0.63	0.39	0.26
Morality -	0.63	1.00	0.32	0.08	0.63	1.00	0.47	0.25
Ability -	0.26	0.32	1.00	0.64	0.39	0.47	1.00	0.67
Assertiveness -	0.14	0.08	0.64	1.00	0.26	0.25	0.67	1.00
	Friendliness	Morality	Ability	Assertiveness	Friendliness	Morality	Ability	Assertiveness

Note. More intense hues indicate larger correlations.

As before, we had hypothesized that the correlation between the friendliness and morality (i.e., horizontal) facets and the correlation between the ability and assertiveness (i.e., vertical) facets would be positive and larger than any correlation between one horizontal facet and one vertical facet. The data confirmed this pattern of correlations in both modes and

both types of social evaluation when we computed within-centered facet scales according to the two-items-per-facet measure, see Table 2.2. The correlations between the latent facets in the four two-items-per-facet measures also confirmed the hypothesized pattern, see Figure 2 and Figures S2.1-3.

Discussion

Based on goodness of absolute and relative model fit as well as correlations between both manifest scales and latent factors, Study 2 largely confirmed the validity our efficient two-items-per-facet measure of four correlated facets of social evaluation (ability, assertiveness, morality, and friendliness; Abele et al., 2021; Koch et al., 2021). We confirmed the internal, convergent, and discriminant validity of the measure across two modes and two types of evaluating several individuals (i.e., separate or joint evaluations, and personal or cultural evaluations). This further corroborated the external validity of the measure.

Studies 3-5 captured cultural evaluations because the measure of cultural (vs. personal, separate, and joint) evaluations fit the data best in Studies 1 and 2

Study 3

We aimed to generalize the two-items-per-facet measure from personal estimates of cultural evaluations (Studies 1a-2) to cultural evaluations proper, and from evaluations of labels of familiar and groups or individuals (Studied 1a-2) to evaluations of unknown individuals that people encountered by seeing a photo of them. The latter generalization matters because people constantly meet strangers (e.g., on the street, at professional and private events, and when browsing social media and news platforms online) that they need to evaluate precisely and swiftly to make opportunities and solve problems.

Method

Participants. We sampled 1,054 people from Prolific. We excluded 69 people who recommended that we do not analyze their data, leaving 985 people (41.5% female, 56.3% male, 1.7% other, 0.4% preferred not to answer; $M_{age} = 41.28$, $SD = 13.09$).

Stimuli. Each person rated a random selection of 100 out of 1,000 individuals based on their public-facing Facebook profile picture in 2021. Most of these pictures provide a glimpse at an individual's real life (e.g., their workplace, social network, habits, hobbies etc.). Thus, seeing someone's Facebook profile picture is a more ecologically valid way of encountering them, compared to a passport photograph. We created the set of 1,000 pictures using a quasi-random procedure. In each of 1,000 trials, we first searched for a randomly selected U.S. city using Facebook's search engine. Second, we selected the first Facebook page that was not the city's page, was located in the U.S., and had at least 300 likes. Third, we selected the profile picture of the individual who liked that page most recently if they were the only (or focal) person in the picture, if their gender, age, and race were discernible, and if they resided in the U.S. as indicated in their profile's "About Info." We coded the gender (woman or man), age (young, middle-aged, or old), and race/ethnicity (White, Black, Latino/a, East Asian, or South Asian) of each picture (62.5% women, 37.5% men; 32.7% young, 51.8% middle-aged, 15.5% old; 81.5% White, 8.9% Black, 5.1% Latino/a, 2.8% East Asian, 1.7% South Asian).

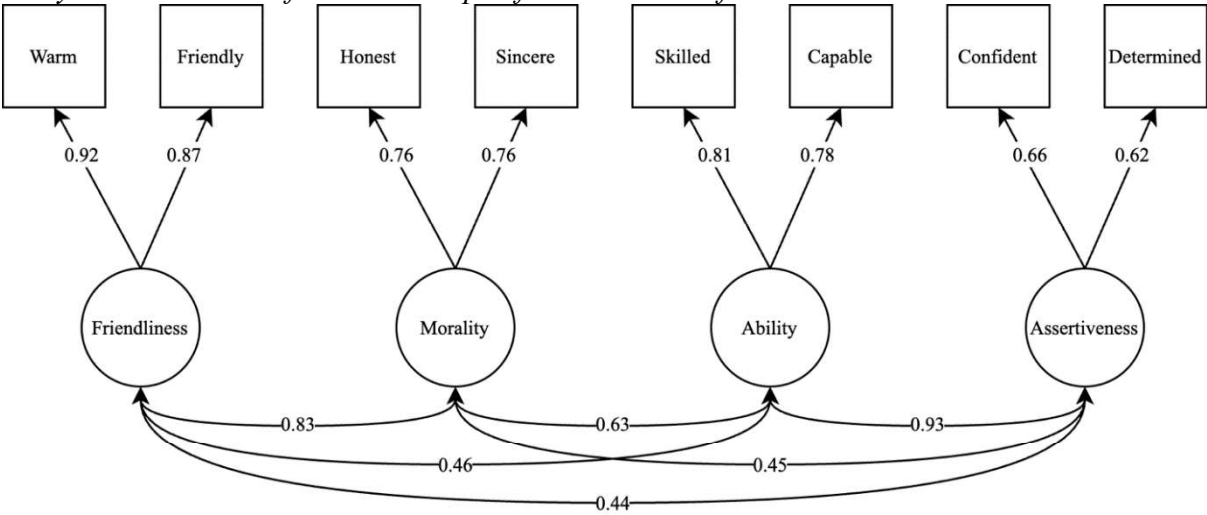
Procedure. On separate survey pages, people saw a picture of an individual and below they read "To what extent do you consider this person to be CAPABLE [or another item]." They used a 7-point scale that ranged from 1 = "not at all" to 7 = "extremely" to rate the individual. They rated 100 individuals in random order and on one and the same item that we randomly selected from the eight items that we examined in Studies 1a-2 plus "positive"

and “good.” At the end of the survey, people provided demographic information including their gender and age.

Results

For each of the 1,000 individuals that people had rated, we computed their mean rating on the eight items that we examined in Studies 1a-2 (e.g., “friendly” and “warm”). Next and based on a matrix that had the individuals as rows, the items as columns, and their mean ratings as the data, we used the cfa function of the R package lavaan (Rosseel, 2012) to run a single-level confirmatory factor analysis (CFA). (It was neither possible nor appropriate to run a multi-level confirmatory factor analysis [MCFA] in Study 3 because Study 3’s data had only one row per individual [vs. several/many rows per group/individual] in Studies 1 and 2 and Supplemental Studies 1 and 2.) The manifest items “friendly” and “warm” loaded on the latent friendliness facet, “honest” and “sincere” loaded on the morality facet, “capable” and “skilled” loaded on the ability facet, and “confident” and “determined” loaded on the assertiveness facet. We allowed no cross-loadings (i.e., we modeled a simple structure), but we allowed the latent facets to correlate with one another. Figure 3 shows the estimated parameters of this two-items-per-facet measure of cultural evaluations.

Figure 3
Study 3: Parameters of a two-items-per-facet measure of cultural evaluations



We report multiple fit indices, and we compare the fit of the measure against the standard cutoffs of ≥ 0.95 for CFI and ≤ 0.06 for RMSEA and SRMR (Hu & Bentler, 1999). Table 3.1 shows that the two-items-per-facet measure fit the data well enough according to the SRMR index. A four-items-per-dimension measure (that we specified as in Studies 1a-2) fit the data worse according to nested chi-square difference tests (all p 's $< .001$). Figure S3.1 shows the estimated parameters of the worse-fitting measure.

Table 3.1

Study 3: Barely satisfactory and superior fit of a two-items-per-facet model

	X^2	CFI	RMSEA	SRMR	AIC	$p(X^2)$
Two-items-per-facet; $df = 14$	230	0.94	0.12	0.06	-12,512	
Four-items-per-dimension; $df = 19$	394	0.90	0.14	0.08	-12,362	$< .001$

Note. $p(X^2)$ represents the p -value from a nested chi-square difference test compared to the respective two-items-per-facet model.

As in Studies 1a-2, we had hypothesized that the correlations between the friendliness and morality (i.e., horizontal) facets and the correlations between the ability and assertiveness (i.e., vertical) facets would be positive and larger than any correlation between one horizontal facet and one vertical facet. The data confirmed this, see Table 3.2. The correlations between the latent facets in the two-items-per-facet measure also confirmed this, see Figure 3.

Table 3.2
Study 3: Correlations between the four facet scales

Two Items per Facet				
Friendliness -	1.00	0.67	0.38	0.30
Morality -	0.67	1.00	0.47	0.28
Ability -	0.38	0.47	1.00	0.62
Assertiveness -	0.30	0.28	0.62	1.00
	Friendliness	Morality	Ability	Assertiveness

Note. More intense hues indicate larger correlations.

Discussion

Based on goodness of absolute and relative model fit as well as correlations between both manifest scales and latent factors, Study 3 largely generalized the internal, convergent, and discriminant validity of the two-items-per-facet measure from personal estimates of cultural evaluations (see Studies 1a-2) to cultural evaluations proper, and from evaluations of labels of familiar groups or individuals (Studies 1a-2) to evaluations of unknown individuals that people encountered by seeing a photo of them.

Study 4

Studies 1a-3 examined items and rating scales that confirmed the facet model’s internal, convergent, and discriminant validity (Abele et al., 2021; Koch et al., 2021) across a variety of to-be-evaluated social entities and two modes and two types of social evaluation.

Studies 4 and 5 aimed to establish the predictive validity of the facet model. Previous and ongoing validations addressed predictive validity by showing that perceivers’

impressions of targets on the four facets predict other impressions, including the targets' life satisfaction, self-efficacy, entitativity, and likeability as seen by the perceivers (Abele et al., 2016; Abele & Hauke, 2019; Barbedor et al., 2024). In Studies 4 and 5, we went beyond these validations by showing that perceivers' impressions of targets on each facet predicted some behavioral intentions of the perceivers towards the targets (i.e., sensitivity). In addition, each facet predicted a unique behavioral intention better than all three other facets in at least one study (i.e., specificity). We measured the behavioral intentions through economic games because they capture people's willingness to act on their impressions (i.e., put their money where their mouth is) in abstracted ways that speak to many real-life situations (Thielmann et al., 2021).

We reasoned that impressions of a target's friendliness, morality, ability, and assertiveness may guide a perceiver to expect loyalty, fairness, performance, and risk-taking from a target, respectively. Accordingly, in four economic games we predicted that the perceivers would decide to make their bonus (for participating in the study) dependent on targets' loyal trickery, fair advice, analytical success, and hubristic gambling, respectively.

Method

Participants. We sampled 399 people from Prolific. We excluded three people who recommended that we do not analyze their data, leaving 396 people (38.9% female, 59.3% male, 0.5% other, 1.3% preferred not to answer; $M_{age} = 41.55$, $SD = 13.25$).

Stimuli. Each person envisioned 16 anonymous individuals that we described in terms of one personality trait. One individual was "a friendly person," whereas another individual was "a person who lacks friendliness." Seven other individuals were "a warm [or honest, sincere, capable, skilled, confident, or determined] person," whereas seven other individuals were "a person who lacks friendliness [or warmth, honesty, sincerity, capability, skill, confidence, or determination]." In sum, each person envisioned one high scorer and one

low scorer on each of the eight items in our brief measure of the four facets, see Studies 1a-3.

Procedure. People played four hypothetical games in random order. Each game was economic (i.e., about winning money) and dyadic, which means that people played each game with each of the 16 anonymous individuals in random order and on the same survey page. In each game played with each individual, people made a choice between two options, see below.

At the end of the survey, people used 7-point scales (1 = “not at all”, 7 = “extremely”) to rate the self on the eight items and “positive” and “good.” Finally, people provided demographic information, including their gender and age.

Loyalty Game. People read “You rolled a 6-sided dice and win \$10 only if you rolled a 5 or 6. As of now, you do not know the outcome of your dice roll. You can publicly reveal the outcome of the dice roll. Or, you can ask each of 16 persons to secretly learn and report the outcome. They can tell the truth or a lie (no one will ever know). If they report that you rolled a 5 or 6, you get \$10. Do you delegate reporting the outcome to each person?” Righteously revealing the outcome was coded as 0, whereas delegating the reporting of the outcome, and thereby expecting loyal trickery, was coded as 1.

Deception Game. People read “Each of 16 persons advises you to choose Option B. They all flipped a coin. If the coin landed on heads, the payoffs of Option B will be \$3 for you and \$7 for them. If the coin landed on tails, the payoffs of Option B will be \$7 for you and \$3 for them. Everyone messages you ‘Tails! Option B pays out \$7 for you and \$3 for me. Pick Option B.’ If you choose Option A, you and them will receive \$5 each. How much do you trust the advice of each person?” Skeptically choosing Option A was coded as 0, whereas expecting true and fair advice, and thus choosing Option B was coded as 1.

Investment Game. People read “Your task is to invest \$1.50 or \$3.50 out of \$5 in each of 16 persons. For each person, if they solve the below problem correctly, you will

FOUR FACETS OF SOCIAL EVALUATION

29

receive double the money you invested plus the money you withheld. If they solve the below problem incorrectly, you will lose the money you invested but receive the money you withheld. They receive \$1.50 regardless of their performance. [A randomized medium-difficulty SAT question was inserted here]. How much do you invest in each person?" Pessimistically investing just \$1.50 was coded as 0, whereas optimistically expecting analytical success and investing \$3.50 was coded as 1.

Hubris Game. Adapted from the Balloon Analogue Risk Task (BART; Lejuez), people read "Each of 16 persons is given a balloon to inflate. They earn \$1 for each time they pump the balloon, but they lose their earnings if the balloon pops. At each turn, they can pump or stop and collect their earnings. Without them knowing, you decide when the balloon pops. If the balloon pops, you keep all their earnings. If the balloon doesn't pop, you earn nothing. You can pop the balloon on the 4th or 6th pump. If they pop the balloon on the 4th pump, you collect \$4 and they collect nothing. If they pop the balloon on the 6th pump, you collect \$6 and they collect nothing. If the balloon does not pop, they keep their earnings. On which pump do you pop the balloon?" Cautiously popping the balloon on the 4th pump already was coded as 0, whereas expecting hubristic gambling, and thus boldly popping the balloon on the 6th pump was coded as 1.

Results

We coded the 16 individuals' high and low scores on the eight items in a way that treated the four facets as orthogonal (i.e., independent). We coded the friendliness of the two individuals who lacked friendliness and warmth as -0.5, who were friendly and warm as 0.5, and we coded the friendliness of the other 12 individuals as 0. We coded the morality of the two individuals who lacked honesty and sincerity as -0.5, who were honest and sincere as 0.5, and we coded the morality of the other 12 individuals as 0. We coded the ability of the two individuals who lacked capability and skill as -0.5, who were capable and skilled as 0.5, and

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

we coded the ability of the other 12 individuals as 0. Finally, we coded the assertiveness of the two individuals who lacked confidence and determination as -0.5, who were confident and determined as 0.5, and we coded the assertiveness of the other 12 individuals as 0.

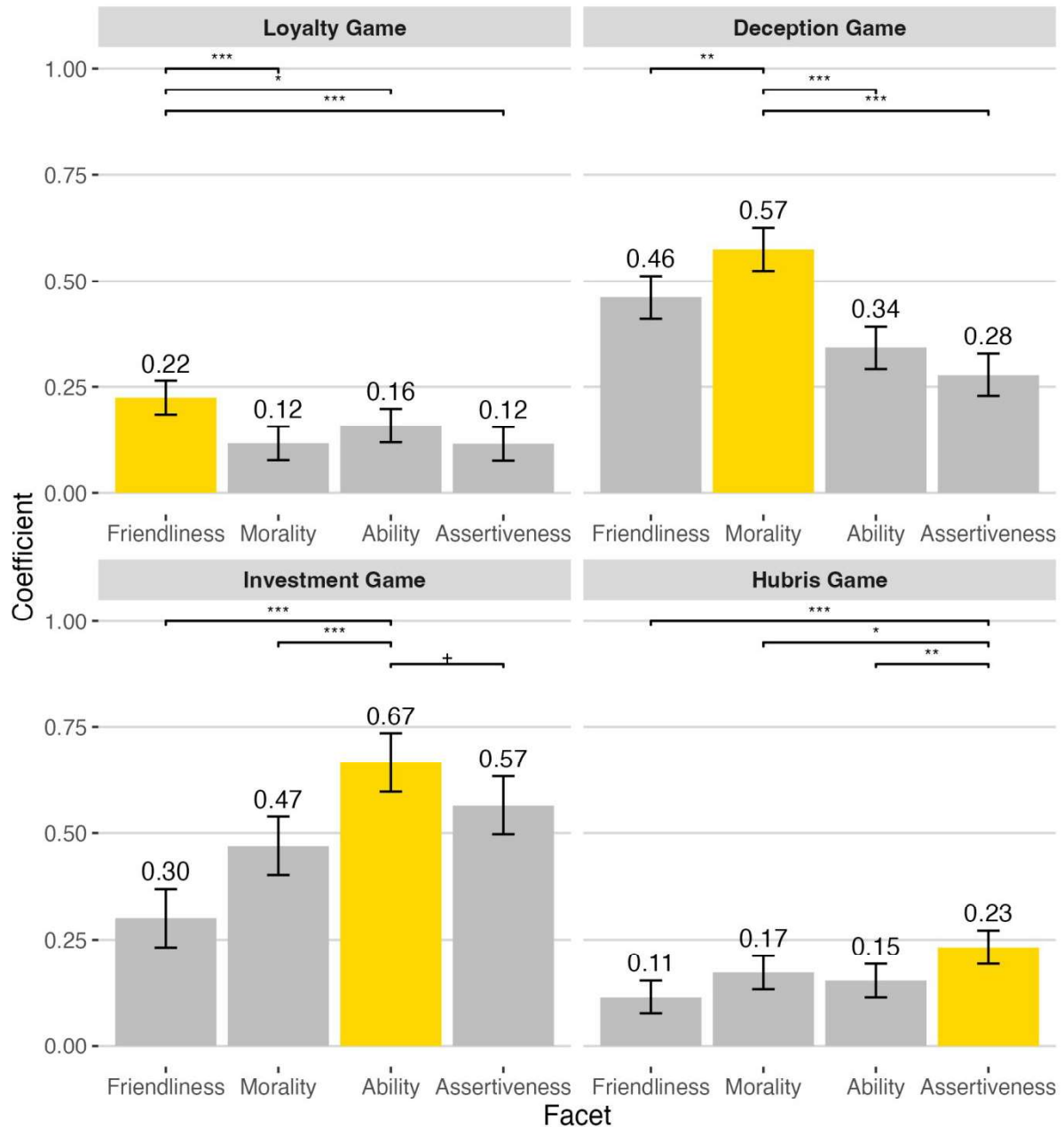
We used the lmer function of the R package lme4 (Bates et al., 2015) to run a linear mixed model with random intercepts for the people and the 16 individuals that they played with. We predicted the people’s binary choices in the loyalty game from the individuals’ orthogonal scores on the four facets. Three additional models were the same except that we predicted people’ binary choice in the deception, investment, and hubris game. We subjected each model to dominance analysis (Azen & Budescu, 2003), which we ran by using the domin function of the R package domir (Lunchman, 2023). Dominance analysis compares regression coefficients based on their sign, size, and scatter.

FOUR FACETS OF SOCIAL EVALUATION

31

Figure 4

Study 4: Expecting loyalty, fairness, success, and hubris from envisioned individuals based on their score on ability, assertiveness, morality, and friendliness



Note. Error bars represent 95% confidence intervals. The highlighted columns represent the strongest predictors of people's binary choice in each game according to dominance analysis. *** means $p < .001$; ** means $p < .010$; * means $p < .050$; + means $p < .100$.

Figure 4 and the pre-registered dominance analysis in Table S4.1 show that in the loyalty game, the friendliness of an individual best predicted that a participant would make their bonus dependent on the individual's loyal trickery. In the deception game, the morality of an individual best predicted that a participant would make their bonus dependent on the

individual’s fair advice. In the investment game, the ability of an individual best predicted that a participant would make their bonus dependent on the individual’s analytical success. In the hubris game, the assertiveness of an individual best predicted that a participant would make their bonus dependent on the individual’s hubristic gambling.

Dominance analysis (Azen & Budescu, 2003) compares regression coefficients descriptively (i.e., no inferential statistics). To substantiate the above interpretations with inferential statistics, we used the linearHypothesis function of the R package car (Fox & Weisberg, 2018). Figure 4 shows in each game, the facet marked in yellow (that we had hypothesized to emerge as the best-predicting facet) actually emerged as the best predictor according to the significance tests reported in Table S4.2. The only exception was the investment game. The ability facet predicted investment better than the assertiveness facet, but this comparison of regression coefficients did not reach statistical significance, $p = .065$.

Discussion

Study 4 validated the two-items-by-facet measure by showing that evaluations on each facet predicted a broadly relevant behavior towards the evaluated individuals better than evaluations on all three other facets (except the ability facet in the investment game). However, we described the individuals solely based on their high or low score on one item/facet (e.g., “is friendly”), and we orthogonalized the individuals’ facet scores. Thus, we take Study 4’s results as a proof-of-concept and acknowledge that the results need to be generalized to real, nuanced social entities whose facet scores vary and covary naturally (i.e., according to the facet model; Abele et al., 2021).

Studies 5a-5d

Studies 5a-5d aimed to validate the facet model by showing that cultural evaluations of real and nuanced individuals on each facet predict a relevant and specific behavior towards

these unknown (vs. familiar in Studies 1a-2) and encountered (vs. envisioned in Studies 1a-4) individuals better than cultural evaluations of the individuals on all three other facets.

Method

Participants. In Studies 5a-5d, we sampled 304, 307, 307, and 304 people from Prolific. We excluded six, three, three, and one people, respectively, who recommended that we do not analyze their data. This left 298 people in Study 5a (43.3% female, 55.0% male, 1.3% other, 0.3% preferred not to answer; $M_{age} = 40.30$, $SD = 13.02$). It left 304 people in Study 5b (44.7% female, 53.0% male, 2.0% other, 0.3% preferred not to answer; $M_{age} = 40.32$, $SD = 12.45$). It left 304 people in Study 5c (50.0% female, 47.4% male, 1.6% other, 1.0% preferred not to answer; $M_{age} = 41.47$, $SD = 13.70$). And it left 303 people in Study 5d (37.6% female, 61.1% male, 1.0% other, 0.3% preferred not to answer; $M_{age} = 40.61$, $SD = 12.04$).

