



Investigating object orientation effects across 18 languages

Sau-Chin Chen¹ · Erin M. Buchanan² · Zoltan Kekecs^{3,4} · Jeremy K. Miller⁵ · Anna Szabelska⁶ · Balazs Aczel³ · Pablo Bernabeu^{7,8} · Patrick Forscher^{9,10} · Attila Szuts³ · Zahir Vally¹¹ · Ali H. Al-Hoorie¹² · Mai Helmy^{13,14} · Caio Santos Alves da Silva¹⁵ · Luana Oliveira da Silva¹⁵ · Yago Luksevicius de Moraes^{15,16} · Rafael Ming Chi Santos Hsu¹⁵ · Anthonieta Looman Mafra¹⁵ · Jaroslava V. Valentova¹⁵ · Marco Antonio Correa Varella¹⁵ · Barnaby Dixon¹⁷ · Kim Peters¹⁷ · Nik Steffens¹⁷ · Omid Ghasemi¹⁸ · Andrew Roberts¹⁸ · Robert M. Ross¹⁹ · Ian D. Stephen^{18,20} · Marina Milyavskaya²¹ · Kelly Wang²¹ · Kaitlyn M. Werner²² · Dawn Liu Holford²³ · Miroslav Sirota²³ · Thomas Rhys Evans²⁴ · Dermot Lynott²⁵ · Bethany M. Lane²⁶ · Danny Riis Sahlholdt²⁶ · Glenn P. Williams²⁷ · Chrystalle B. Y. Tan²⁸ · Alicia Foo²⁹ · Steve M. J. Janssen²⁹ · Nwadiogo Chisom Arinze³⁰ · Izuchukwu Lawrence Gabriel Ndukaihe³⁰ · David Moreau³¹ · Brianna Jurosic³² · Brynna Leach³² · Savannah Lewis³³ · Peter R. Mallik³⁴ · Kathleen Schmidt³² · William J. Chopik³⁵ · Leigh Ann Vaughn³⁶ · Manyu Li³⁷ · Carmel A. Levitan³⁸ · Daniel Storage³⁹ · Carlota Batres⁴⁰ · Tyler McGee⁴⁰ · Janina Enachescu⁴¹ · Jerome Olsen⁴² · Martin Voracek⁴³ · Claus Lamm⁴³ · Ekaterina Pronizius⁴³ · Tilli Ripp⁴⁴ · Jan Philipp Röer⁴⁴ · Roxane Schnepfer⁴⁴ · Marietta Papadatou-Pastou⁴⁵ · Aviv Mokady⁴⁶ · Niv Reggev⁴⁶ · Priyanka Chandel⁴⁷ · Pratibha Kujur⁴⁷ · Babita Pande⁴⁷ · Arti Parganiha⁴⁷ · Noorshama Parveen⁴⁷ · Sraddha Pradhan⁴⁷ · Margaret Messiah Singh⁴⁷ · Max Korbmacher⁴⁸ · Jonas R. Kunst⁴⁹ · Christian K. Tamnes⁵⁰ · Frederike S. Woelfert⁵⁰ · Kristoffer Klevjer⁵¹ · Sarah E. Martiny⁵¹ · Gerit Pfuhl^{51,52} · Sylwia Adamus⁵³ · Krystian Barzykowski⁵³ · Katarzyna Filip⁵³ · Patricia Arriaga⁵⁴ · Vasilije Gvozdenovic⁵⁵ · Vanja Kovic⁵⁵ · Fei Gao⁵⁶ · Jingxiang Li⁵⁶ · Jozef Bavoľár⁵⁷ · Monika Hricová⁵⁷ · Pavol Kačmár⁵⁷ · Matúš Adamkovič^{58,59,60} · Peter Babinčák⁶¹ · Gabriel Baník⁶² · Ivan Ropovik^{63,64} · Danilo Zambrano Ricaurte⁶⁵ · Sara Álvarez-Solas⁶⁶ · Harry Manley^{67,68} · Panita Suavansri⁶⁷ · Chun-Chia Kung⁶⁹ · Belemir Çoktok⁷⁰ · Asil Ali Özdoğru^{70,71} · Çağlar Solak⁷² · Sinem Söylemez⁷² · Sami Çoksarı^{73,74} · İlker Dalgat⁷⁵ · Mahmoud Elsherif⁷⁶ · Martin Vasilev⁷⁷ · Vinka Mlakic⁷⁸ · Elisabeth Oberzaucher⁷⁹ · Stefan Stieger⁷⁸ · Selina Volsa⁷⁸ · Erica D. Musser⁸⁰ · Janis Zickfeld⁸¹ · Christopher R. Chartier³²