Stimuli. Each person encountered a random selection of 100 out of 1,000 individuals as they appeared in their public-facing Facebook profile picture in 2021. Study 3 describes their gender, age, and race distribution and the quasi-random inclusion of each individual in the set.

Procedure. On each of 100 randomly ordered survey pages, people encountered one picture of an individual. Below, they read the instructions of one hypothetical economic game, namely the loyalty, deception, investment, and hubris game that we described in Study 4. Below, they made the same binary choice as in Study 4. That is, they delegated the reporting of a secret dice roll to the individual (vs. revealed it themselves), trusted the advice of the individual (vs. ignored it), invested a larger (vs. smaller) sum in the individual's analytical success, or bet against the individual boldly (vs. cautiously). So, in Studies 5a-5d we used a between-subjects design (vs. the within-subjects design of Study 4 in which people played all four games with the 16 individuals that they envisioned in that study).

Results

We used the lmer function of the R package lme4 (Bates et al., 2015) to run a linear mixed model with random intercepts for the people and the individuals that they played with. We predicted the people’s binary choices in the loyalty game from the individuals’ correlated scores on the four facets. The individuals’ friendliness scores were their mean ratings on the “friendly” and “warm” items that we computed across all people in Study 3 who had rated one of the two items. The individuals’ morality, ability, and assertiveness scores were their mean ratings on “honest” and “sincere”, “capable” and “skilled”, and “confident” and “determined” across all people in Study 3 who had rated one of the two respective items. In sum, we predicted the binary choices of Study 5a-d’s people from the cultural evaluations that we had captured with the help of Study 3’s people. In three additional models, we predicted people’ binary choice in the deception, investment, and hubris game. We subjected each model to dominance analysis (Azen & Budescu, 2003), which we ran by using the domin function of the R package domir (Lunchman, 2023). Again, dominance analysis compares regression coefficients based on their sign, size, and scatter.

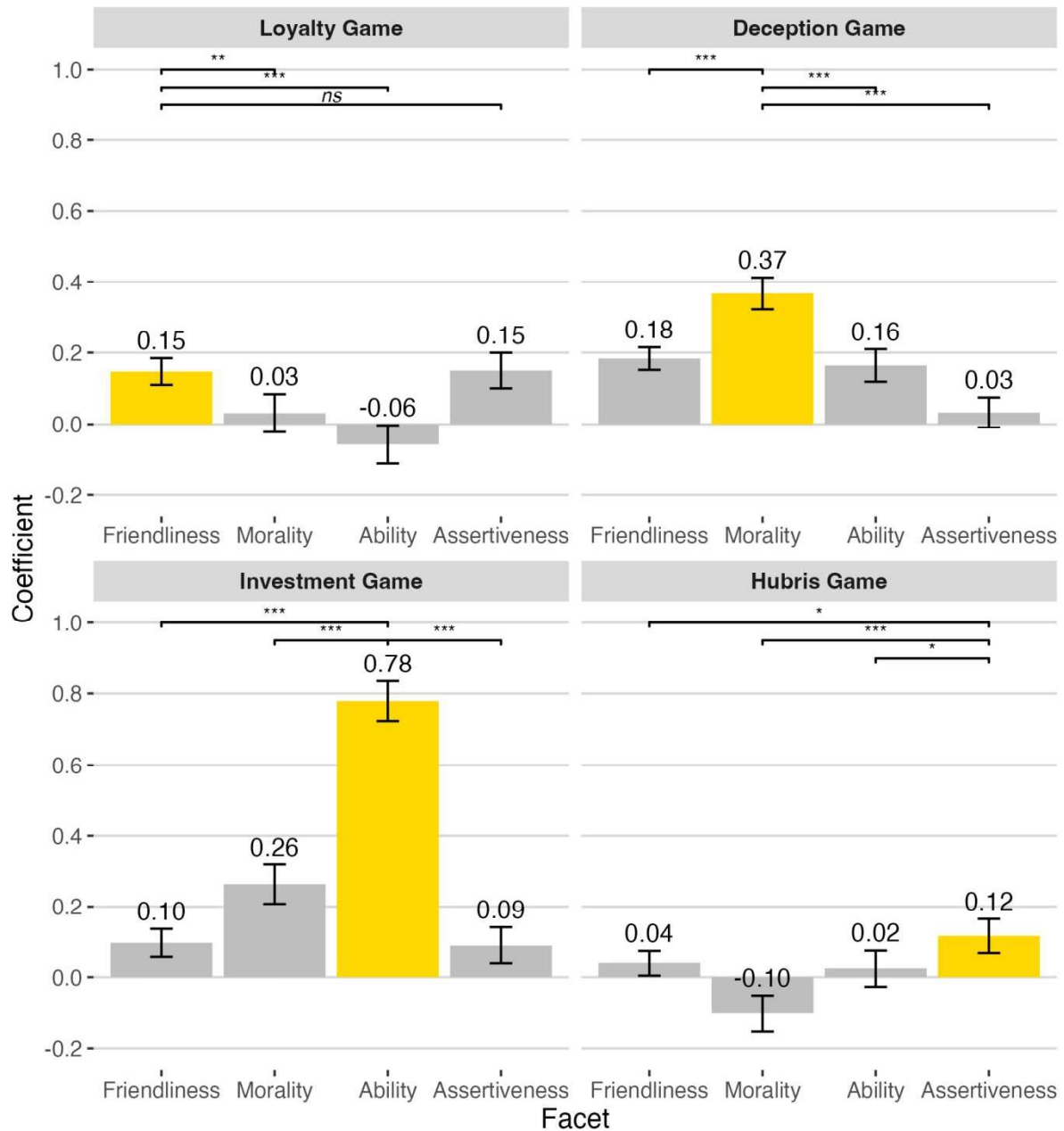
Figure 5 and the pre-registered dominance analysis in Table S5.1 show that in the loyalty game, the friendliness of an individual best predicted that a participant would make their bonus dependent on the individual’s loyal trickery. In the deception game, the morality of an individual best predicted that a participant would make their bonus dependent on the individual’s fair advice. In the investment game, the ability of an individual best predicted that a participant would make their bonus dependent on the individual’s analytical success. In the hubris game, the assertiveness of an individual best predicted that a participant would make their bonus dependent on the individual’s hubristic gambling.

FOUR FACETS OF SOCIAL EVALUATION

35

Figure 5

Studies 5a-5d: Expecting loyalty, fairness, success, and hubris from envisioned individuals based on their score on ability, assertiveness, morality, and friendliness



Note. Error bars represent 95% confidence intervals. Highlighted columns represent the relatively most important predictor of each game according to dominance analysis.

*** means $p < .001$; ** means $p < .010$; * means $p < .050$; + means $p < .100$.

Again, dominance analysis (Azen & Budescu, 2003) compares regression coefficients descriptively (i.e., no inferential statistics). To substantiate the above interpretations with inferential statistics, we used the linearHypothesis function of the R package car (Fox & Weisberg, 2018). Figure 5 shows in each game, the facet marked in yellow (that we had

hypothesized to emerge as the best-predicting facet) actually emerged as the best predictor according to the significance tests reported in Table S5.2. The only exception was the loyalty game. The friendliness facet did not predict loyalty better than the assertiveness facet, $p = .930$.

Discussion

Studies 5a-5d aimed validated the facet model by showing that cultural evaluations of real and nuanced individuals on each facet predicted a relevant and specific behavior towards these unknown (vs. familiar in Studies 1-2a) and encountered (vs. envisioned in Studies 1a-4) individuals better than cultural evaluations of the individuals on all three other facets. The only exception was the friendliness facet in the loyalty game, which predicted loyalty but did not predict loyalty better than the assertiveness facet.

General Discussion

The present research validated a recently endorsed model of social-evaluative dimensions that distinguishes the two horizontal facets friendliness and morality from the two vertical facets ability and assertiveness (Abele et al., 2021; Koch et al., 2021). In three studies and two supplemental studies, perceivers evaluated different types and numbers of targets in different modes (i.e., separate or joint evaluation) and ways (i.e., personal or cultural impressions). Across these five studies, an efficient two-items-per-facet measure fit the data well enough and better than a four-items-per-dimension measure (i.e., Big Two model) and five-items-per-facet measure (see Supplemental Studies 1 and 2). In addition, the efficient measure confirmed the hypothesized pattern of statistical relations between the facets, namely larger correlations between the friendliness and morality facets, and between the ability and assertiveness facets, compared to the correlations between any one horizontal facet and any one vertical facet. Apart from corroborating the external, internal, convergent, discriminant validity of the facet model in these ways, the present research corroborated its predictive

validity across two additional studies. Evaluations on each facet predicted a broadly relevant behavior intention towards the evaluated social entities (i.e., making one's bonus dependent on their loyalty, fairness, success, and hubris) better than evaluations on all other facets in at least one of the two studies.

The present findings advance the validation of the facet model in several important ways. We validated the facet model for the within-perceiver perspective by factor-analyzing differences between targets within the perceivers. Other research (Abele et al., 2016; Abele & Hauke, 2019; Barbedor et al., 2024) validated the facet model for the within-target perspective by factor-analyzing differences between perceivers within the targets. Validating the facet model for the within-perceiver perspective required an efficient two-items-per-facet measure (vs. the exhaustive five-items-per-facet measure in the previous work) that is more feasible in studies in which perceivers evaluate several or even many social entities on the facets. The tabular analyses in the online supplement (Tables S1.5-7, Table SS2.1, and see Supplemental Studies 1 and 2 as well) show that the efficient measure validated the facet model for the within-target perspective as well.

In addition, we validated the two-items-per-facet measure across envisioned and encountered (i.e., actually seen) individuals and groups as the targets of evaluation, and across several modes and types of social evaluation. These modes include separate, joint, personal, and cultural evaluation, which covers not just various real-life situations, but also the standard paradigms in various research programs that examine social evaluation (Abele & Wojciszke, 2014; Ellemers, 2017; Fiske, 2018; Koch et al., 2016; Yzerbyt, 2018). We note that the efficient two-items-per-facet measure described cultural evaluations better than personal evaluations, and separate evaluations slightly better than joint evaluations, according to the fit indices that we considered in Studies 1 and 2.

Limitations and future research

First, we validated the facet model for evaluations of U.S. targets by U.S. perceivers. Other research validated the facet model for European contexts as well as China and Australia (Abele et al., 2016; Barbedor et al., 2024). Future research should generalize the two-items-per-facet measure to other national contexts, especially those that are not WEIRD (White, educated, industrialized, rich, and democratic; Muthukrishna et al., 2020).

Second, in Supplemental Studies 1 and 2, we reduced an initial list of five items per facet to two items per facet to achieve a satisfactory level of internal, convergent, and discriminant validity. Future research may examine whether this gain incurred a loss in content validity (i.e., covering every nuance of the content of each facet).

Third, we validated the two-items-per-facet measure through multi-level confirmatory factor analyses that treated the evaluators as data clusters. Thus, we can infer from our results that the two-items-per-facet measure was an adequate description of how an evaluator differentiated the targets, on average. We did not run multi-group confirmatory factor analysis that treated the evaluators as groups, not least because the ratio of targets to items was (very) low in our studies (e.g., four individuals versus eight items in Study 2). Thus, we cannot infer from our results that the two-items-per-facet measure was an adequate description of how each evaluator differentiated the targets. In fact, previous research showed that sign and size of the perceived correlations between various social-evaluative dimensions vary across evaluators (Hehman et al., 2019; Stolier et al., 2018; 2020). Accordingly, it is unlikely that our two-items-per-facet measure is an adequate description of how each evaluator differentiates the targets that we examined.

Fourth, we predicted personal behaviors towards the targets from cultural evaluations rather than personal evaluations by those that showed the behaviors, which would have been less bold but more straightforward.

Finally, we predicted hypothetical behaviors and not actually incentivized behaviors towards the targets. Future research may address the above shortcomings.

Conclusion

The present endeavor secured considerable progress with validating a recent model of four social-evaluative facets (friendliness, morality, ability, and assertiveness) endorsed by an ongoing adversarial collaboration (Ellemers et al., 2020). Our validation work is sufficient to comfortably begin treating the two-items-per-facet measure that we contribute as a useful tool to capture social evaluation. We hope to facilitate more research that decomposes the Big Two into the four basic facets of social evaluation (e.g., studies that aim to resolve the controversies between different research programs; Abele et al., 2021; Koch et al., 2021).

References

Abele, A. E. (2022). Evaluation of the self on the big two and their facets: Exploring the model and its nomological network. *International Review of Social Psychology*, 35(1), 1-15. <https://doi.org/10.5334/irsp.688>

Abele, A. E., Cuddy, A. J. C., Judd, C. M., & Yzerbyt, V. Y. (2008). Fundamental dimensions of social judgment: A view from different perspectives. *European Journal of Social Psychology*, 38, 1063-1065. <https://doi.org/10.1037/0022-3514.89.6.899>

Abele, A. E., Ellemers, N., Fiske, S. T., Koch, A., & Yzerbyt, V. (2021). Navigating the social world: Toward an integrated framework for evaluating self, individuals, and groups. *Psychological Review*, 128(2), 290-314. <https://doi.org/10.1037/rev0000262>

Abele, A. E., & Hauke, N. (2020). Comparing the facets of the big two in global evaluation of self versus other people. *European Journal of Social Psychology*, 50(5), 969-982. <https://doi.org/10.1002/ejsp.2639>

Abele, A. E., Hauke, N., Peters, K., Louvet, E., Szymkow, A., & Duan, Y. (2016). Facets of the fundamental content dimensions: Agency with competence and assertiveness—Communion with warmth and morality. *Frontiers in Psychology*, 7, 1810. <https://doi.org/10.3389/fpsyg.2016.01810>

Abele, A. E., & Wojciszke, B. (2014). Communal and agentic content in social cognition: A dual perspective model. In J. M. Olson & M. P. Zanna (Eds.), *Advances in Experimental Social Psychology* (Vol. 50, pp. 195-255). Cambridge, MA: Academic Press. <https://doi.org/10.1016/B978-0-12-800284-1.00004-7>

Barbedor, J., Schneider, J., Yzerbyt, V., & Abele, A. (2024). A novel approach to the evaluation of groups: Type of group and facet of evaluation matter. *Manuscript under review*.

Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67, 1-48.

<https://doi.org/10.18637/jss.v067.i01>

Bazerman, M. H., Moore, D. A., Tenbrunsel, A. E., Wade-Benzoni, K. A., & Blount, S. (1999). Explaining how preferences change across joint versus separate evaluation. *Journal of Economic Behavior & Organization*, 39(1), 41-58.

[https://doi.org/10.1016/S0167-2681\(99\)00025-6](https://doi.org/10.1016/S0167-2681(99)00025-6)

Brambilla, M., Sacchi, S., Rusconi, P., Cherubini, P., & Yzerbyt, V. Y. (2012). You want to give a good impression? Be honest! Moral traits dominate group impression formation. *British Journal of Social Psychology*, 51(1), 149-166.

<https://doi.org/10.1111/j.2044-8309.2010.02011.x>

Carrier, A., Louvet, E., Chauvin, B., & Rohmer, O. (2014). The primacy of agency over competence in status perception. *Social Psychology*, 45(5), 1-10.

<https://doi.org/10.1027/1864-9335/a000176>

Douglas, B. D., Ewell, P. J., & Brauer, M. (2023). Data quality in online human-subjects research: Comparisons between MTurk, Prolific, CloudResearch, Qualtrics, and SONA. *PLOS One*, 18(3), e0279720. <https://doi.org/10.1371/journal.pone.0279720>

Ellemers, N., Fiske, S. T., Abele, A. E., Koch, A., & Yzerbyt, V. (2020). Adversarial alignment enables competing models to engage in cooperative theory building toward cumulative science. *Proceedings of the National Academy of Sciences*, 117(14), 7561-7567. <https://doi.org/10.1073/pnas.1906720117>

Fiske, S. T. (1992). Thinking is for doing: portraits of social cognition from daguerreotype to laserphoto. *Journal of Personality and Social Psychology*, 63(6), 877-889.

<https://doi.org/10.1037/0022-3514.63.6.877>

- 1
2
3 Fiske, S. T. (2018). Stereotype content: Warmth and competence endure. *Current Directions*
4 *in Psychological Science*, 27(2), 67-73. <https://doi.org/10.1177/0963721417738825>
5
6
7
8 Fiske, S. T. (2002). What we know now about bias and intergroup conflict, the problem of
9 the century. *Current Directions in Psychological Science*, 11(4), 123-128.
10
11 <https://doi.org/10.1111/1467-8721.00183>
12
13
14 Fox, J., & Weisberg, S. (2018). *An R companion to applied regression* (3rd ed.). Sage
15 publications.
16
17
18 Gligorić, V., van Kleef, G. A., & Rutjens, B. T. (2022). Social evaluations of scientific
19 occupations. *Scientific Reports*, 12(1), 18339. [https://doi.org/10.1038/s41598-022-](https://doi.org/10.1038/s41598-022-23197-7)
20 [23197-7](https://doi.org/10.1038/s41598-022-23197-7)
21
22
23
24
25
26 Gneezy, U. (2005). Deception: The role of consequences. *American Economic Review*, 95(1),
27 384-394. <https://doi.org/10.1257/0002828053828662>
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

- Imhoff, R., & Koch, A. (2017). How orthogonal are the Big Two of social perception? On the curvilinear relation between agency and communion. *Perspectives on Psychological Science*, 12(1), 122-137. <https://doi.org/10.1177/1745691616657334>
- Imhoff, R., Koch, A., & Flade, F. (2018). (Pre)occupations: A data-driven model of jobs and its consequences for categorization and evaluation. *Journal of Experimental Social Psychology*, 77, 76-88. <https://doi.org/10.1016/j.jesp.2018.04.001>
- Judd, C. M., Garcia-Marques, T., & Yzerbyt, V. (2019). The complexity of relations between dimensions of social perception: Decomposing bivariate associations with crossed random factors. *Journal of Experimental Social Psychology*, 82, 200-207. <https://doi.org/10.1016/j.jesp.2019.01.008>
- Kline, R. B. (2005). *Principles and practice of structural equation modeling*. (2nd ed.). Guilford.
- Koch, A., Dorrough, A., Glöckner, A., & Imhoff, R. (2020). The ABC of society: Perceived similarity in agency/socioeconomic success and conservative-progressive beliefs increases intergroup cooperation. *Journal of Experimental Social Psychology*, 90, 103996. <https://doi.org/10.1016/j.jesp.2020.103996>
- Koch, A., Imhoff, R., Dotsch, R., Unkelbach, C., & Alves, H. (2016). The ABC of stereotypes about groups: Agency/socioeconomic success, conservative–progressive beliefs, and communion. *Journal of Personality and Social Psychology*, 110(5), 675-709. <https://doi.org/10.1037/pspa0000046>
- Koch, A., Imhoff, R., Unkelbach, C., Nicolas, G., Fiske, S., Terache, J., Carrier, A., & Yzerbyt, V. (2020). Groups' warmth is a personal matter: Understanding consensus on stereotype dimensions reconciles adversarial models of social evaluation. *Journal of Experimental Social Psychology*, 89, 103995. <https://doi.org/10.1016/j.jesp.2020.103995>

- Koch, A., Yzerbyt, V., Abele, A., Ellemers, N., & Fiske, S. T. (2021). Social evaluation: Comparing models across interpersonal, intragroup, intergroup, several-group, and many-group contexts. In B. Gawronski (Ed.), *Advances in Experimental Social Psychology* (Vol. 63, pp. 1-68). Cambridge, MA: Academic Press.
<https://doi.org/10.1016/bs.aesp.2020.11.001>
- Leach, C. W., Ellemers, N., & Barreto, M. (2007). Group virtue: The importance of morality (vs. competence and sociability) in the positive evaluation of in-groups. *Journal of Personality and Social Psychology*, 93(2), 234-249.
<https://doi.org/10.1037/0022-3514.93.2.234>
- Lejuez, C. W., Read, J. P., Kahler, C. W., Richards, J. B., Ramsey, S. E., Stuart, G. L., Strong, D. R., & Brown, R. A. (2002). Evaluation of a behavioral measure of risk taking: The Balloon Analogue Risk Task (BART). *Journal of Experimental Psychology: Applied*, 8(2), 75–84. <https://doi.org/10.1037/1076-898X.8.2.75>
- Lickel, B., Hamilton, D. L., Wierzchowska, G., Lewis, A., Sherman, S. J., & Uhles, A. N. (2000). Varieties of groups and the perception of group entitativity. *Journal of Personality and Social Psychology*, 78(2), 223–246. <https://doi.org/10.1037/0022-3514.78.2.223>
- Luchman, M. J. (2023). Package ‘domir’. Available from <https://bioconductor.statistik.tu-dortmund.de/cran/web/packages/domir/domir.pdf>
- McNeish, D., & Wolf, M. G. (2023). Dynamic fit index cutoffs for one-factor models. *Behavior Research Methods*, 55(3), 1157-1174.
<https://doi.org/10.3758/s13428-022-01847-y>

- Muthukrishna, M., Bell, A. V., Henrich, J., Curtin, C. M., Gedranovich, A., McInerney, J., & Thue, B. (2020). Beyond western, educated, industrial, rich, and democratic (WEIRD) psychology: Measuring and mapping scales of cultural and psychological distance. *Psychological Science*, 31(6), 678-701. <https://doi.org/10.1177/0956797620916782>
- Nicolas, G., Bai, X., & Fiske, S. T. (2022). A spontaneous stereotype content model: Taxonomy, properties, and prediction. *Journal of Personality and Social Psychology*, 123(6), 1243-1263. <https://doi.org/10.1037/pspa0000312>
- Oliveira, M., Garcia-Marques, T., Garcia-Marques, L., & Dotsch, R. (2020). Good to Bad or Bad to Bad? What is the relationship between valence and the trait content of the Big Two?. *European Journal of Social Psychology*, 50(2), 463-483. <https://doi.org/10.1002/ejsp.2618>
- Peer, E., Brandimarte, L., Samat, S., & Acquisti, A. (2017). Beyond the Turk: Alternative platforms for crowdsourcing behavioral research. *Journal of Experimental Social Psychology*, 70, 153-163. <https://doi.org/10.1016/j.jesp.2017.01.006>
- Pornprasertmanit, S., Lee, J., & Preacher, K. J. (2014). Ignoring clustering in confirmatory factor analysis: Some consequences for model fit and standardized parameter estimates. *Multivariate Behavioral Research*, 49(6), 518-543. <https://doi.org/10.1080/00273171.2014.933762>
- Rau, R., Carlson, E. N., Back, M. D., Barranti, M., Gebauer, J. E., Human, L. J., Leising, D., & Nestler, S. (2021). What is the structure of perceiver effects? On the importance of global positivity and trait-specificity across personality domains and judgment contexts. *Journal of Personality and Social Psychology*, 120(3), 745-764. <https://doi.org/10.1037/pspp0000278>

- Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48(2), 1-36. <https://doi.org/10.18637/jss.v048.i02>
- Stolier, R. M., Hehman, E., & Freeman, J. B. (2018). A dynamic structure of social trait space. *Trends in Cognitive Sciences*, 22(3), 197-200.
<https://doi.org/10.1016/j.tics.2017.12.003>
- Stolier, R. M., Hehman, E., & Freeman, J. B. (2020). Trait knowledge forms a common structure across social cognition. *Nature Human Behaviour*, 4(4), 361-371.
<https://doi.org/10.1038/s41562-019-0800-6>
- Thielmann, I., Böhm, R., Ott, M., & Hilbig, B. (2021). Economic Games: An Introduction and Guide for Research. *Collabra: Psychology*, 7(1), 19004.
<https://doi.org/10.1525/collabra.19004>
- U.S. Census Bureau. (n.d.). United States [data table]. *Profile*. Retrieved from
https://data.census.gov/profile/United_States?g=010XX00US
- Walsh, J., Vaida, N., & Fiske, S. (2023). Stereotype content predicts economic discrimination even under incentives. *Manuscript under review*.
- Yzerbyt, V. Y. (2018). The dimensional compensation model: Reality and strategic constraints on Warmth and Competence in intergroup perceptions. In A. E. Abele & B. Wojciszke (Eds.), *The agency-communion framework* (pp.126–141). London, UK: Routledge.
- Yzerbyt, V., Barbedor, J., Carrier, A., & Rohmer, O. (2022). The facets of social hierarchy: How judges' legitimacy beliefs and relative status shape their evaluation of assertiveness and ability. *International Review of Social Psychology*, 35(1), 18.
<https://doi.org/10.5334/irsp.695>

Supplementary Information for “A brief measure of four basic facets of social evaluation”

Table of Contents

TABLE OF CONTENTS	1
SUPPLEMENTAL STUDY 1.....	3
METHOD	3
RESULTS	3
TABLE SS1.1	4
TABLE SS1.2	5
FIGURE SS1.1.....	6
DISCUSSION.....	6
FIGURE SS1.2.....	7
FIGURE SS1.3.....	8
TABLE SS1.1	8
SUPPLEMENTAL STUDY 2.....	9
METHOD	9
TABLE SS2.1	10
RESULTS	10
DISCUSSION.....	12
FIGURE SS2.2.....	13
FIGURE SS2.3.....	13
TABLE SS2.4	14
STUDIES 1A AND 1B.....	15
FIGURE S1.1	15
FIGURE S1.2	15
FIGURE S1.3	16
FIGURE S1.4	16
FIGURE S1.5	17
FIGURE S1.6	17
FIGURE S1.7	18
FIGURE S1.8	18
FIGURE S1.9	19
FIGURE S1.10	19
FIGURE S1.11	20
FIGURE S1.12	20
FIGURE S1.13	21
FIGURE S1.14	21
FIGURE S1.15	22
FIGURE S1.16	22
FIGURE S1.17	23
FIGURE S1.18	23
FIGURE S1.19	24
FIGURE S1.20	24
FIGURE S1.21	25
FIGURE S1.22	25
FIGURE S1.23	26
TABLE S1.1	26
TABLE S1.2	27
TABLE S1.3	27

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

TABLE S1.4	28
TABLE S1.5	29
TABLE S1.6	30
TABLE S1.7	30
TABLE S1.8	31
STUDY 2.....	32
FIGURE S2.1	32
FIGURE S2.2	32
FIGURE S2.3	33
FIGURE S2.4	33
FIGURE S2.5	34
FIGURE S2.6	34
FIGURE S2.7	35
TABLE S2.1	35
TABLE S2.2	36
STUDY 3.....	37
FIGURE S5.1	37
STUDY 4.....	38
TABLE S4.1	38
TABLE S4.2	39
STUDY 5.....	40
TABLE S5.1	40
TABLE S5.2	41

For Review Only

Supplemental Study 1

To make progress towards our overall aim to validate a brief measure of four correlated facets of social evaluation (friendliness, morality, ability, and assertiveness; Abele et al., 2021; Koch et al., 2021), we reexamined and collectively discussed in a virtual meeting the measurement habits across the core publications of our five models of social evaluation. This adversarial alignment led us to settle on five items per facet that seemed to capture the content of the facets satisfactorily from the perspective of each model.

Method

Participants. We sampled 1,010 people from Prolific. We excluded 11 people who recommended that we do not analyze their data, leaving 999 people (48.2% female, 50.8% male, 1.0% other; $M_{age} = 43.33$, $SD = 15.91$).

Stimuli. Each person rated a random selection of six out of the 42 groups that other people in previous research had listed most frequently when instructed to list the groups that together form today's U.S. society (see Study 5a in Koch et al., 2016). The 42 groups were defined by their gender, age, race, status, beliefs etc. We list them in alphabetical order: Asian people, atheist people, athletes, Black people, blue collar workers, celebrities, Christians, conservatives, Democrats, drug addicts, elderly people, gay people, goths, hippies, hipsters, Hispanic people, homeless people, homosexual people, immigrants, Jews, jocks, lesbian people, liberals, lower class people, men, middle class people, Muslims, nerds, parents, politicians, poor people, religious people, Republicans, rich people, students, teenagers, transgender people, upper class people, white collar workers, White people, women, and working class people. These groups are social categories, according to a categorization by Lickel and colleagues (2000; they also mention intimacy groups [e.g., friends], task groups [e.g., a committee], and loose associations [e.g., riders on a bus]; ongoing research by Barbedor and colleagues [2024] validates the facet model for intimacy and task groups).

Procedure. People used 7-point scales that ranged from 1 = "not at all" to 7 = "extremely" to rate each of the six familiar and envisioned groups on 20 items. The items "friendly," "warm," "nice," "sociable," and "agreeable" aimed to measure the perceived friendliness of the groups. "Honest," "sincere," "moral," "reliable," and "trustworthy" aimed to measure the perceived morality of the groups. "Skilled," "capable," "competent," "intelligent," and "smart" aimed to measure the perceived ability of the groups. And "confident," "determined," "persistent," "self-assured," and "assertive" aimed to measure the perceived assertiveness of the groups.

People rated their personal impressions of the groups. In addition, people rated one group on all 20 items (in random order) before rating another group on the next survey page (i.e., sequential evaluation), with the order of the six groups being random as well. People read "to what extent do you think of [group; e.g. Asian people] as [item; e.g., FRIENDLY]?" This separate mode of social evaluation prioritizes depth over breadth and represents real-life situations in which a perceiver encounters targets sequentially rather than simultaneously (e.g., when they interview job candidates; Abele & Wojciszke, 2014; Nicolas et al., 2022).

At the end of the study, people provided demographic information, including their gender and age.

Results

We used the cfa function of the R package lavaan (Rosseel, 2012) to run a multi-level confirmatory factor analysis (MCFA; Pornprasertmanit et al., 2014) that accounted for dependencies in our data (i.e., every perceiver rated multiple targets). We defined the raters as

data clusters (i.e., we treated the groups as nested within the raters), to model the differences between the groups within the raters (instead of modeling the differences between the raters within the groups as in previous work; Abele et al., 2016; Abele & Hauke, 2019).

Level 1 – our main interest – had the groups as rows, the items as columns, and individual ratings as the data. The five items meant to capture friendliness loaded on the latent friendliness facet, the five morality items loaded on the latent morality facet, the five ability items loaded on the latent ability facet, and the five assertiveness items loaded on the latent assertiveness facet. We allowed no cross-loadings (i.e., we modeled a simple structure), but we allowed the latent facets to correlate with one another. We additionally examined level 2, which had the raters as rows, the items as columns, and mean ratings across the groups as the data. All eight manifest items loaded on a latent general factor because we had no hypothesis for the factor structure of level 2, and thus decided to keep it simple. Level 2’s latent general factor captured that some people gave higher ratings to all groups on all items due to acquiescence, complaisance (i.e., wanting to be liked), philanthropy, or a combination of these. Figure SS1.2 shows the estimated parameters for Levels 1 and 2.

We report multiple fit indices (χ^2 , CFI, RMSEA, SRMR, and AIC) to account for the limitations inherent in relying on a single index (Kline, 2005). We compare the fit of the measure against the standard cutoffs of ≥ 0.95 for CFI and ≤ 0.06 for RMSEA and SRMR (Hu & Bentler, 1999). The five-items-per-facet measure did not fit the data (see Table SS1.1), presumably because some items captured more facets than their designated facet, but the model’s simple structure (that we had imposed for internal validity) ignored cross-loadings.

Table SS1.1
Supplemental Study 1: Satisfactory and Superior Model Fit of a Two-Items-Per-Facet Measure

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; df = 34	382	0.99	0.04	0.09	132,521	
Two-items-per-facet & Two-facets-per-dimension; df = 35	408	0.99	0.04	0.09	132,545	<.001
Four-items-per-dimension; df = 39	2,336	0.92	0.10	0.12	134,465	<.001
Five-items-per-facet; df = 334	8,867	0.92	0.07	0.16	308,187	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

To conclude satisfactory convergent and discriminant validity, the correlation between the latent friendliness and morality (i.e., horizontal) facets and the correlation between the latent ability and assertiveness (i.e., vertical) facets would have to be positive and larger than all correlations between one latent horizontal facet and one latent vertical facet. Figure SS1.2 shows that the correlations between the latent facets disconfirmed the sought-after pattern.

Next, we computed a friendliness scale by averaging the five friendliness items (i.e., “friendly,” “warm,” “nice,” “sociable,” and “agreeable”). Likewise, we computed a morality, ability, and assertiveness scale by averaging the five morality, ability, and assertiveness items, respectively. For each scale, we centered the differences between the groups within each rater, to account for individual differences in acquiescence, complaisance, and/or philanthropy as in the MCFA. Table SS1.2 shows the correlations between the within-centered facet scales, which

again disconfirmed the sought-after pattern, presumably also because of fuzzy items that captured more facets than their designated facet.

Table SS1.2

Supplemental Study 1: Correlations Between the Facet Scales Centered Within the Raters

	Two Items per Facet				Five Items per Facet			
Friendliness -	1.00	0.73	0.48	0.24	1.00	0.77	0.60	0.26
Morality -	0.73	1.00	0.56	0.21	0.77	1.00	0.70	0.23
Ability -	0.48	0.56	1.00	0.55	0.60	0.70	1.00	0.50
Assertiveness -	0.24	0.21	0.55	1.00	0.26	0.23	0.50	1.00
	Friendliness	Morality	Ability	Assertiveness	Friendliness	Morality	Ability	Assertiveness

Note. More intense hues indicate larger correlations.

Rejecting the five-items-per-facet measure on the grounds of insufficient internal, convergent, and discriminant validity, we resorted to exploring the data by estimating several two-items-per-facet models one at a time. We selected two items per facet that, first, formed a well-fitting (see Table SS1.1) simple-structured model, and, second, whose four latent facets correlated more strongly within (vs. between) the Big Two dimensions (see Figure SS1.1, which shows the estimated parameters for this two-items-per-facet measure). This outcome dependent item selection was a problem that we resolved by replicating the (internal, convergent, and discriminant) validity of the emerging two items-per-facet model in the subsequent studies. The two manifest items “friendly” and “warm” loaded on the latent friendliness facet, “honest” and “sincere” loaded on the morality facet, “capable” and “skilled” loaded on the ability facet, and “confident” and “determined” loaded on the assertiveness facet.

Figure SS1.1
Supplemental Study 1: Parameters of a Two-Items-Per-Facet Measure of Personal Evaluations

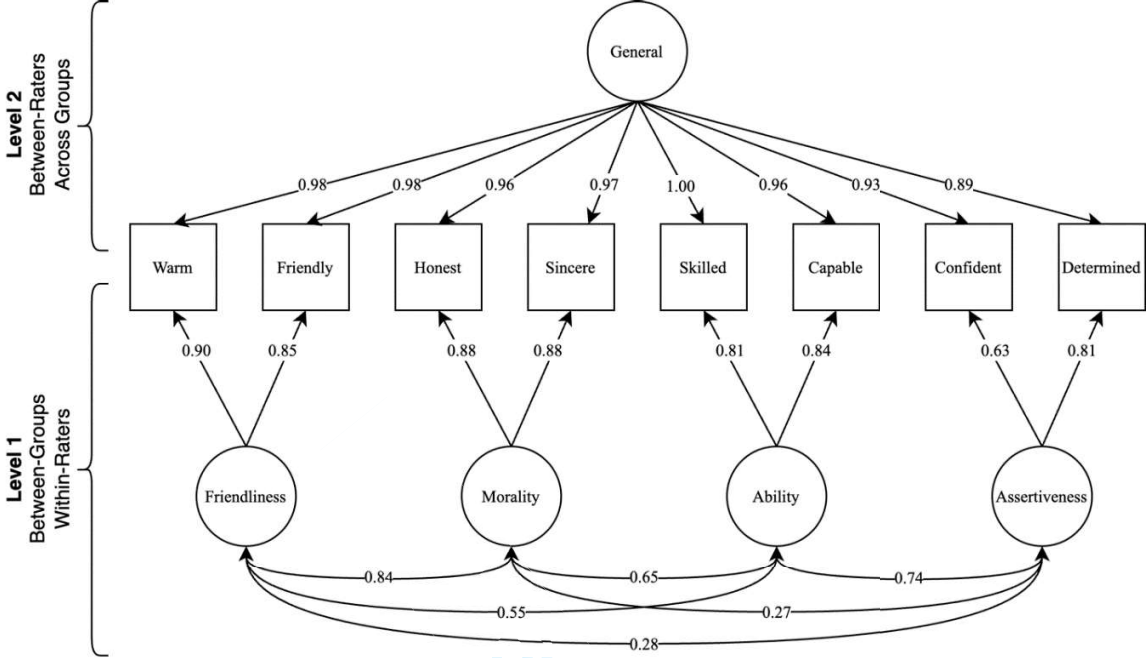


Table SS1.1 also shows the fit of a measure in which the manifest items “friendly,” “warm,” “honest,” and “sincere” loaded on a latent horizontal dimension, and “capable,” “skilled,” “confident,” and “determined” loaded on a vertical dimension that we allowed to correlate with the horizontal dimension. The fit of the four-items-per-dimension measure was worse than the fit of the two-items-per-facet measure according to nested chi-square difference tests (all p ’s $< .001$). Figure SS1.3 shows the estimated parameters for the worse-fitting four-items-per-dimension measure, which represented the standard Big Two model that does not differentiate its two dimensions into the facets.

Discussion

Based on goodness of absolute and relative model fit as well as correlations between manifest and latent factors, Study 1 rejected an exhaustive five-items-per-facet measure and instead explored an efficient two-items-per-facet measure of four correlated facets of social evaluation (ability, assertiveness, morality, and friendliness; Abele et al., 2021).

An important question is how our brief measure compares to the exhaustive measures of the facets in previous work. Only four, four, and six of the eight items in our brief measure are included in the exhaustive 20-item, 16-item, and 20-item measures contributed by Abele and colleagues (2016), Abele and Hauke (2019), and Barbedor and colleagues (2024), respectively. Because our brief measure is not nested in their exhaustive ones, we cannot compare their goodness of model fit.

In our view, the unique strengths of their measures are content and cross-cultural validity. As to content validity, their exhaustive 16/20-item measures likely cover more nuances of the facets than our brief eight-item measure, which, for example, captures the morality facet in terms of honesty but not fairness etc. As to cross-cultural validity, Abele and colleagues’ facet measure of impressions of individuals generalizes across several languages and national contexts, and the same is true for Barbedor and colleagues’ (2024) facet measure of impressions of groups.

Besides efficiency, the unique strengths of our brief measure include convergent, discriminant, ecological, and external validity in English as the lingua franca of the U.S. national context. As to convergent and discriminant validity, we confirm higher correlations between the facets within (vs. across) the Big Two dimensions. As to ecological and external validity, our facet measure of individuals and groups generalizes across the separate and joint mode of social evaluation (see Studies 1 and 2), across personal and cultural social evaluation (see Studies 1 and 2), and across one target as seen by different perceivers and different targets as seen by one perceiver.”

Figure SS1.2
Supplemental Study 1: Parameters of a Five-Items-Per-Facet Measure of Personal Evaluations

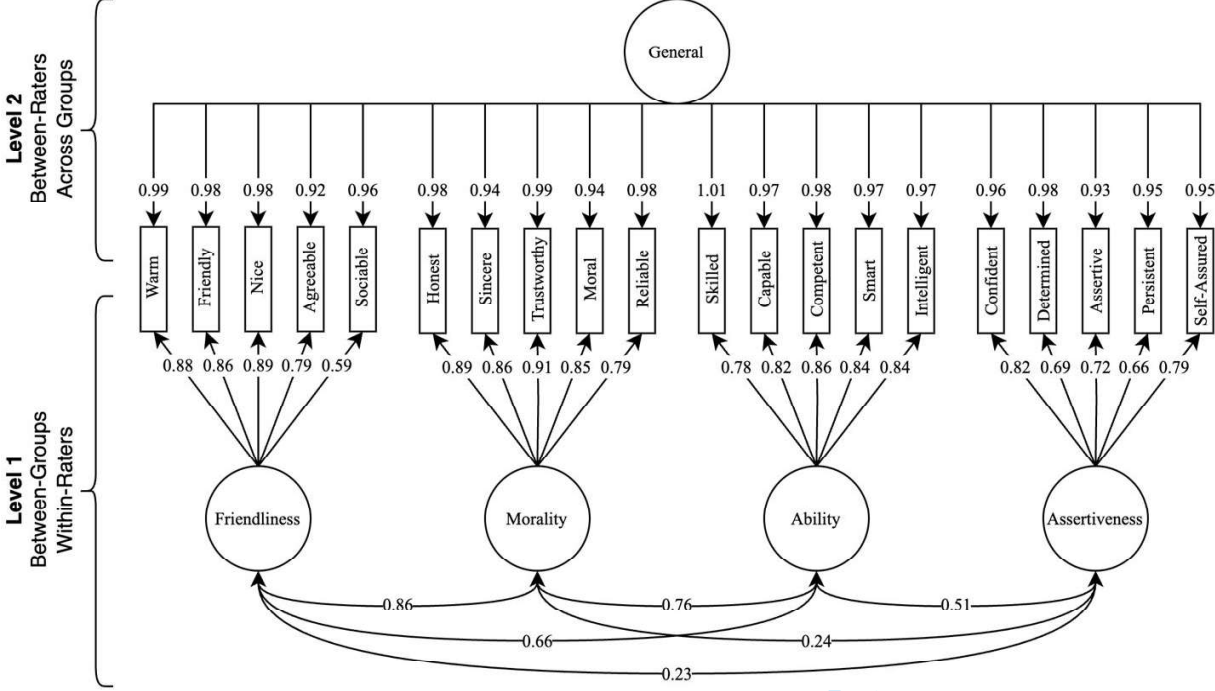


Figure SS1.3
Supplemental Study 1: Parameters of a Four-Items-Per-Dimension Measure of Personal Evaluation

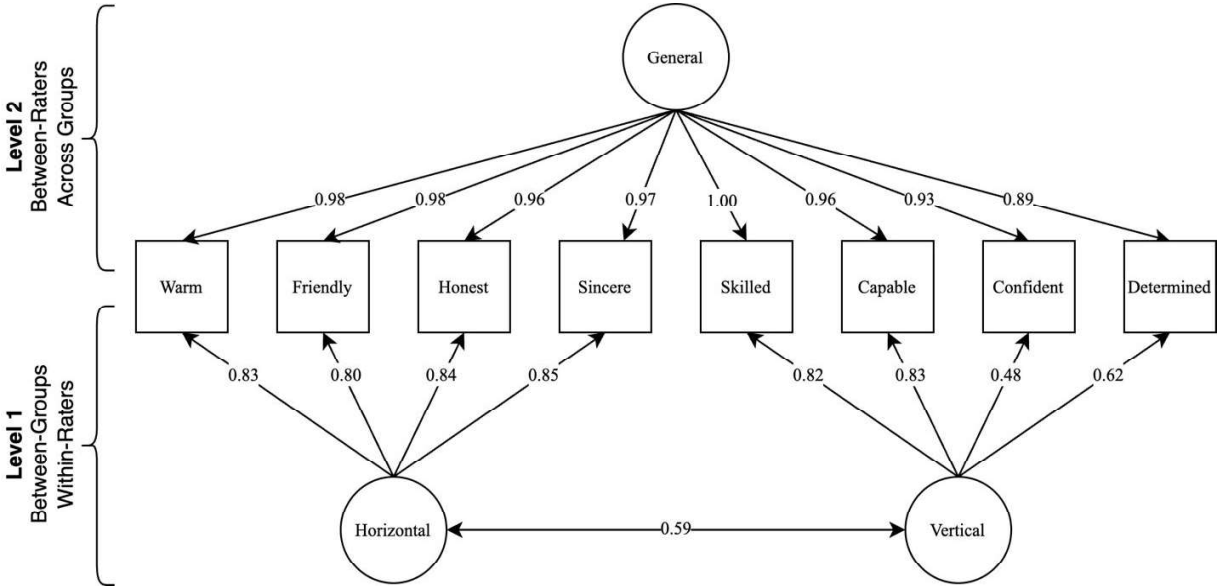


Table SS1.1
Supplemental Study 1: Tabular confirmatory factor analyses, averaging across the 42 target groups

	X ²	CFI	RMSEA	SRMR	AIC
Two-items-per-facet; $df=34$	29.39	0.98	0.08	0.03	3,047
Four-items-per-dimension; $df=39$	70.91	0.94	0.13	0.06	3,078
Five-items-per-facet; $df=334$	442.38	0.90	0.11	0.07	7,068

Note. These were factor analyses that we ran separately for each group, with the perceivers as rows, the eight or twenty items as columns, and the perceivers' ratings of the target in a given factor analysis on a given item as the data points. The table reports fit indices averaged across all 42 groups that we examined in Supplemental Study 1.

Supplemental Study 2

We aimed to generalize the two-items-per-facet measure that we explored in Supplemental Study 1 and confirmed in Studies 1a and 1b from evaluations of groups (Studies 1 and 2) to evaluations of individuals. Accordingly, in Study 3 each person evaluated several individuals as in previous work (Abele et al., 2016; Abele & Hauke, 2019). Supplemental Study 2 also aimed to compare once again the adequacy of our efficient two-items-per-facet measure to the adequacy of the exhaustive five-items-per-facet measure that we examined in Supplemental Study 1.

Method

Participants. We sampled 1,009 people from Prolific. We excluded 19 people who recommended that we do not analyze their data, leaving 990 people (48.6% female, 51.1% male, 0.3% other; $M_{age} = 45.67$, $SD = 16.25$).

Stimuli and Procedure. People began by typing into text boxes the names of a same-sex close friend, acquaintance, and celebrity. These targets of social evaluation differ in terms of terms of their closeness to the perceiver and positivity in the eyes of the perceiver, two continua that characterize the focus of many, if not all, models of social evaluation (Abele et al., 2021; Koch et al., 2021). As in Supplemental Study 1, people separately and personally evaluated these individuals, themselves, and these groups on the same 20 items as in Supplemental Study 1. On each survey page, they read “to what extent do you think of [individual or group] as [items]?” At the end of the study, people provided demographic information, including their gender and age.

In addition to the individual targets reported above, participants typed the names of an in-group and out-group of theirs and evaluated these groups on the same variables as the individual targets. We excluded in-group and out-group targets because this study was focused on individual targets, and we had already confirmed the factor structure for groups in Supplemental Study 1.

We additionally measured perceived competitiveness, cooperativeness, shared values, similarity (between the participant and target), prestige, status, ideology, and traditionalism on 7-point scales. The exact wording of each question and their respective scales are reported in Table SS2.1. Participants also rated themselves on each of the variables reported in the manuscript and supplement except for similarity.

Table SS2.1
Supplemental Study 2: Additional Variables Collected

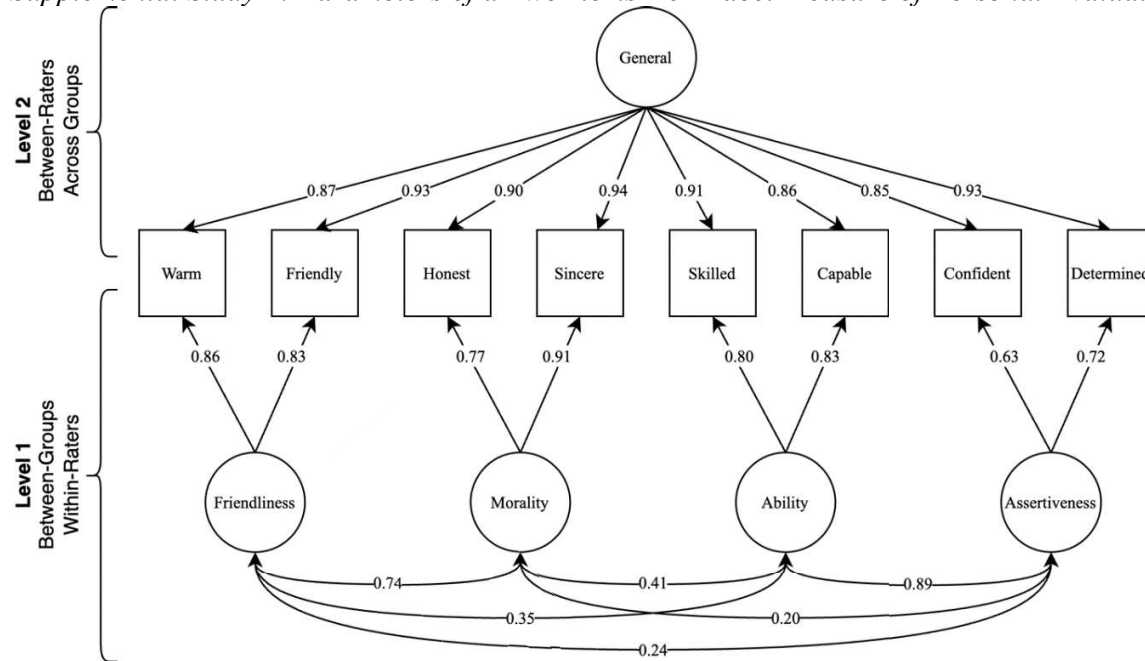
Measure	Question	Scale
Competitiveness	<i>“To what extent do you think of [individual or group] as competitive with other people (for example, exploiting limited resources or not)?”</i>	1 (“Not at all”) to 7 (“Extremely”)
Cooperativeness	<i>“To what extent do you think of [individual or group] as cooperative with the rest of society (for example, sharing limited resources or not)?”</i>	1 (“Not at all”) to 7 (“Extremely”)
Shared Values	<i>“To what extent do you think of [individual or target] as sharing goals or not sharing goals with the rest of society?”</i>	1 (“Not at all”) to 7 (“Extremely”)
Similarity	<i>“To what extent do you think that [individual or group] is similar to you?”</i>	1 (“Not at all”) to 7 (“Extremely”)
Prestige	<i>“How would you rate [individual or group] on having prestigious jobs?”</i>	1 (“Not at all”) to 7 (“Extremely”)
Status	<i>“How would you rate [individual or group] on economic success?”</i>	1 (“Not at all”) to 7 (“Extremely”)
Ideology	<i>“In politics today, do you think of [individual or group] as progressive or conservative?”</i>	1 (“Progressive”) to 7 (“Conservative”)
Traditionalism	<i>“Do you think of [individual or group] as modern or traditional?”</i>	1 (“Modern”) to 7 (“Traditional”)

Results

We used the cfa function of the R package lavaan (Rosseel, 2012) to run a multi-level confirmatory factor analysis (MCFA; Pornprasertmanit et al., 2014) on people’s evaluations of the four individuals. As in Supplemental Study 1 and Studies 1a and 1b, we defined the raters as data clusters (i.e., we treated the individuals as nested within the raters), to model the differences between the individuals within the raters. We specified Levels 1 and 2 in the same way as for the two-items-per-facet measures in Supplemental Study 1 and Studies 1a and 1b. Figure SS2.1 shows the parameter estimates for Levels 1 and 2.

Figure SS2.1

Supplemental Study 2: Parameters of a Two-Items-Per-Facet Measure of Personal Evaluations



We report multiple fit indices, and we compare the fit of the measure against the standard cutoffs of ≥ 0.95 for CFI and ≤ 0.06 for RMSEA and SRMR (Hu & Bentler, 1999). Table 3.1 shows that the two-items-per-facet measure fit the data well enough.

Table SS2.2

Supplemental Study 2: Satisfactory and Superior Model Fit of a Two-Items-Per-Facet Measure

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; $df = 34$	191	0.99	0.03	0.05	85,992	
Two-items-per-facet & Two-facets-per-dimension; $df = 35$	201	0.99	0.03	0.05	86,000	.002
Four-items-per-dimension; $df = 39$	1,904	0.90	0.11	0.12	87,695	<.001
Five-items-per-facet; $df = 334$	5,795	0.92	0.06	0.14	202,479	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

Table SS2.2 also shows that the more exhaustive five-items-per-facet measure and a four-items-per-dimension measure (that we specified as in Supplemental Study 1 and Studies 1a and 1b) fit the data worse than the hypothesized two-items-per-facet measure. Figures SS2.2 and SS2.2 show the estimated parameters for the two worse-fitting measures.

Again, we had hypothesized that the correlation between the friendliness and morality (i.e., horizontal) facets and the correlation between the ability and assertiveness (i.e., vertical) facets would be positive and larger than any correlation between one horizontal facet and one vertical facet. The data confirmed this for the within-centered facet scales computed

according to the two-items-per-facet measure but not the five-items-per-facet measure, see Table SS2.3. The correlations between the latent facets in the two-items-per-facet measure also confirmed the hypothesized pattern. This was not the case for the five-items-per-facet measure, compare Figure SS2.1 and Figure SS2.2.

Table SS2.3
Supplemental Study 2: Correlations Between the Four Facet Scales Centered Within the Raters

	Two Items per Facet				Five Items per Facet			
Friendliness -	1.00	0.62	0.40	0.19	1.00	0.63	0.49	0.22
Morality -	0.62	1.00	0.46	0.16	0.63	1.00	0.60	0.15
Ability -	0.40	0.46	1.00	0.60	0.49	0.60	1.00	0.54
Assertiveness -	0.19	0.16	0.60	1.00	0.22	0.15	0.54	1.00
	Friendliness	Morality	Ability	Assertiveness	Friendliness	Morality	Ability	Assertiveness

Note. More intense hues indicate larger correlations.

Discussion

Based on goodness of absolute and relative model fit as well as correlations between both manifest scales and latent factors, Supplemental Study 2 generalized our efficient two-items-per-facet measure of four correlated facets of social evaluation (ability, assertiveness, morality, and friendliness; Abele et al., 2021; Koch et al., 2021) from evaluations of several or many groups to evaluations of several individuals.

Figure SS2.2

Supplemental Study 2: Parameters of a Five-Items-Per-Facet Measure of Personal Evaluations

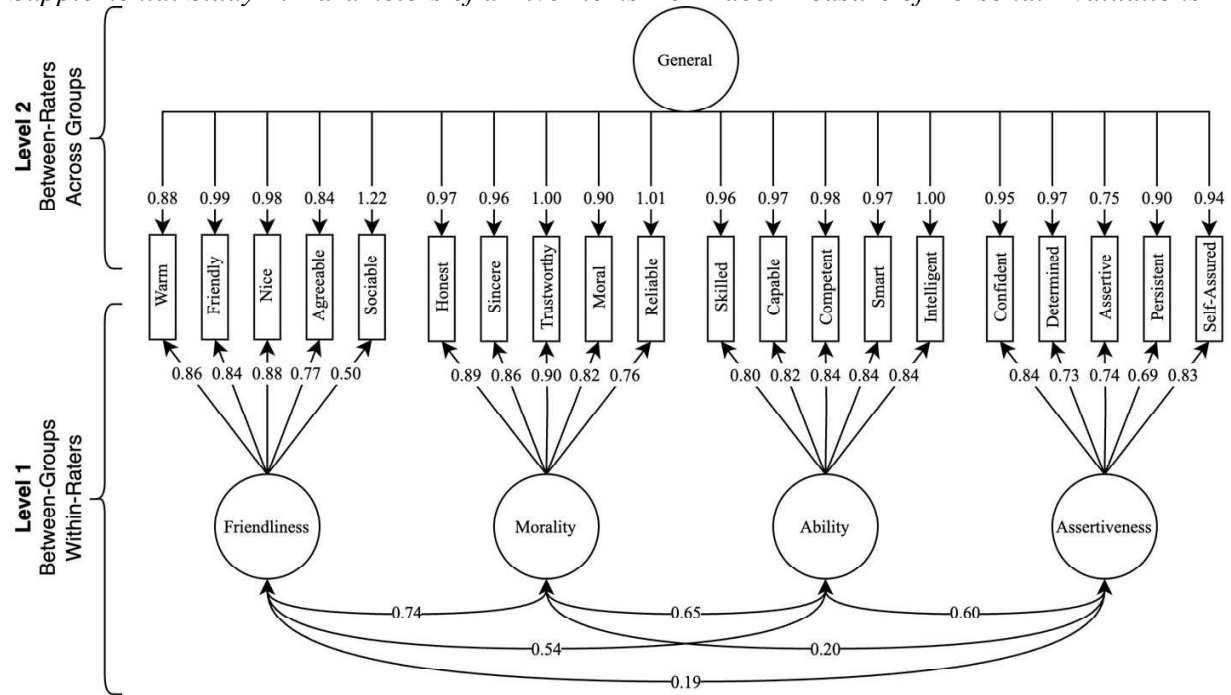
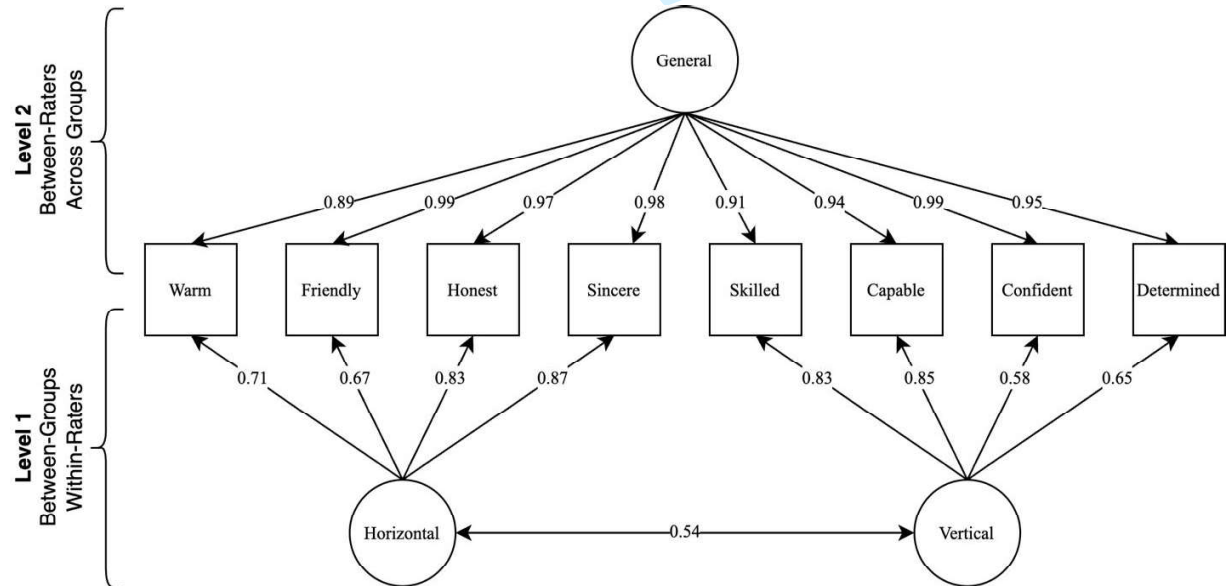


Figure SS2.3

Supplemental Study 2: Parameters of a Four-Items-Per-Dimension Measure of Personal Evaluation



1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

Table SS2.4
Supplemental Study 2: Tabular confirmatory factor analyses, averaging across the 4 target individuals

	X ²	CFI	RMSEA	SRMR	AIC
Two-items-per-facet; <i>df</i> = 34	39.78	0.99	0.04	0.02	21,058
Four-items-per-dimension; <i>df</i> = 39	388.41	0.92	0.14	0.05	21,397
Five-items-per-facet; <i>df</i> = 334	1,420.24	0.92	0.09	0.05	49,778

Note. These were factor analyses that we ran separately for each individual, with the perceivers as rows, the eight or twenty items as columns, and the perceivers' ratings of the target in a given factor analysis on a given item as the data points. The table reports fit indices averaged across all 4 individuals that we examined in Supplemental Study 2.

For Review Only

Studies 1a and 1b

Figure S1.1

Study 1: Parameters of a Two-Items-Per-Facet Measure of Separate Evaluations

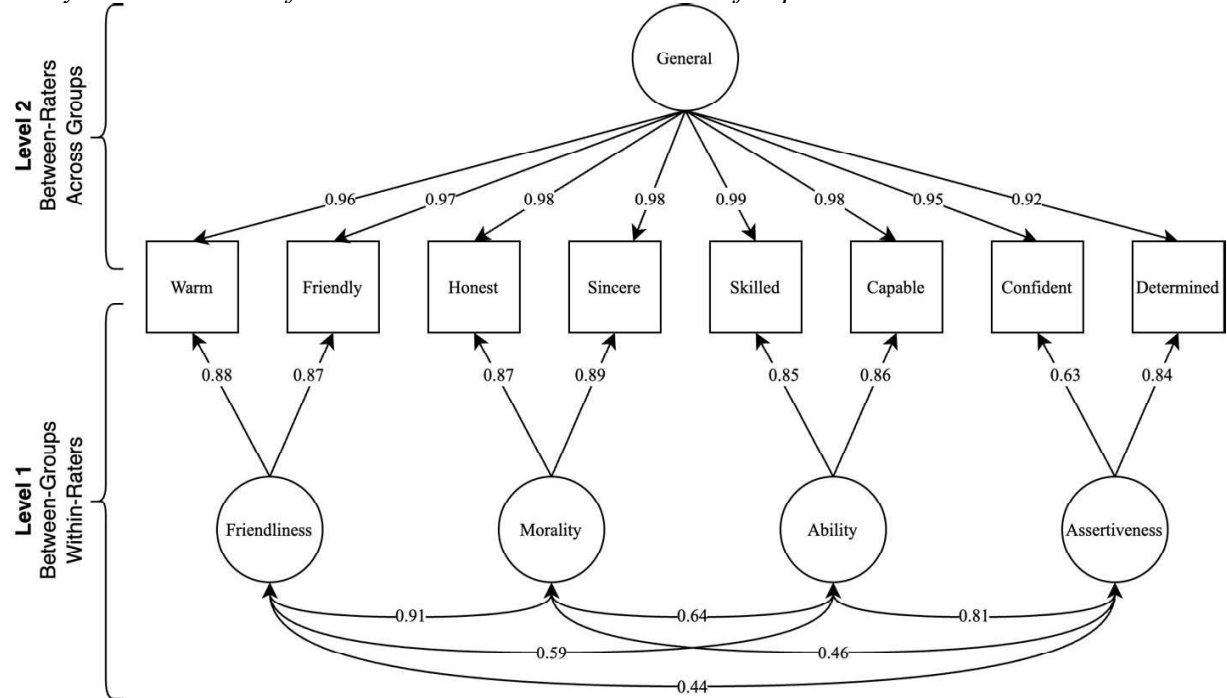


Figure S1.2

Study 1: Parameters of a Two-Items-Per-Facet Measure of Joint Evaluations

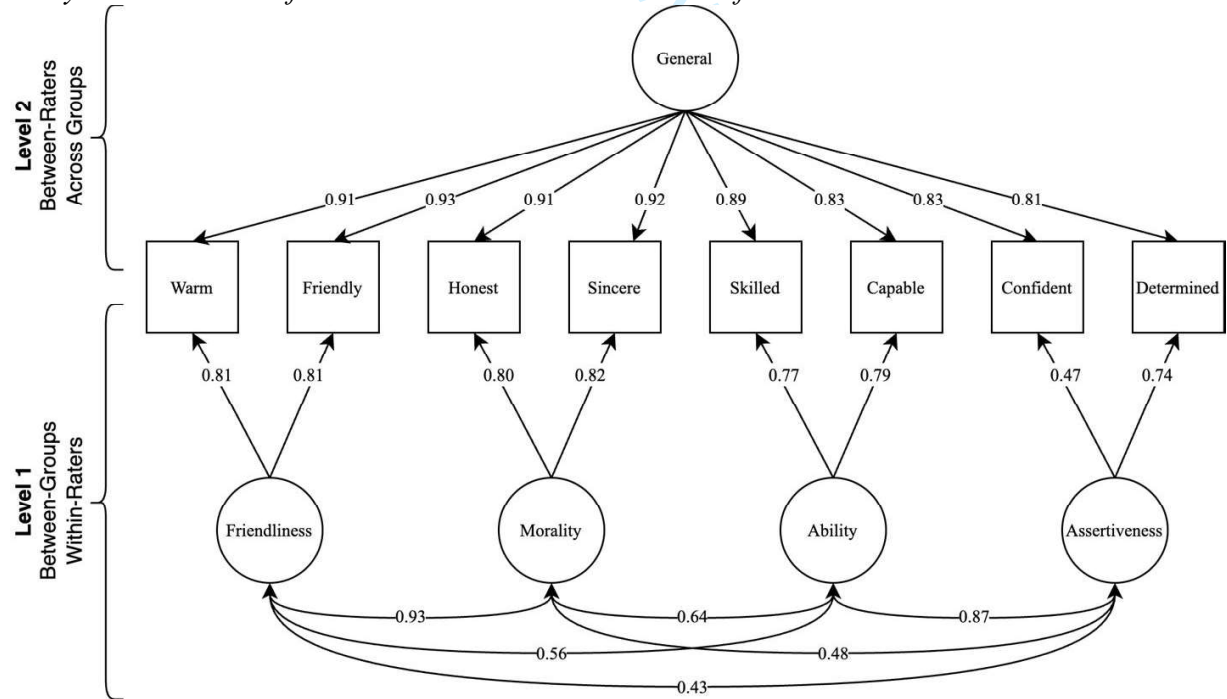


Figure S1.3

Study 1: Parameters of a Two-Items-Per-Facet Measure of Personal Evaluations

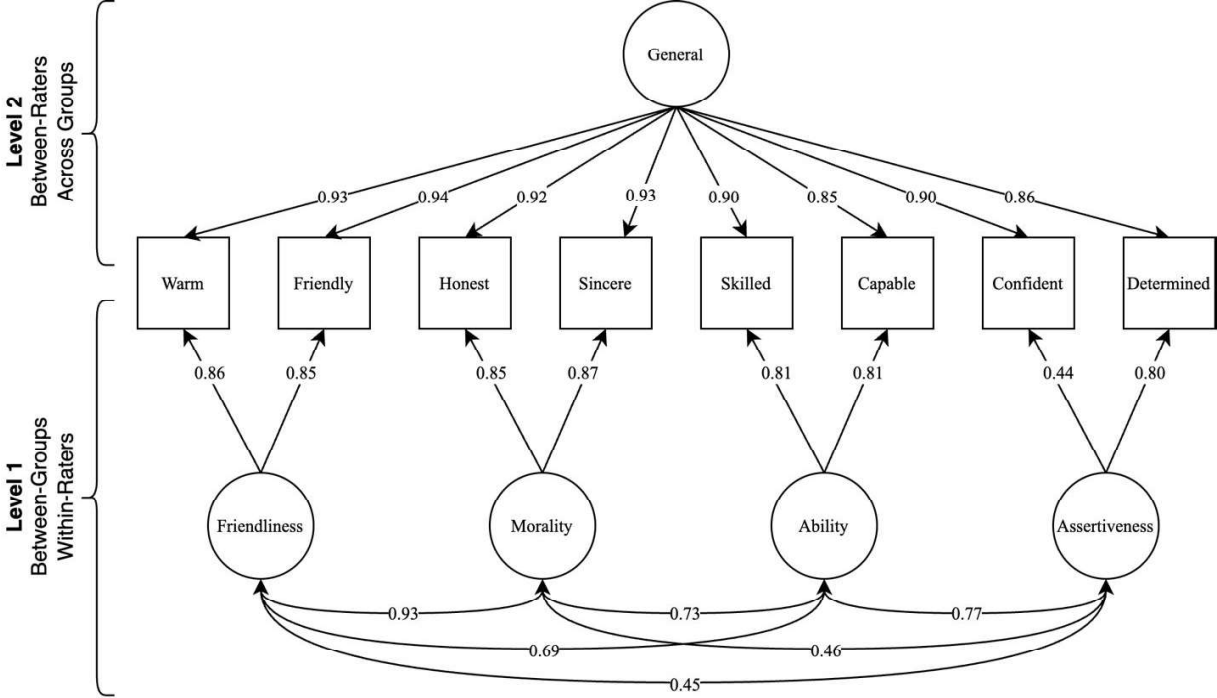


Figure S1.4

Study 1: Parameters of a Four-Items-Per-Dimension Measure of Separate Evaluations

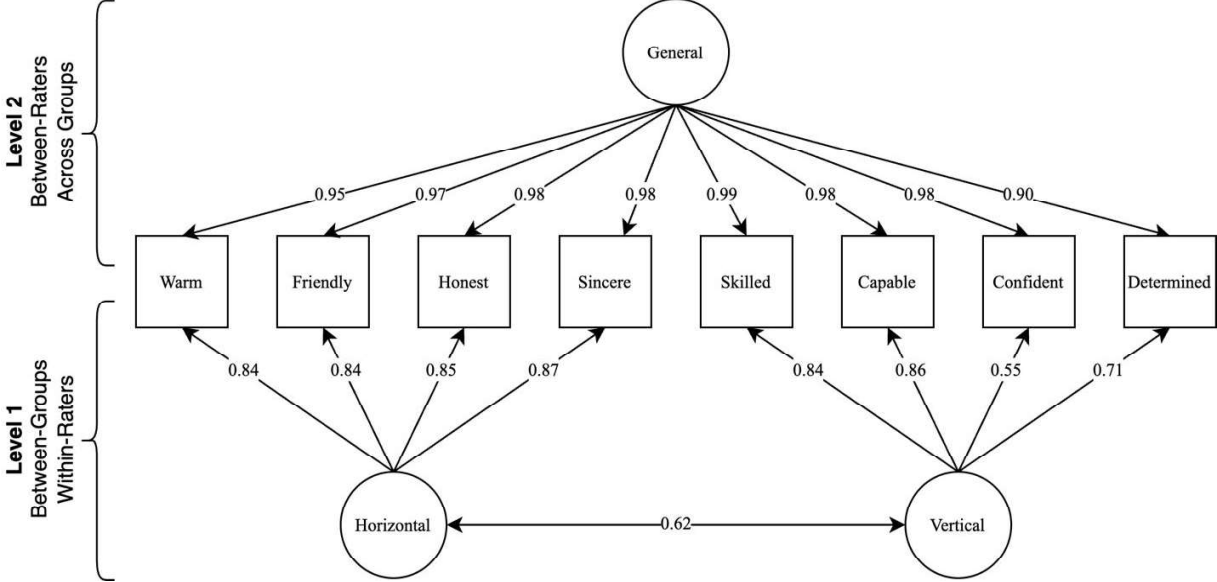


Figure S1.5

Study 1: Parameters of a Four-Items-Per-Dimension Measure of Joint Evaluations

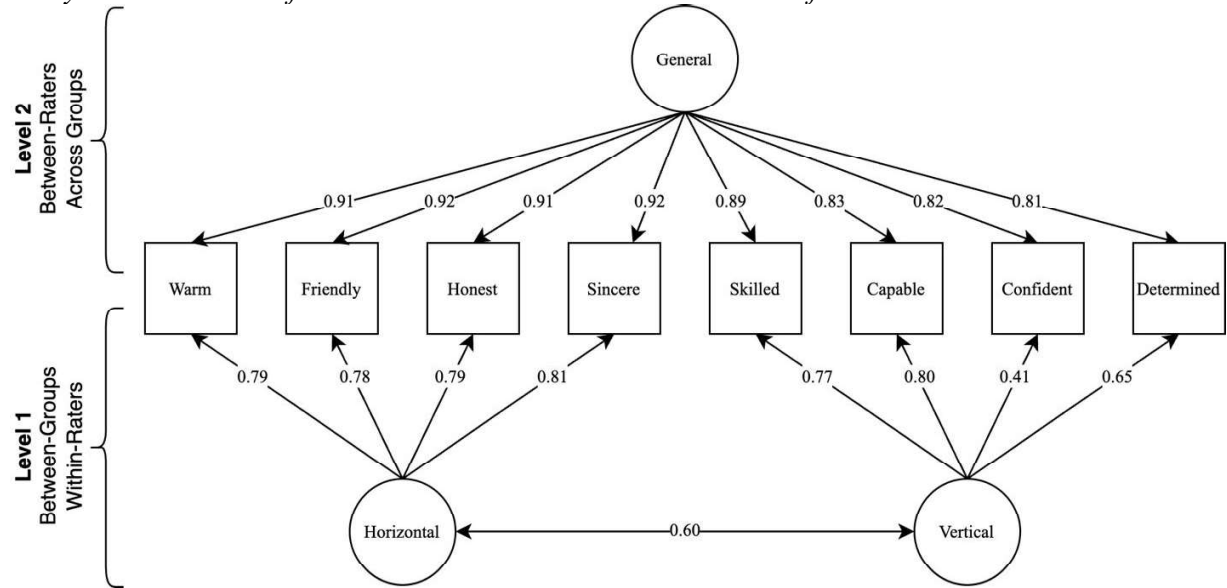


Figure S1.6

Study 1: Parameters of a Four-Items-Per-Dimension Measure of Personal Evaluations

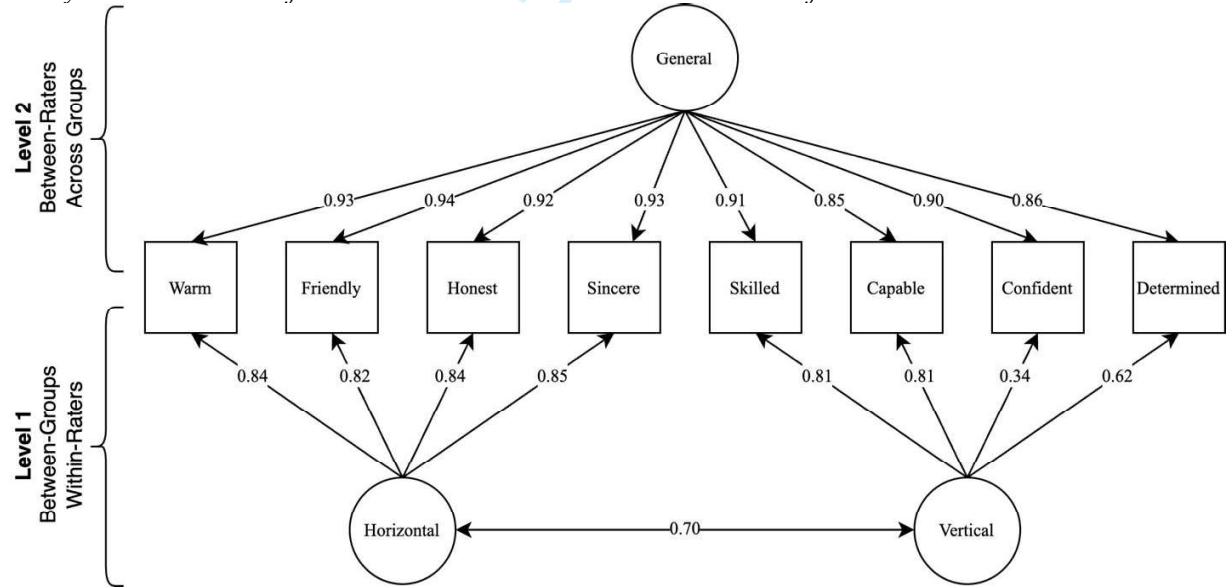


Figure S1.7

Study 1: Parameters of a Four-Items-Per-Dimension Measure of Cultural Evaluations

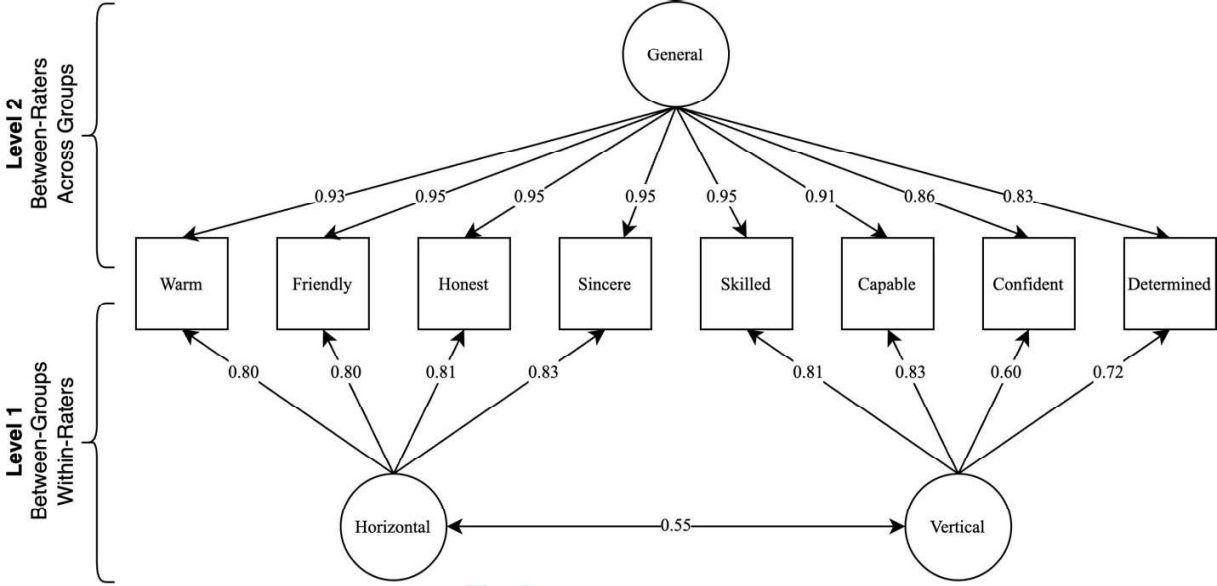


Figure S1.8

Study 1a: Parameters of a Two-Items-Per-Facet Measure of Separate Evaluations

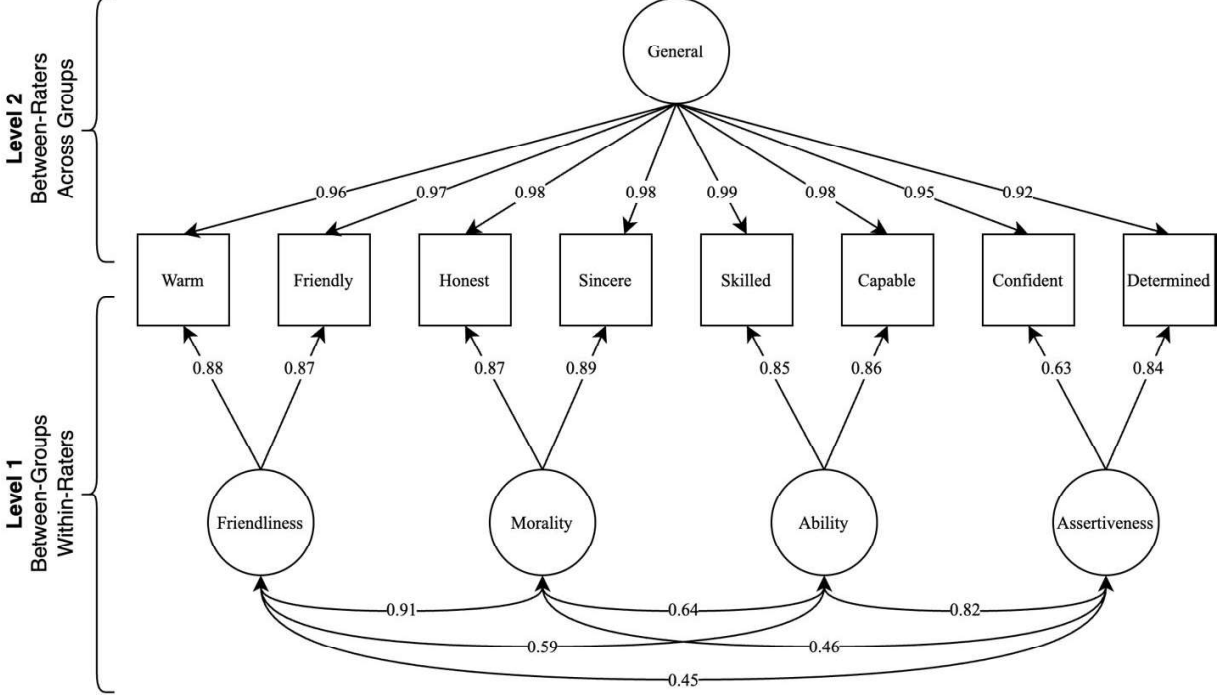


Figure S1.9

Study 1a: Parameters of a Two-Items-Per-Facet Measure of Joint Evaluations

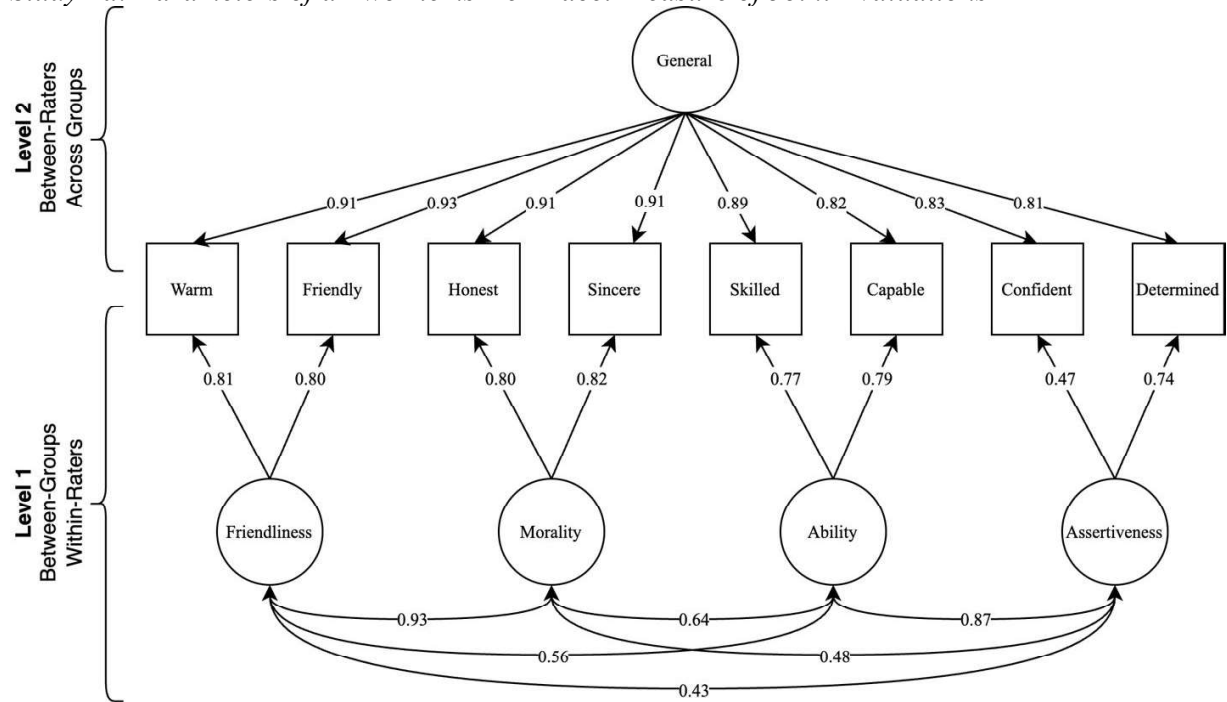


Figure S1.10

Study 1a: Parameters of a Two-Items-Per-Facet Measure of Personal Evaluations

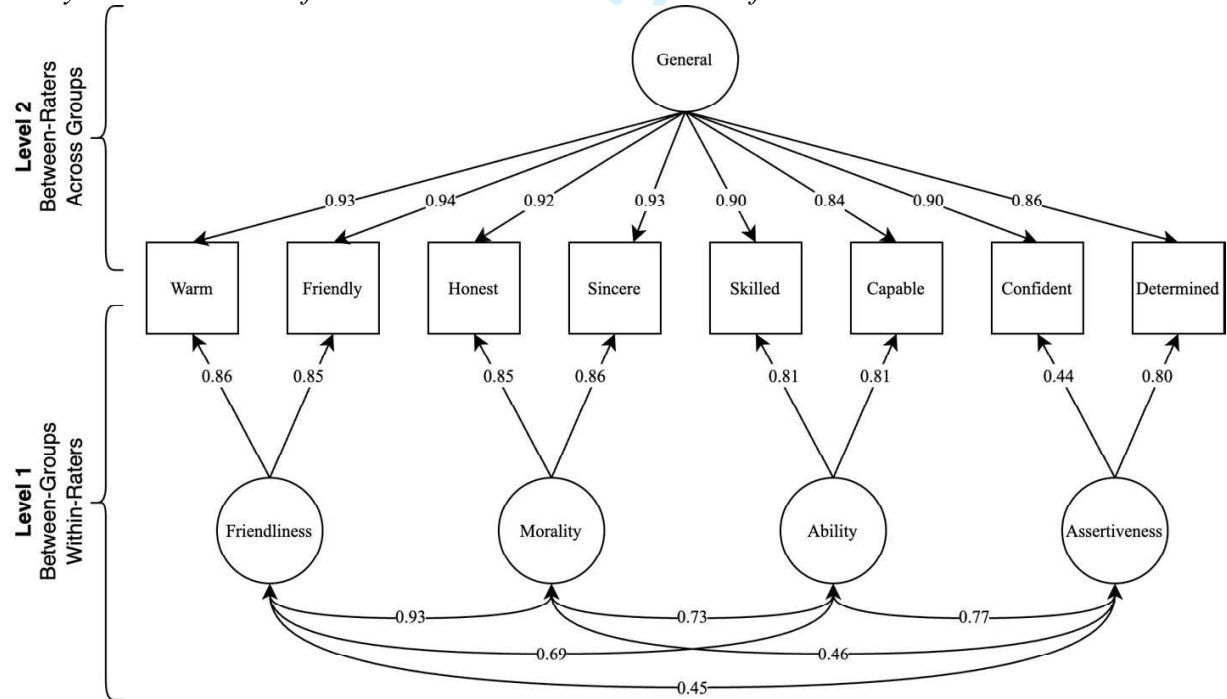


Figure S1.11

Study 1a: Parameters of a Two-Items-Per-Facet Measure of Cultural Evaluations

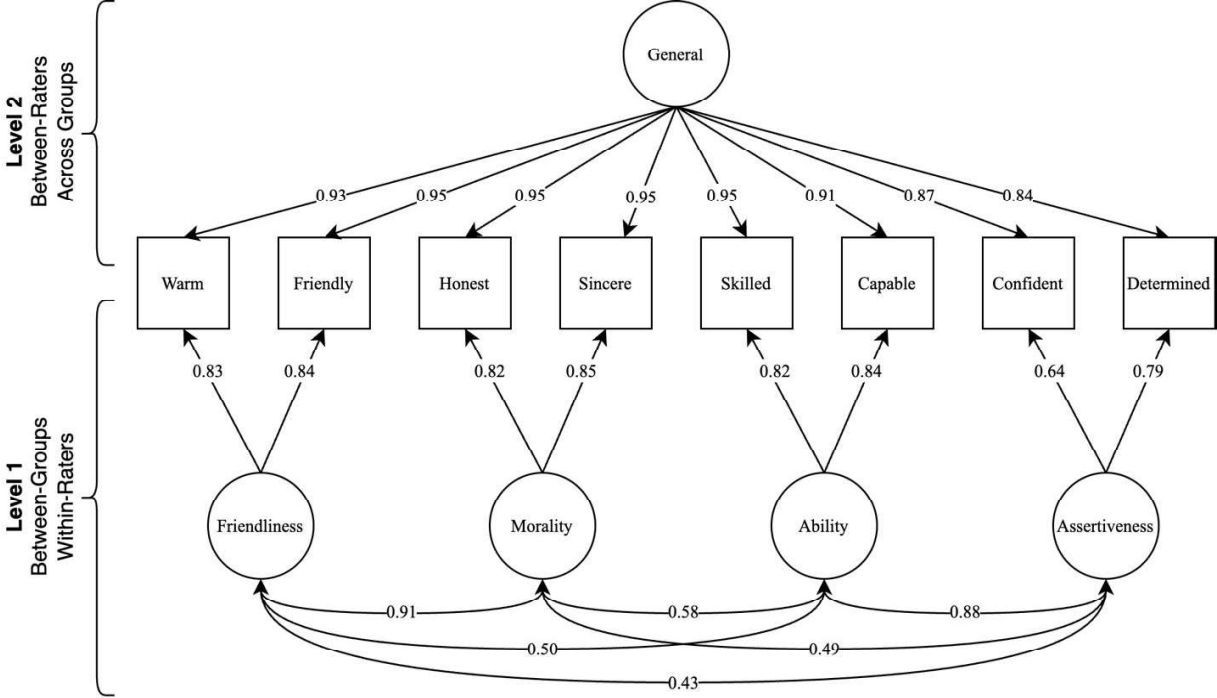


Figure S1.12

Study 1a: Parameters of a Four-Items-Per-Dimension Measure of Separate Evaluations

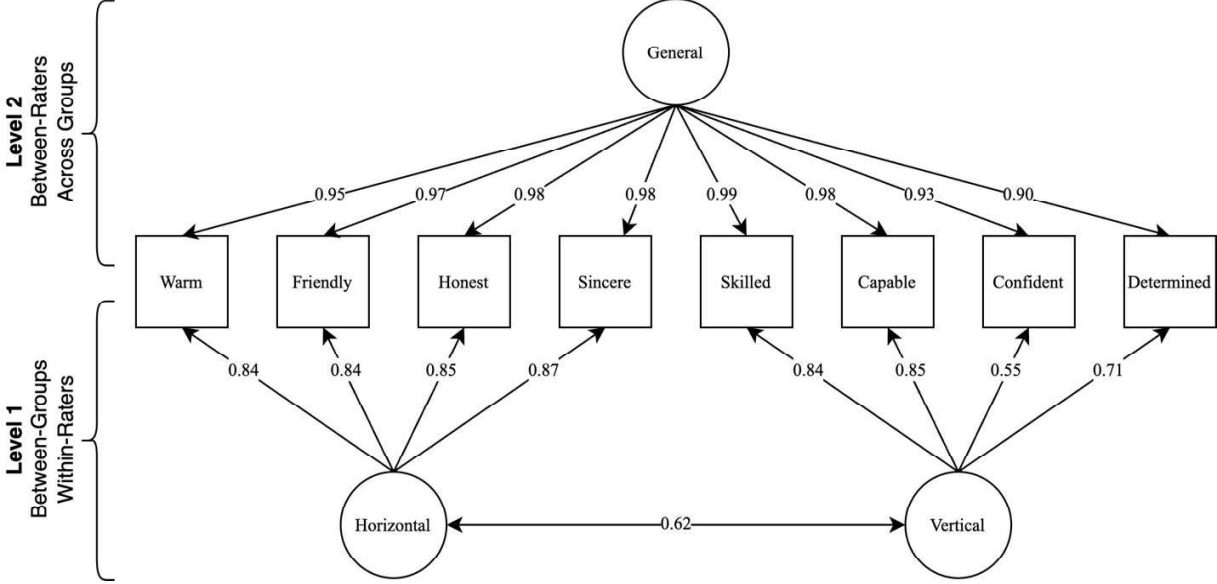


Figure S1.13

Study 1a: Parameters of a Four-Items-Per-Dimension Measure of Joint Evaluations

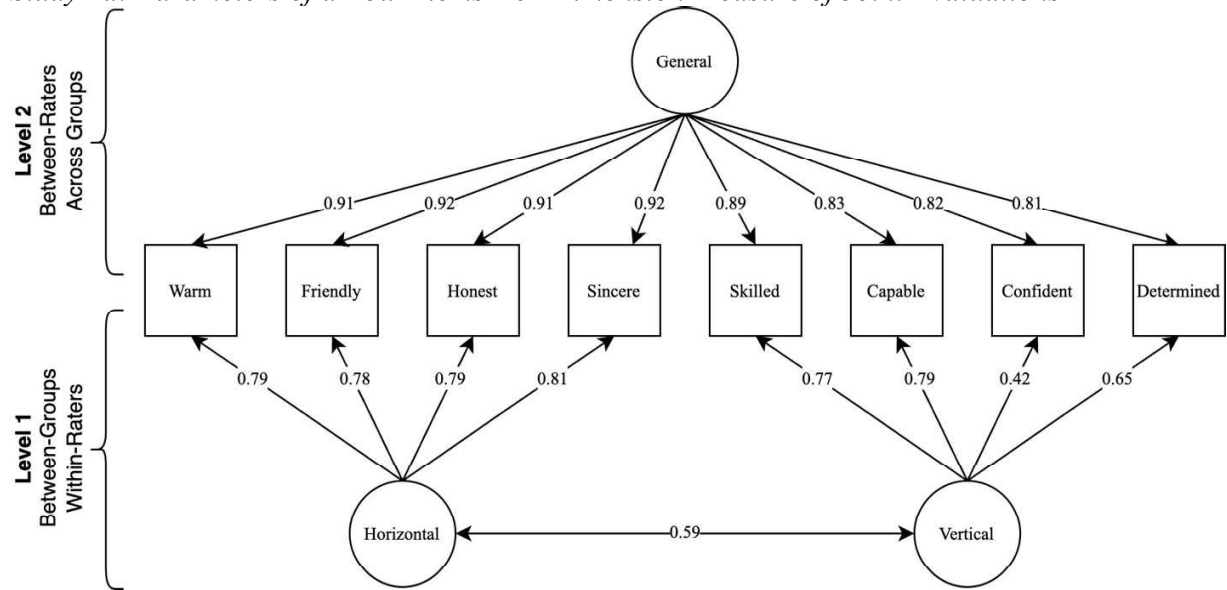


Figure S1.14

Study 1a: Parameters of a Four-Items-Per-Dimension Measure of Personal Evaluations

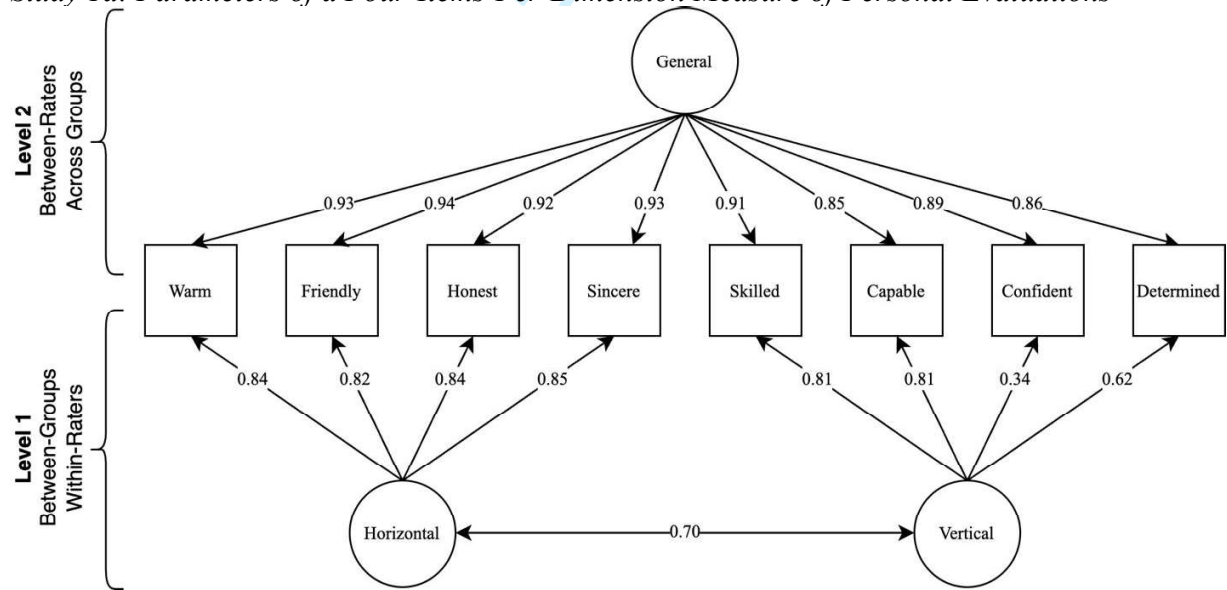


Figure S1.15
Study 1a: Parameters of a Four-Items-Per-Dimension Measure of Cultural Evaluations

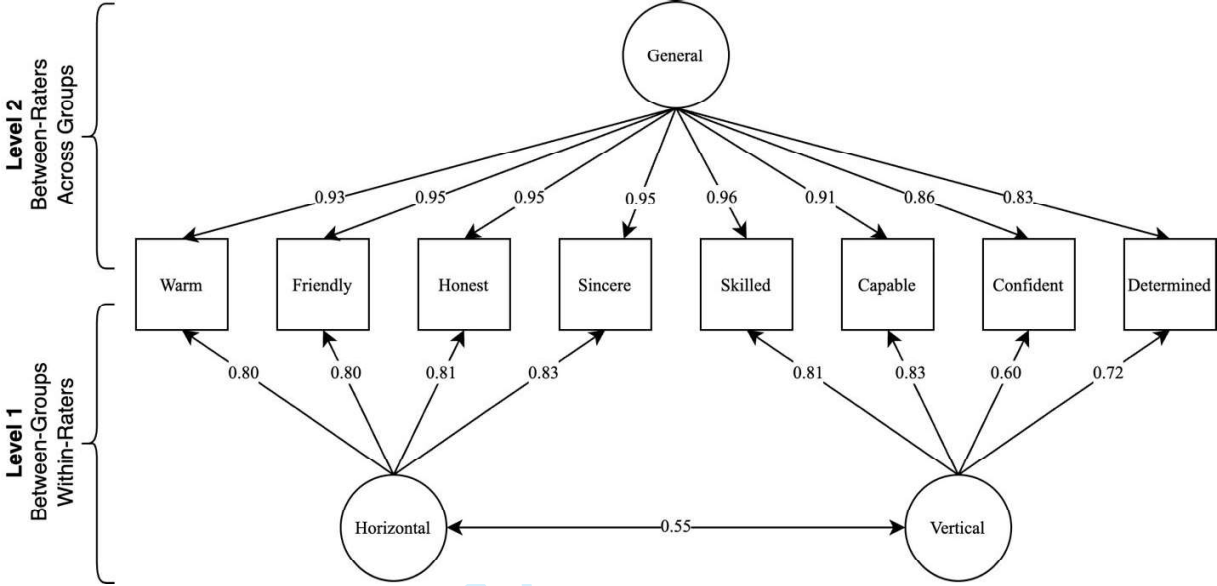


Figure S1.16
Study 1b: Parameters of a Two-Items-Per-Facet Measure of Separate Evaluations

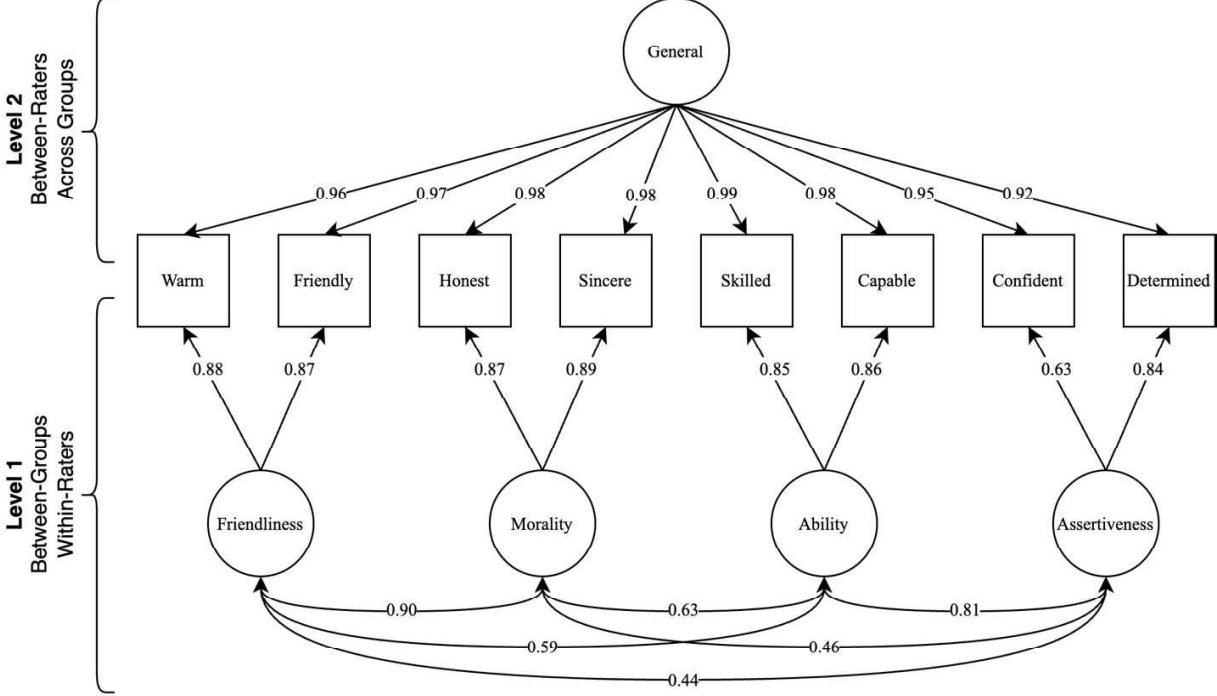


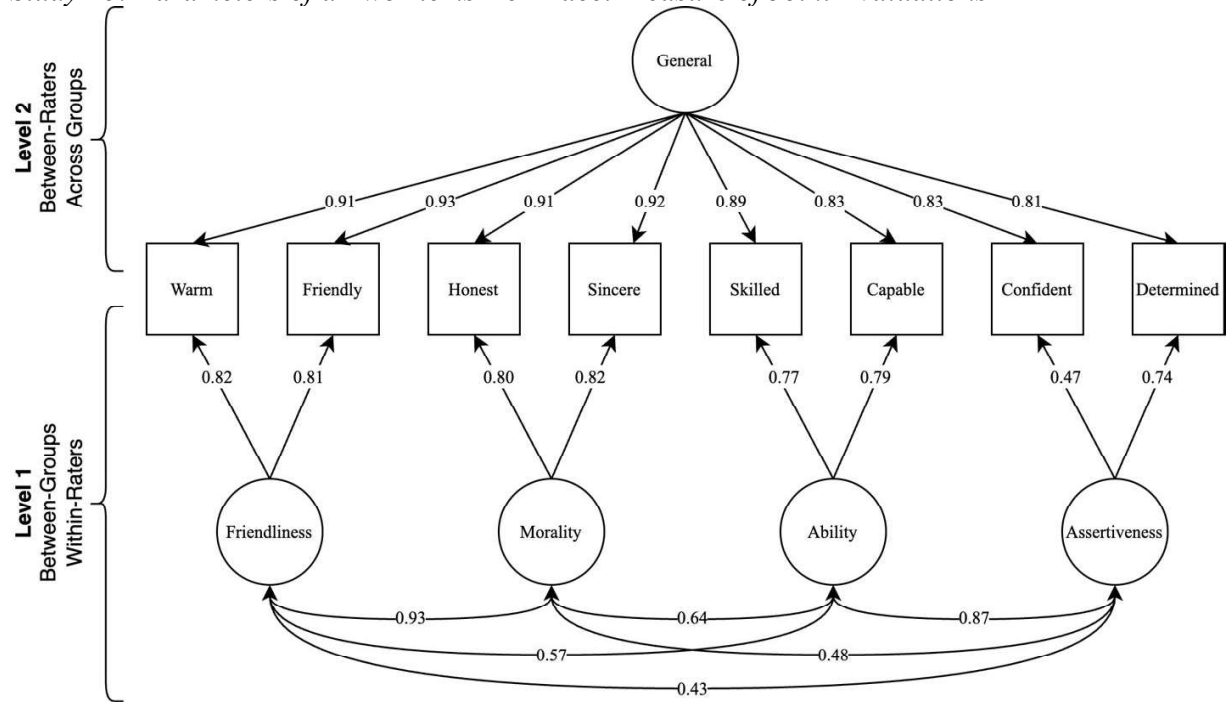
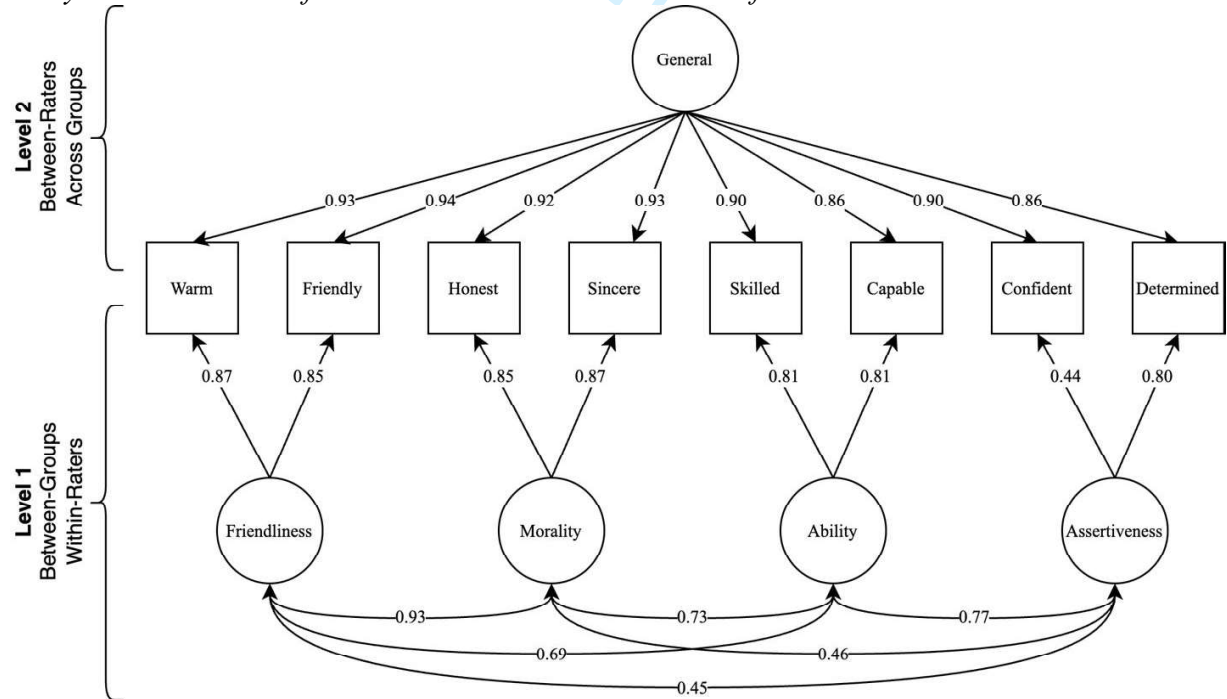
Figure S1.17*Study 1b: Parameters of a Two-Items-Per-Facet Measure of Joint Evaluations***Figure S1.18***Study 1b: Parameters of a Two-Items-Per-Facet Measure of Personal Evaluations*

Figure S1.19

Study 1b: Parameters of a Two-Items-Per-Facet Measure of Cultural Evaluations

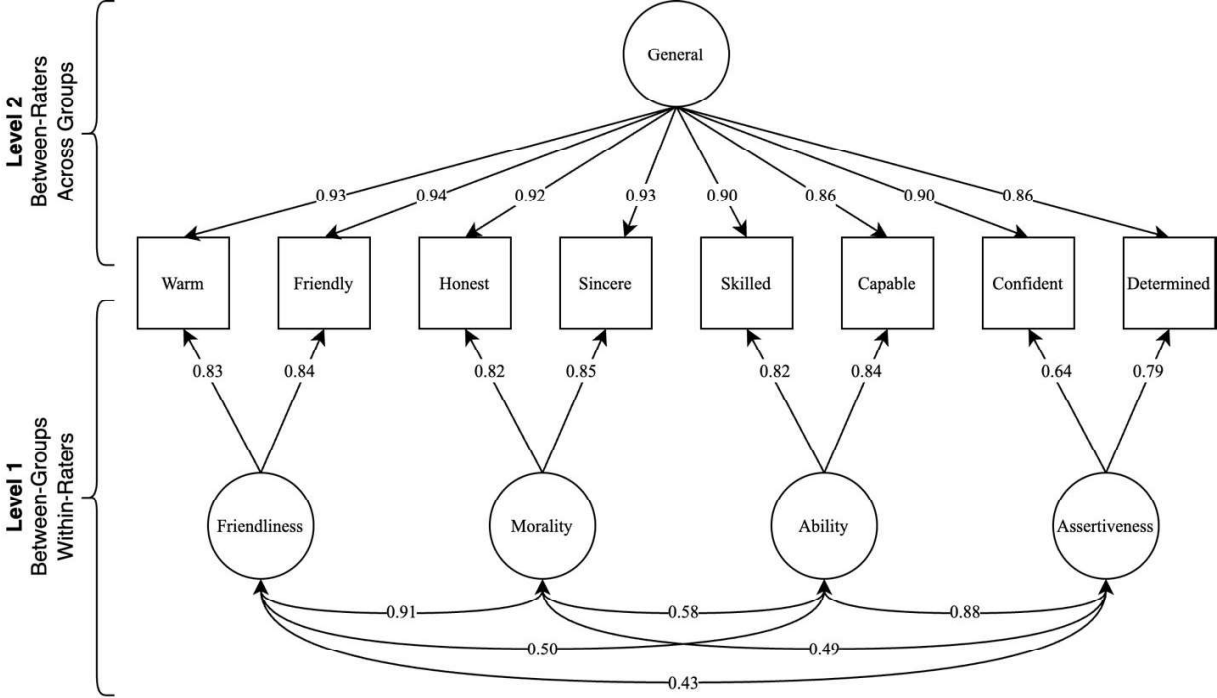


Figure S1.20

Study 1b: Parameters of a Four-Items-Per-Dimension Measure of Separate Evaluations

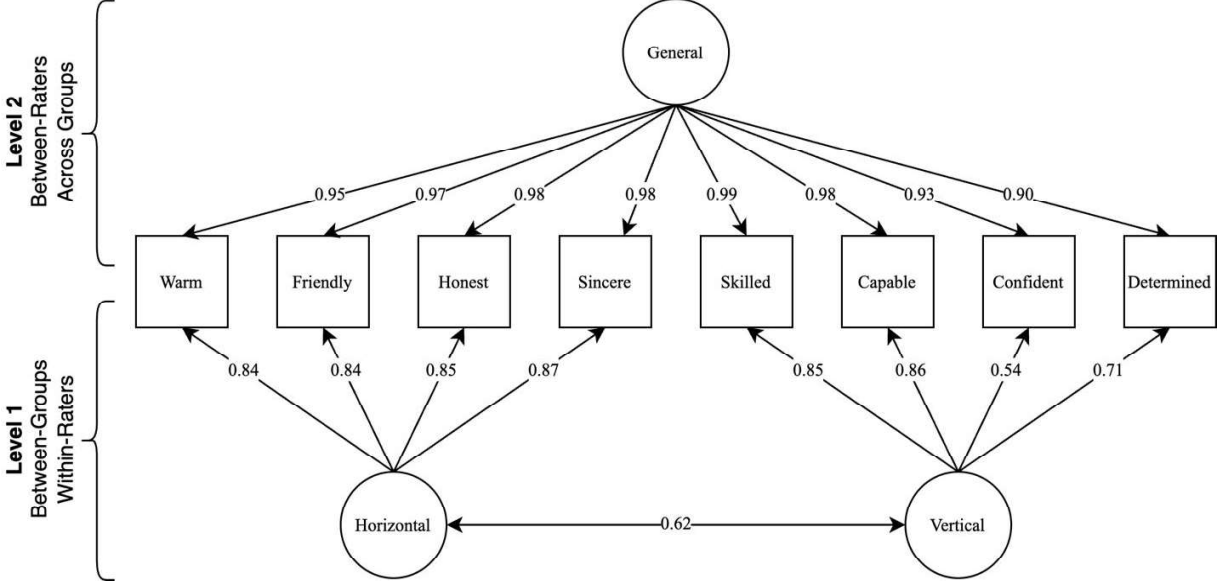


Figure S1.21

Study 1b: Parameters of a Four-Items-Per-Dimension Measure of Joint Evaluations

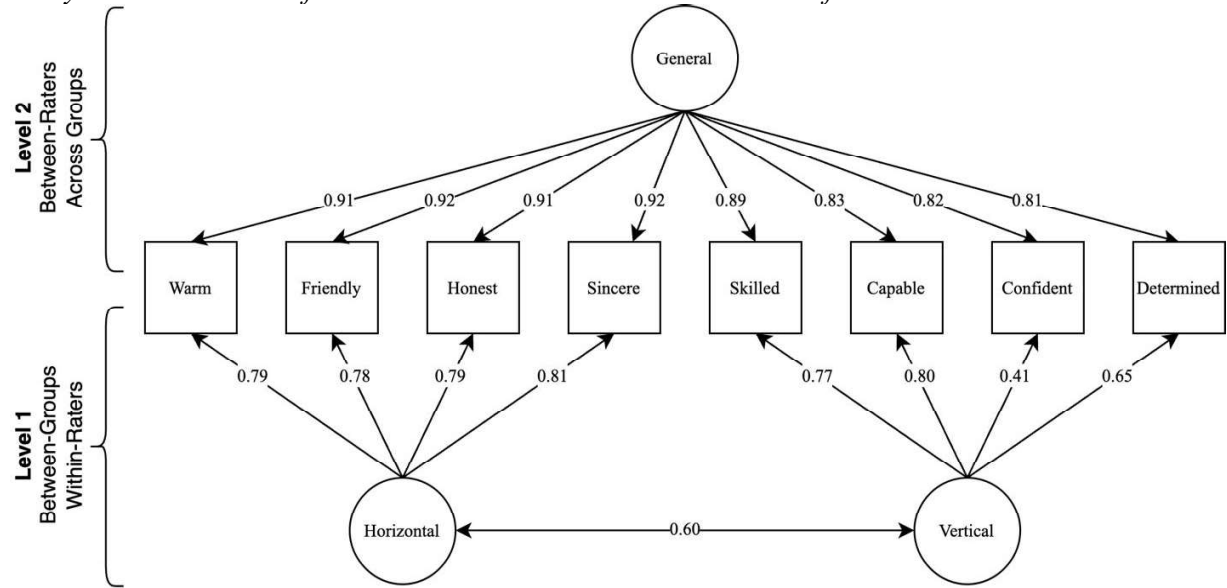


Figure S1.22

Study 1b: Parameters of a Four-Items-Per-Dimension Measure of Personal Evaluations

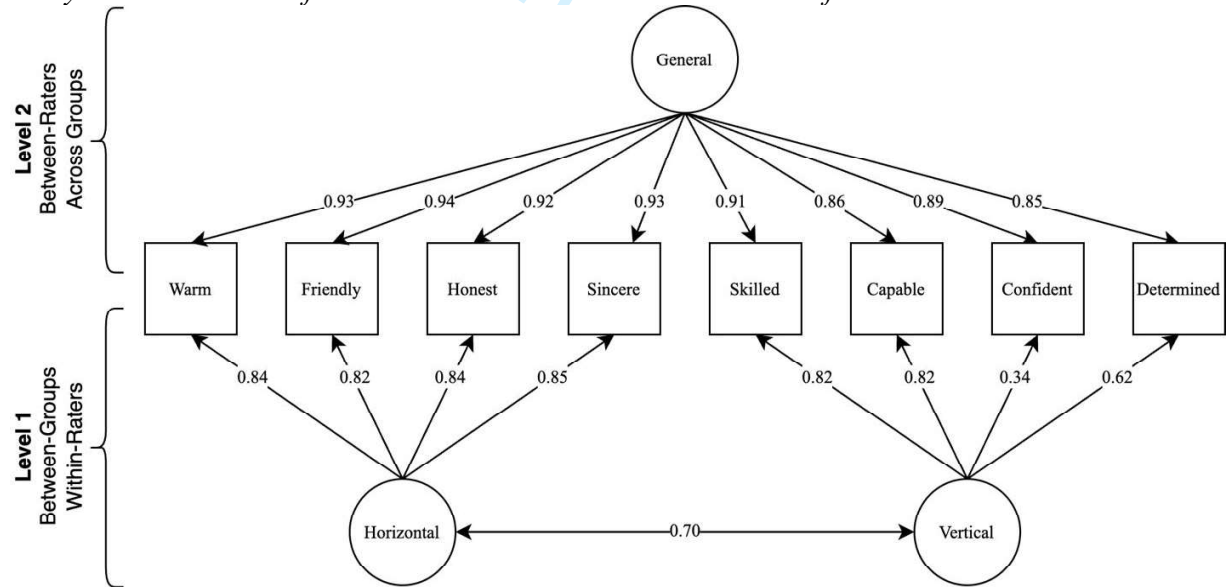


Figure S1.23

Study 1b: Parameters of a Four-Items-Per-Dimension Measure of Cultural Evaluations

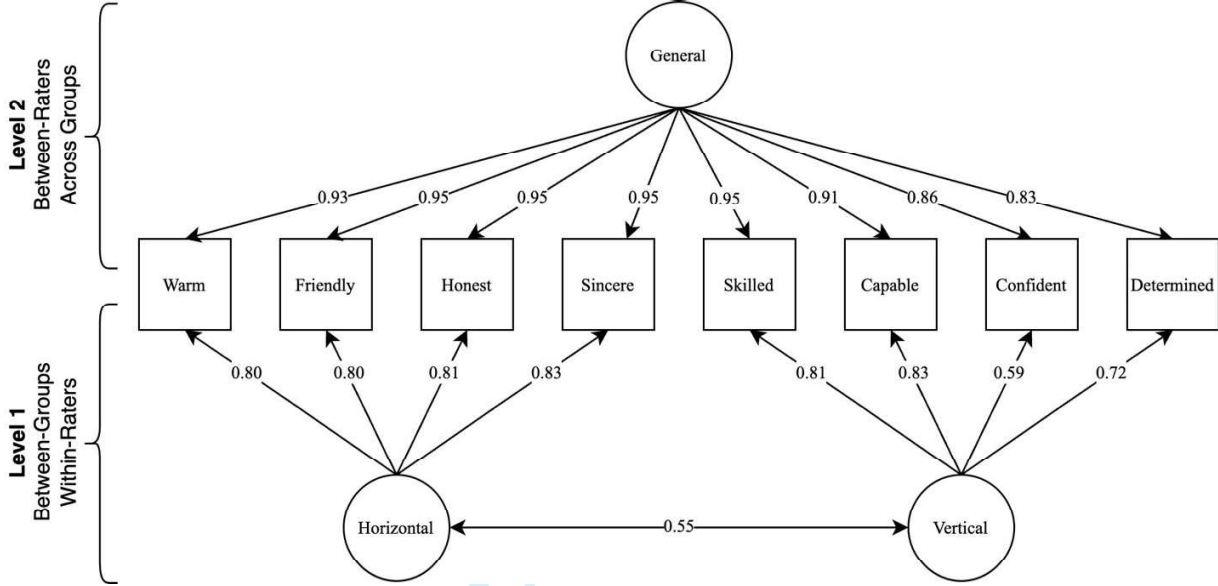


Table S1.1

Studies 1a and 1b: Satisfactory and superior fit of a two-items-per-facet model

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; df = 34						
Separate evaluations	1,649	0.99	0.03	0.06	885,536	
Joint evaluations	2,021	0.99	0.04	0.07	935,804	
Personal evaluations	1,938	0.99	0.04	0.08	876,224	
Cultural evaluations	1,042	0.99	0.03	0.05	944,668	
Two-items-per-facet & Two-facets-per-dimension; df = 35						
Separate evaluations	1,670	0.99	0.03	0.06	885,555	<.001
Joint evaluations	2,023	0.99	0.04	0.07	935,804	0.227
Personal evaluations	1,966	0.99	0.04	0.08	876,250	<.001
Cultural evaluations	1,046	0.99	0.03	0.05	944,670	0.041
Four-items-per-dimension; df = 39						
Separate evaluations	7,923	0.96	0.07	0.08	891,800	<.001
Joint evaluations	3,790	0.97	0.05	0.08	937,562	<.001
Personal evaluations	5,547	0.97	0.06	0.09	879,823	<.001
Cultural evaluations	4,077	0.98	0.05	0.06	947,692	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

Table S1.2*Study 1a: Fit Measures*

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; df = 34						
Separate evaluations	831	0.99	0.03	0.06	450,219	
Joint evaluations	1,022	0.99	0.04	0.07	476,422	
Personal evaluations	980	0.99	0.04	0.08	445,190	
Cultural evaluations	526	0.99	0.03	0.05	481,206	
Two-items-per-facet & Two-facets-per-dimension; df = 35						
Separate evaluations	841	0.99	0.03	0.06	450,226	0.002
Joint evaluations	1,023	0.99	0.04	0.07	476,421	0.395
Personal evaluations	995	0.99	0.04	0.08	445,202	<.001
Cultural evaluations	527	0.99	0.03	0.05	481,206	0.232
Four-items-per-dimension; df = 39						
Separate evaluations	3,962	0.96	0.07	0.08	453,339	<.001
Joint evaluations	1,921	0.97	0.05	0.08	477,311	<.001
Personal evaluations	2,791	0.97	0.06	0.09	446,991	<.001
Cultural evaluations	2,059	0.98	0.05	0.06	482,729	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

Table S1.3*Study 1b: Fit Measures*

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; df = 34						
Separate evaluations	819	0.99	0.03	0.06	435,405	
Joint evaluations	1,001	0.99	0.04	0.07	459,470	
Personal evaluations	960	0.99	0.04	0.08	431,120	
Cultural evaluations	517	0.99	0.03	0.05	463,550	
Two-items-per-facet & Two-facets-per-dimension; df = 35						
Separate evaluations	830	0.99	0.03	0.06	435,414	<.001
Joint evaluations	1,002	0.99	0.04	0.07	459,468	0.389
Personal evaluations	973	0.99	0.04	0.08	431,131	<.001
Cultural evaluations	520	0.99	0.03	0.05	463,551	0.088
Four-items-per-dimension; df = 39						
Separate evaluations	3,963	0.96	0.07	0.08	458,540	<.001
Joint evaluations	1,871	0.97	0.05	0.08	460,329	<.001
Personal evaluations	2,759	0.97	0.06	0.09	432,909	<.001
Cultural evaluations	2,020	0.98	0.05	0.06	465,042	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

Table S1.4
Study 1a: Correlations Between the Four Facet Scales Centered Within the Raters

Separate					Joint				
Friendliness -	1.00	0.79	0.51	0.35	1.00	0.74	0.44	0.24	
Morality -	0.79	1.00	0.55	0.36	0.74	1.00	0.50	0.27	
Ability -	0.51	0.55	1.00	0.63	0.44	0.50	1.00	0.56	
Assertiveness -	0.35	0.36	0.63	1.00	0.24	0.27	0.56	1.00	
Personal					Cultural				
Friendliness -	1.00	0.78	0.56	0.28	1.00	0.75	0.41	0.31	
Morality -	0.78	1.00	0.60	0.27	0.75	1.00	0.48	0.36	
Ability -	0.56	0.60	1.00	0.51	0.41	0.48	1.00	0.66	
Assertiveness -	0.28	0.27	0.51	1.00	0.31	0.36	0.66	1.00	
	Friendliness	Morality	Ability	Assertiveness	Friendliness	Morality	Ability	Assertiveness	

Table S1.5*Study 1b: Correlations Between the Four Facet Scales Centered Within the Raters*

Separate					Joint				
Friendliness -	1.00	0.79	0.51	0.35	1.00	0.74	0.44	0.24	
Morality -	0.79	1.00	0.55	0.36	0.74	1.00	0.51	0.27	
Ability -	0.51	0.55	1.00	0.63	0.44	0.51	1.00	0.56	
Assertiveness -	0.35	0.36	0.63	1.00	0.24	0.27	0.56	1.00	

Personal					Cultural				
Friendliness -	1.00	0.78	0.57	0.28	1.00	0.75	0.41	0.31	
Morality -	0.78	1.00	0.60	0.27	0.75	1.00	0.48	0.36	
Ability -	0.57	0.60	1.00	0.51	0.41	0.48	1.00	0.66	
Assertiveness -	0.28	0.27	0.51	1.00	0.31	0.36	0.66	1.00	
	Friendliness	Morality	Ability	Assertiveness	Friendliness	Morality	Ability	Assertiveness	

Table S1.6

Study 1: Tabular confirmatory factor analyses, averaging across the 20 target groups

	X ²	CFI	RMSEA	SRMR	AIC
Two-items-per-facet; <i>df</i> = 34					
Separate evaluations	94.97	0.99	0.05	0.02	43,121
Joint evaluations	68.06	0.99	0.04	0.02	47,136
Personal evaluations	86.39	0.99	0.05	0.02	43,814
Cultural evaluations	69.72	1.00	0.04	0.02	46,446
Four-items-per-dimension; <i>df</i> = 39					
Separate evaluations	566.36	0.95	0.12	0.03	43,582
Joint evaluations	193.03	0.98	0.06	0.05	47,251
Personal evaluations	310.37	0.97	0.08	0.04	44,028
Cultural evaluations	335.85	0.97	0.09	0.04	46,703

Note. These were factor analyses that we ran separately for each group, with the perceivers as rows, the eight or twenty items as columns, and the perceivers' ratings of the target in a given factor analysis on a given item as the data points. The table reports fit indices averaged across all 20 groups that we examined in Study 1.

Table S1.7

Study 1a: Tabular confirmatory factor analyses, across the 20 target groups

	X ²	CFI	RMSEA	SRMR	AIC
Two-items-per-facet; <i>df</i> = 34					
Separate evaluations	48.52	0.99	0.04	0.02	24,022
Joint evaluations	34.16	0.99	0.03	0.02	21,947
Personal evaluations	44.17	0.99	0.04	0.02	22,293
Cultural evaluations	35.31	1.00	0.03	0.02	23,678
Four-items-per-dimension; <i>df</i> = 39					
Separate evaluations	285.60	0.95	0.11	0.05	24,076
Joint evaluations	98.16	0.98	0.06	0.03	22,174
Personal evaluations	157.53	0.97	0.08	0.04	23,803
Cultural evaluations	170.51	0.97	0.08	0.04	22,397

Note. These were factor analyses that we ran separately for each group, with the perceivers as rows, the eight or twenty items as columns, and the perceivers' ratings of the target in a given factor analysis on a given item as the data points. The table reports fit indices averaged across all 20 groups that we examined in Study 1a.

Table S1.8*Study 1b: Tabular confirmatory factor analyses, across the 20 target groups*

	X ²	CFI	RMSEA	SRMR	AIC
Two-items-per-facet; $df = 34$					
Separate evaluations	47.38	0.99	0.04	0.02	21,216
Joint evaluations	34.34	0.99	0.03	0.02	23,157
Personal evaluations	42.97	0.99	0.04	0.02	21,564
Cultural evaluations	34.89	1.00	0.03	0.02	22,812
Four-items-per-dimension; $df = 39$					
Separate evaluations	281.99	0.95	0.11	0.05	21,441
Joint evaluations	95.51	0.98	0.06	0.03	23,208
Personal evaluations	153.76	0.97	0.08	0.04	21,664
Cultural evaluations	166.07	0.97	0.08	0.04	22,933

Note. These were factor analyses that we ran separately for each group, with the perceivers as rows, the eight or twenty items as columns, and the perceivers' ratings of the target in a given factor analysis on a given item as the data points. The table reports fit indices averaged across all 20 groups that we examined in Study 1b.

Study 2

Figure S2.1

Study 2: Parameters of a Two-Items-Per-Facet Measure of Separate Evaluations

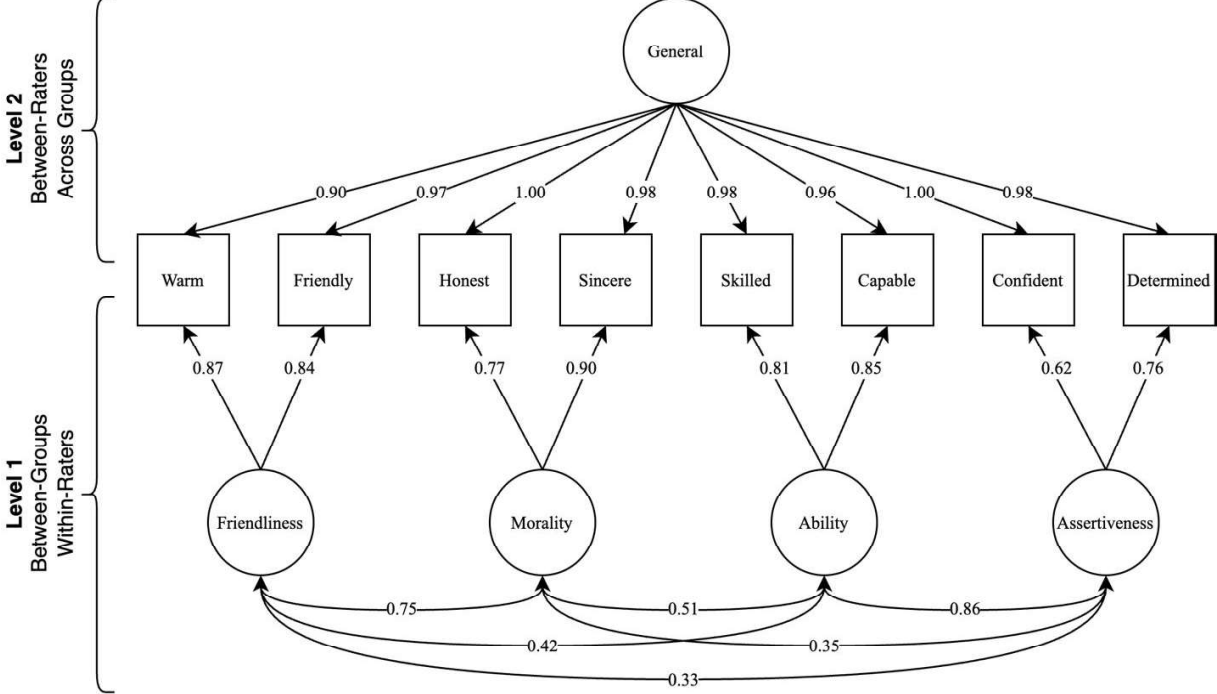


Figure S2.2

Study 2: Parameters of a Two-Items-Per-Facet Measure of Joint Evaluations

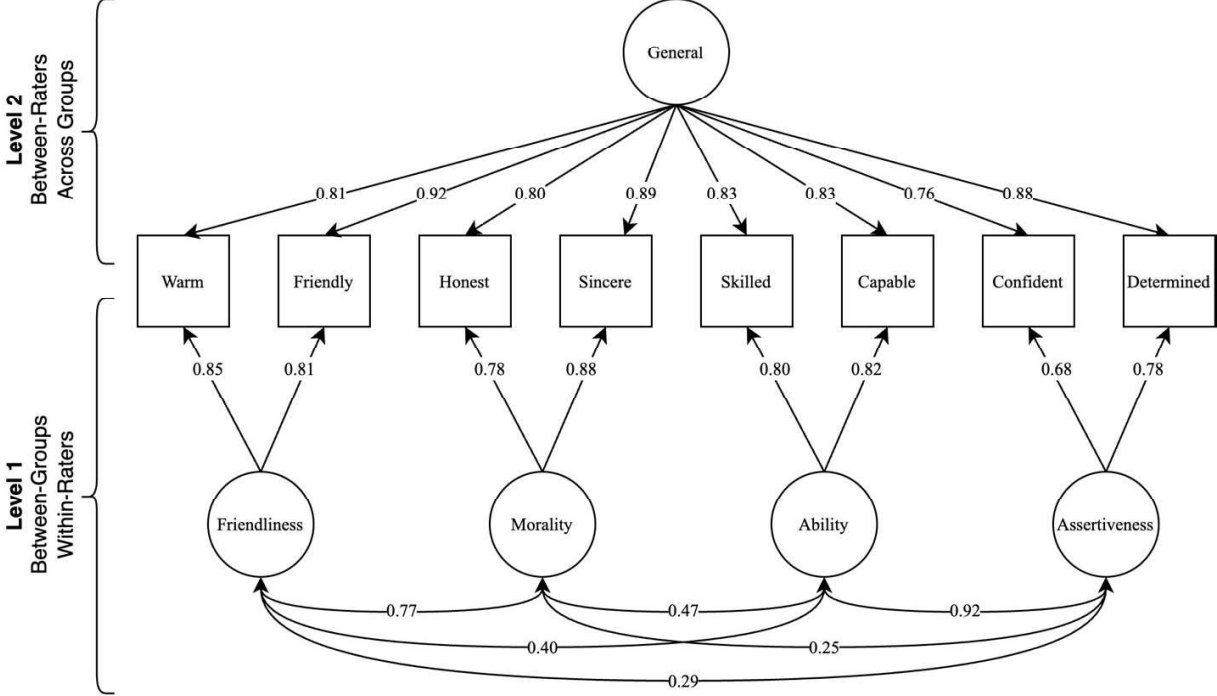


Figure S2.3

Study 2: Parameters of a Two-Items-Per-Facet Measure of Personal Evaluations

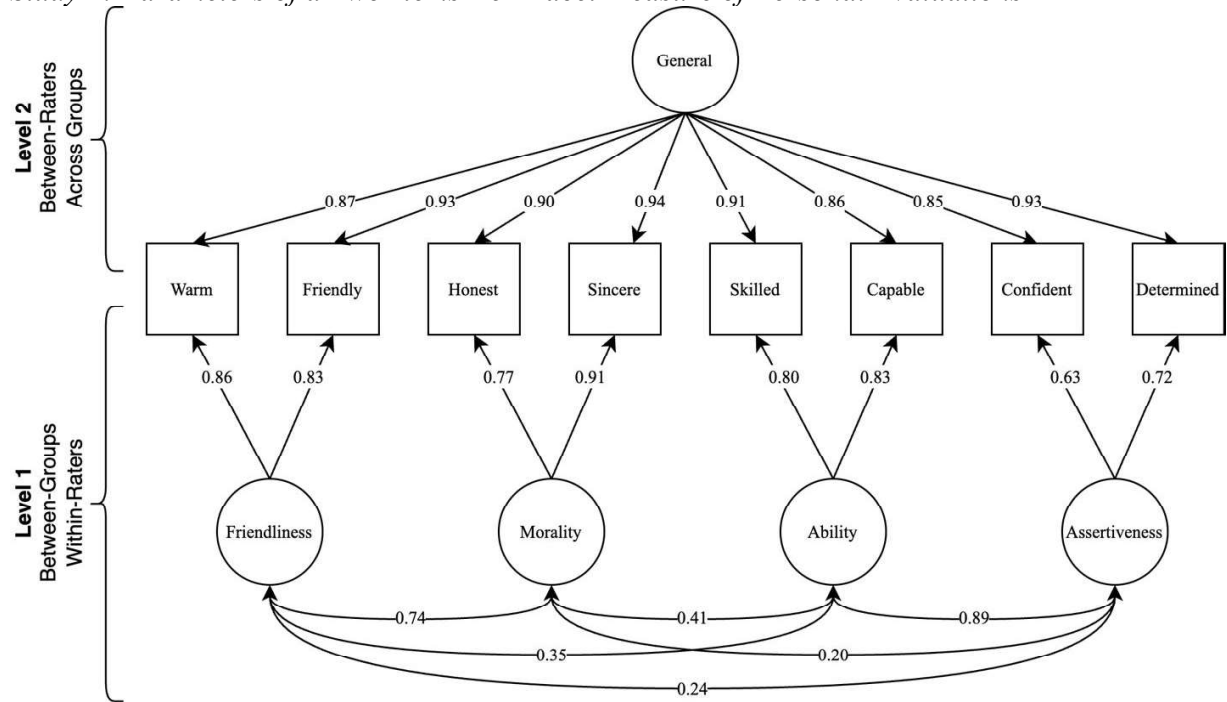


Figure S2.4

Study 2: Parameters of a Four-Items-Per-Dimension Measure of Separate Evaluations

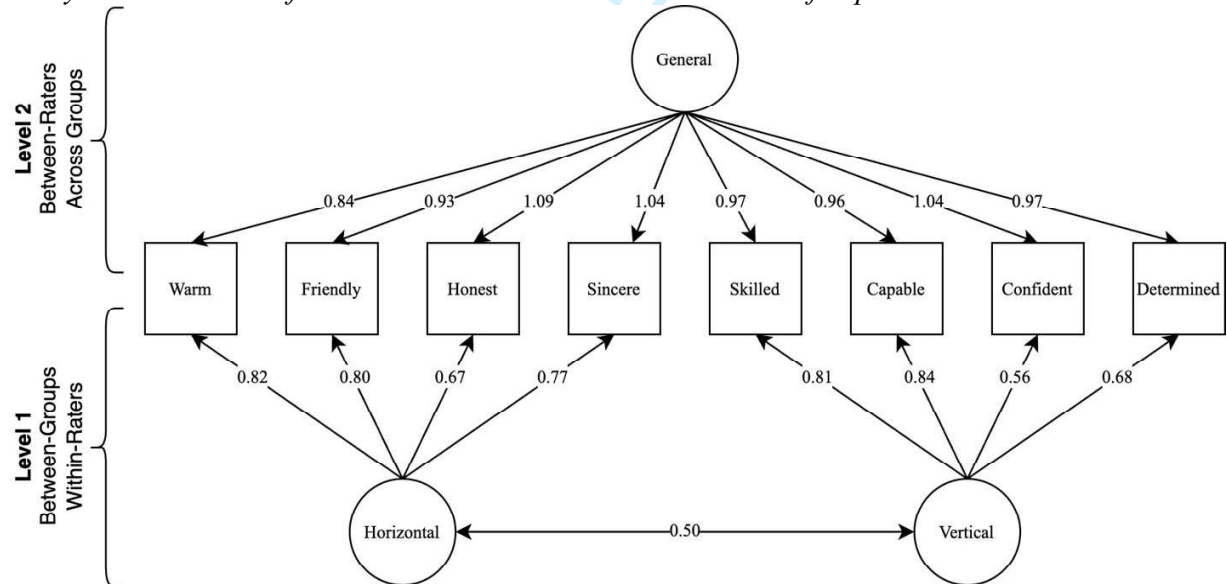


Figure S2.5

Study 2: Parameters of a Four-Items-Per-Dimension Measure of Joint Evaluations

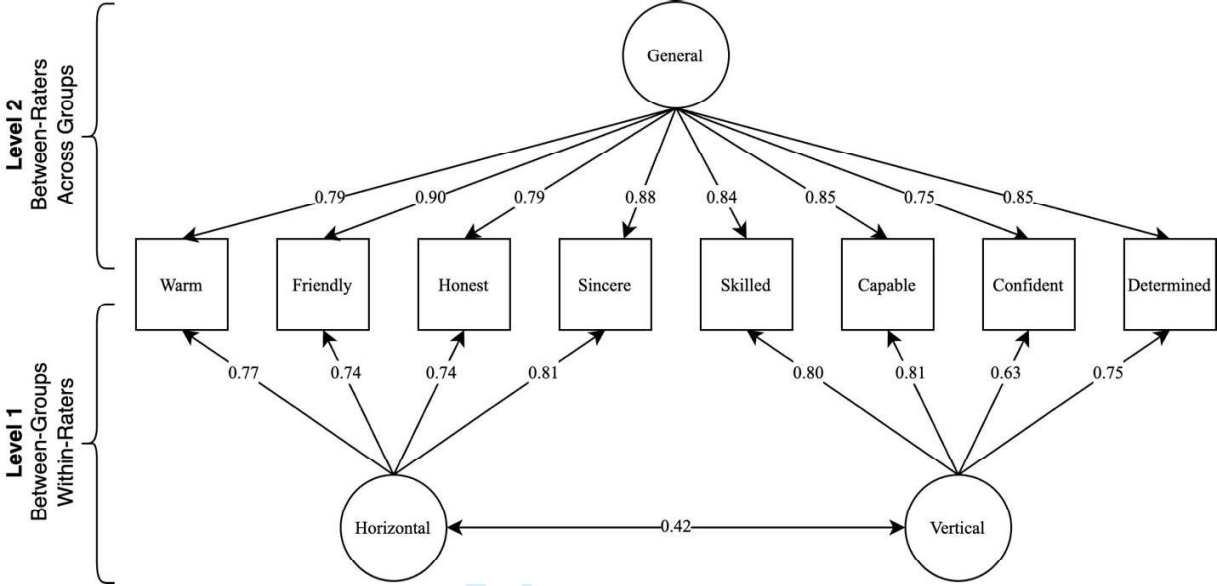


Figure S2.6

Study 2: Parameters of a Four-Items-Per-Dimension Measure of Personal Evaluations

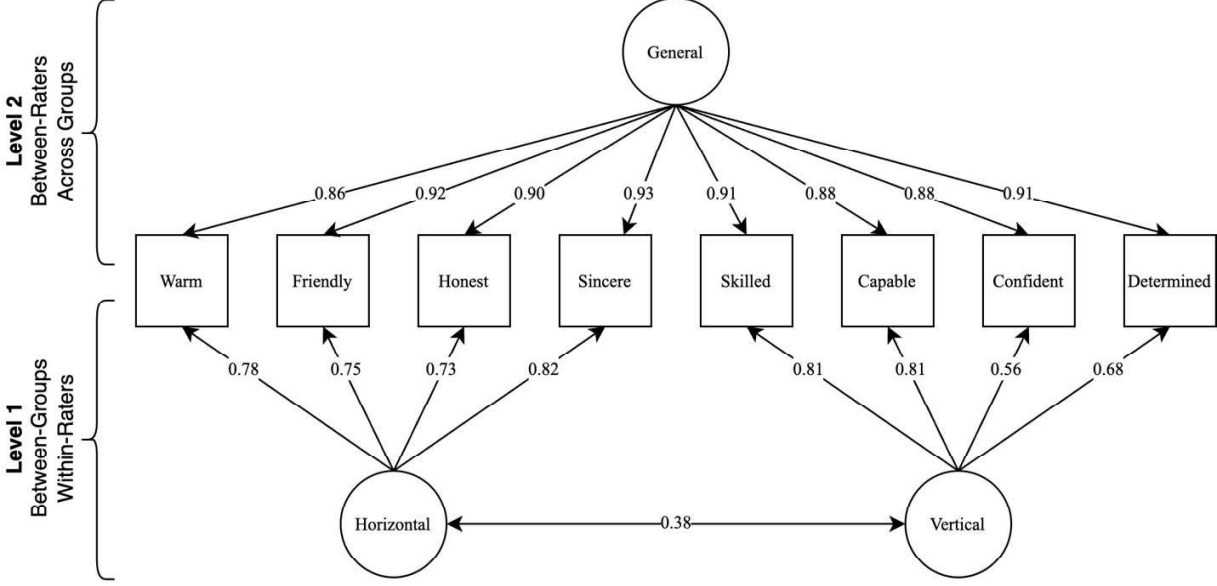
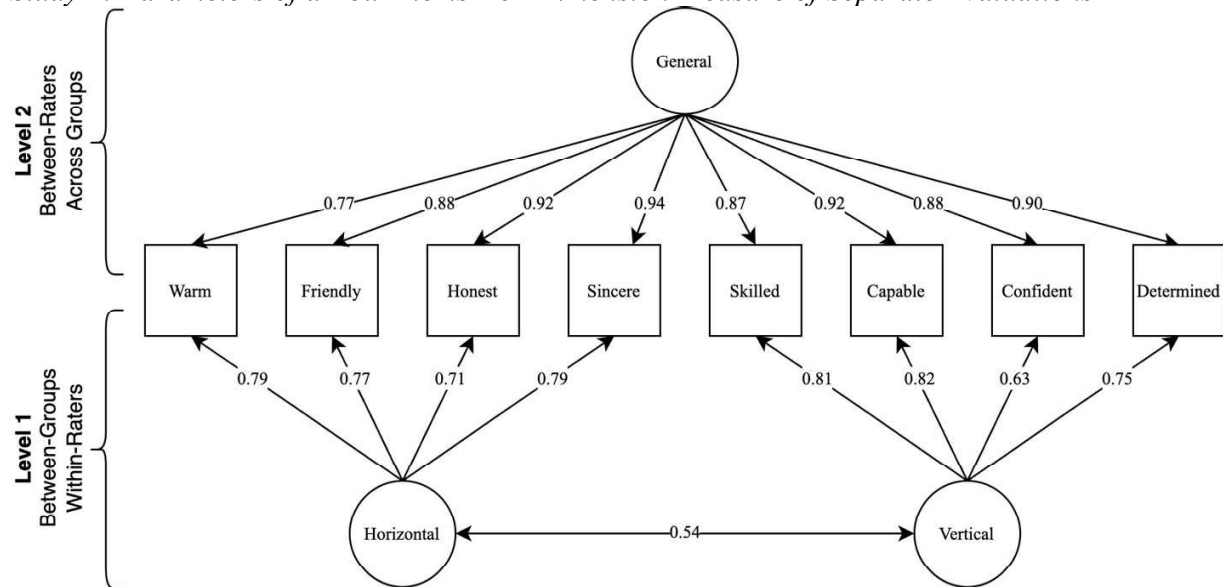


Figure S2.7

Study 2: Parameters of a Four-Items-Per-Dimension Measure of Separate Evaluations

**Table S2.1**

Study 2: Satisfactory and superior fit of a two-items-per-facet model

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; df = 34						
Separate evaluations	216	0.99	0.04	0.08	65,349	
Joint evaluations	204	0.99	0.04	0.09	65,000	
Personal evaluations	248	0.98	0.05	0.09	65,150	
Cultural evaluations	164	0.99	0.04	0.07	65,306	
Two-items-per-facet & Two-facets-per-dimension; df = 35						
Separate evaluations	222	0.98	0.04	0.08	65,352	0.017
Joint evaluations	228	0.98	0.04	0.09	65,022	<.001
Personal evaluations	263	0.98	0.05	0.09	65,162	<.001
Cultural evaluations	174	0.99	0.04	0.07	65,315	0.001
Four-items-per-dimension; df = 39						
Separate evaluations	1,005	0.92	0.09	0.12	66,127	<.001
Joint evaluations	779	0.93	0.08	0.12	65,565	<.001
Personal evaluations	967	0.92	0.09	0.11	65,858	<.001
Cultural evaluations	834	0.93	0.08	0.11	65,967	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

Table S2.2
Study 2: Tabular confirmatory factor analyses, averaging across the 4 target individuals

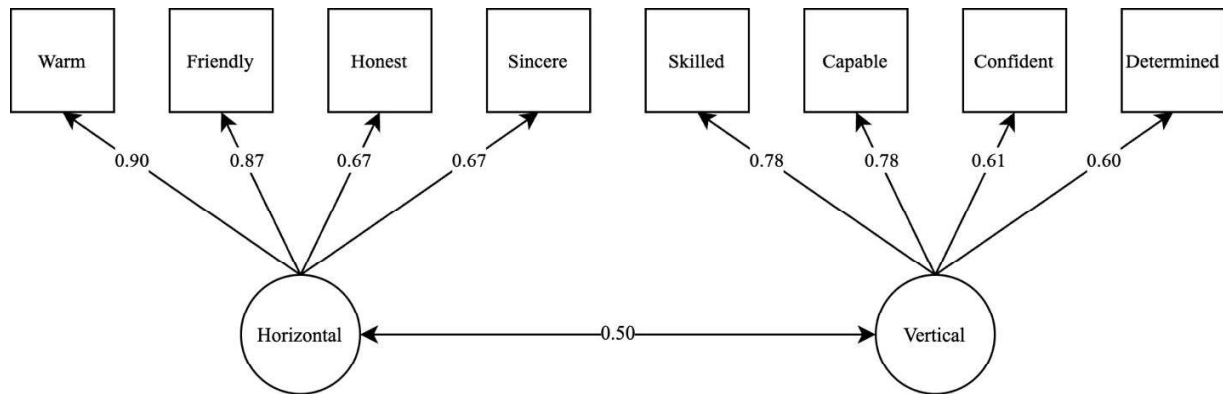
	X ²	CFI	RMSEA	SRMR	AIC
Two-items-per-facet; <i>df</i> = 34					
Separate evaluations	60.96	0.98	0.07	0.03	16,054
Joint evaluations	21.15	1.00	0.03	0.02	16,154
Personal evaluations	42.25	0.99	0.05	0.02	15,999
Cultural evaluations	40.97	0.99	0.05	0.02	16,161
Four-items-per-dimension; <i>df</i> = 39					
Separate evaluations	244.44	0.93	0.13	0.05	16,227
Joint evaluations	145.09	0.95	0.10	0.04	16,268
Personal evaluations	187.87	0.94	0.11	0.05	16,134
Cultural evaluations	199.21	0.94	0.11	0.05	16,309

Note. These were factor analyses that we ran separately for each individual, with the perceivers as rows, the eight or twenty items as columns, and the perceivers' ratings of the target in a given factor analysis on a given item as the data points. The table reports fit indices averaged across all 4 individuals that we examined in Study 2.

Study 3

Figure S5.1

Study 3: Parameters of a Four-Items-Per-Dimension Measure of Cultural Evaluations

**Table S3.1**

Study 3: Fit Measures with Simple-Structured Model with Two Correlated Dimensions and Four Items per Dimension (i.e., Combined Approach)

	X ²	CFI	RMSEA	SRMR	AIC	p(X ²)
Two-items-per-facet; <i>df</i> = 14	230	0.94	0.12	0.06	-12,512	
Two-items-per-facet & Two-facets-per-dimension; <i>df</i> = 15	243	0.94	0.12	0.07	-12,505	<.001
Four-items-per-dimension; <i>df</i> = 19	394	0.90	0.14	0.08	-12,362	<.001

Note. p(X²) represents the p-value from a nested chi-square difference test compared to the respective two-items-per-facet model.

Study 4
Table S4.1
Study 4: Results of Multilevel Linear Regression and Dominance Analysis

	Loyalty Game				Deception Game				Investment Game				Hubris Game			
	b	t	p	DA	b	t	p	DA	b	t	p	DA	b	t	p	DA
Morality	0.574*** [0.524, 0.625]	20.369	<.001	0.447 (1)	0.117*** [0.078, 0.157]	5.448	<.001	0.134 (3)	0.470*** [0.401, 0.538]	12.120	<.001	0.205 (3)	0.173*** [0.134, 0.212]	8.635	<.001	0.248 (2)
Friendliness	0.462*** [0.412, 0.512]	16.384	<.001	0.057 (2)	0.225*** [0.186, 0.264]	10.428	<.001	0.490 (1)	0.301*** [0.232, 0.369]	7.754	<.001	0.082 (4)	0.115*** [0.076, 0.154]	5.736	<.001	0.109 (4)
Ability	0.343*** [0.293, 0.394]	12.176	<.001	0.032 (3)	0.159*** [0.120, 0.199]	7.382	<.001	0.246 (2)	0.667*** [0.598, 0.735]	17.202	<.001	0.415 (1)	0.154*** [0.115, 0.193]	7.690	<.001	0.196 (3)
Assertiveness	0.279*** [0.229, 0.329]	9.893	<.001	0.104 (4)	0.116*** [0.077, 0.156]	5.390	<.001	0.131 (4)	0.566*** [0.497, 0.634]	14.596	<.001	0.298 (2)	0.232*** [0.193, 0.271]	11.598	<.001	0.447 (1)
Constant	0.368*** [0.341, 0.396]	25.355	<.001		0.378*** [0.350, 0.406]	25.970	<.001		0.390*** [0.362, 0.418]	25.842	<.001		0.524*** [0.493, 0.554]	33.984	<.001	
Observations		396				396				396				396		
Log Likelihood		-2,799.46				-3,350.48				-2,636.08				3,606.19		
AIC		5,614.93				7,116.96				5,288.16				7,228.37		
BIC		5,668.96				7,170.99				5,342.19				7,2282.40		

Note. Table S4.1 shows the results of a linear mixed model with random intercepts for participant and evaluated target for each of the four games in Study 4. 95% confidence intervals shown beneath slope estimate, b. DA = dominance analysis standardized general dominance (rank in parentheses); AIC = Akaike Information Criteria; BIC = Bayesian Information Criteria. *** means $p < .001$; ** means $p < .010$; * means $p < .050$; + means $p < .100$.

Table S4.2*Study 4: Pairwise Comparisons of the Magnitude of Regression Coefficients*

Game	Hypothesized best predictor	β	Rivalling predictor	β	χ^2	p
Loyalty Game	Friendliness	0.22	Morality	0.12	12.4	< .001
Loyalty Game	Friendliness	0.22	Ability	0.16	4.64	.031
Loyalty Game	Friendliness	0.22	Assertiveness	0.12	12.7	< .001
Deception Game	Morality	0.57	Friendliness	0.46	7.94	.005
Deception Game	Morality	0.57	Ability	0.34	33.6	< .001
Deception Game	Morality	0.57	Assertiveness	0.28	54.9	< .001
Investment Game	Ability	0.67	Friendliness	0.30	12.9	< .001
Investment Game	Ability	0.67	Morality	0.47	44.6	< .001
Investment Game	Ability	0.67	Assertiveness	0.57	3.4	.065
Hubris Game	Assertiveness	0.23	Friendliness	0.11	4.39	.036
Hubris Game	Assertiveness	0.23	Morality	0.17	17.2	< .001
Hubris Game	Assertiveness	0.23	Ability	0.15	7.64	.006

Study 5

Table S5.1

Study 5: Results of Multilevel Linear Regression and Dominance Analysis

	Loyalty Game				Deception Game				Investment Game				Hubris Game			
	b	t	p	DA	b	t	p	DA	b	t	p	DA	b	t	p	DA
Morality	0.36*** [0.32, 0.41]	16.01	<.001	0.44 (1)	0.03 [-0.02, 0.08]	1.12	.263	0.18 (3)	0.26*** [0.21, 0.32]	9.38	<.001	0.20 (2)	-0.10*** [-0.15, -0.05]	-3.69	<.001	0.18 (2)
Friendliness	0.18*** [0.15, 0.22]	11.45	<.001	0.34 (2)	0.15*** [0.11, 0.18]	7.75	<.001	0.51 (1)	0.10*** [0.06, 0.13]	4.80	<.001	0.13 (4)	0.04* [0.00, 0.07]	1.99	.047	0.06 (4)
Ability	0.16*** [0.11, 0.20]	6.92	<.001	0.17 (3)	-0.06* [-0.11, -0.01]	-2.29	.022	0.06 (4)	0.76*** [0.71, 0.82]	27.31	<.001	0.51 (1)	0.020 [-0.03, 0.07]	0.73	.467	0.14 (3)
Assertiveness	0.04† [-0.01, 0.08]	1.72	.085	0.06 (4)	0.15*** [0.10, 0.20]	6.00	<.001	0.25 (2)	0.10*** [0.05, 0.15]	3.80	<.001	0.16 (3)	0.12*** [0.07, 0.17]	4.93	<.001	0.63 (1)
Constant	0.25*** [0.21, 0.28]	14.17	<.001		0.50*** [0.46, 0.53]	25.97	<.001		0.37*** [0.34, 0.40]	23.22	<.001		0.38*** [0.36, 0.42]	20.72	<.001	
Observations		304				298				304				303		
Log Likelihood		-9,445.11				-13,904.71				-15,657.38				-13,361.88		
AIC		18,904.22				27,823.43				31,328.76				26,737.76		
BIC		18,962.47				27,881.54				31,387.02				26,795.99		

Note. Table S5.1 shows the results of a linear mixed model with random intercepts for participant and evaluated target for each of the four games in Study 5. 95% confidence intervals shown beneath slope estimate, b. DA = dominance analysis standardized general dominance (rank in parentheses); AIC = Akaike Information Criteria; BIC = Bayesian Information Criteria. *** means $p < .001$; ** means $p < .010$; * means $p < .050$; + means $p < .100$.

Table S5.2*Studies 5a-5d: Pairwise Comparisons of the Magnitude of Regression Coefficients*

Game	Hypothesized best predictor	β	Rivalling predictor	β	χ^2	p
Loyalty Game	Friendliness	0.15	Morality	0.03	7.84	.005
Loyalty Game	Friendliness	0.15	Ability	-0.06	37.3	< .001
Loyalty Game	Friendliness	0.15	Assertiveness	0.15	0.01	.930
Deception Game	Morality	0.37	Friendliness	0.18	27.4	< .001
Deception Game	Morality	0.37	Ability	0.16	29.8	< .001
Deception Game	Morality	0.37	Assertiveness	0.03	125	< .001
Investment Game	Ability	0.78	Friendliness	0.10	127	< .001
Investment Game	Ability	0.78	Morality	0.26	380	< .001
Investment Game	Ability	0.78	Assertiveness	0.09	201	< .001
Hubris Game	Assertiveness	0.12	Friendliness	0.04	41.5	< .001
Hubris Game	Assertiveness	0.12	Morality	-0.10	5.85	.016
Hubris Game	Assertiveness	0.12	Ability	0.02	4.25	.039