Accepted: 20 July 2025 / Published online: 22 October 2025

© The Author(s), under exclusive licence to Springer Science+Business Media, LLC, part of Springer Nature 2025

Abstract

Mental simulation theories of language comprehension propose that people automatically create mental representations of objects mentioned in sentences. Mental representation is often measured with the sentence-picture verification task, wherein participants first read a sentence that implies the object property (i.e., shape and orientation). Participants then respond to an image of an object by indicating whether it was an object from the sentence or not. Previous studies have shown matching advantages for shape, but findings concerning object orientation have not been robust across languages. This registered report investigated the match advantage of object orientation across 18 languages in nearly 4,000 participants. The preregistered analysis revealed no compelling evidence for a match advantage for orientation across languages. Additionally, the match advantage was not predicted by mental rotation scores. In light of these findings, we discuss the implications for current theory and methodology surrounding mental simulation.

Keywords Cross-lingual research · Language comprehension · Mental rotation · Mental simulation

Extended author information available on the last page of the article

Mental simulation of object properties is a major topic in conceptual processing research (Ostarek & Huettig, 2019; Scorolli, 2014). Theoretical frameworks of conceptual processing describe the integration of linguistic representations and situated simulation (e.g., reading about bicycles integrates the situation in which bicycles would be used, Barsalou, 2008; Zwaan, 2014a). Proponents of situated cognition contend that perceptual representations can be generated during language processing (Barsalou, 1999; Wilson, 2002), as cognition is thought to be an interaction of the body, environment, and processing (Barsalou, 2020). Given this definition of situated cognition, it is important to investigate previously established embodied cognition effects across multiple environments (in this case, languages and cultures), especially as the credibility revolution has indicated that not all published findings are replicable (Vazire, 2018).

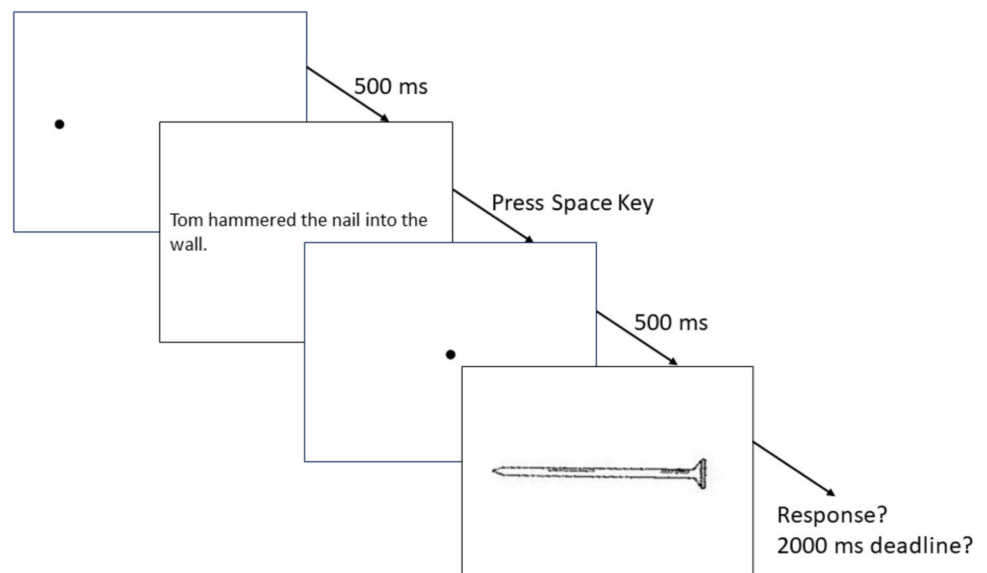
One empirical index of situated simulation is the mental simulation effect measured in the sentence-picture verification task (see Figure 1). This task requires participants to read a probe sentence displayed on the screen. On the following screen, participants see a picture of an object and must verify whether the object was mentioned in the probe sentence. Verification response times are used to test the mental simulation effect, which occurs when people are faster to respond to pictures that match the properties implied by the probe sentences. For example, the orientation implied by the sentence “*Tom hammered the nail into the wall*” would be matched if the following picture showed a horizontally-oriented nail rather than a vertically-oriented one. The opposite would be true of the sentence “*Tom hammered the nail into the floor plank*”.

Mental simulation effects have been demonstrated for object shape (Zwaan et al., 2002), color (Connell, 2007), and orientation (Stanfield & Zwaan, 2001). Subsequent

replication studies revealed consistent results for shape but inconsistent findings for color and orientation effects (De Koning et al., 2017a; Rommers et al., 2013; Zwaan & Pecher, 2012). Existing theoretical frameworks do not provide much guidance regarding the potential causes for this discrepancy. With the accumulating concerns about the lack of reproducibility (e.g., Kaschak & Madden, 2021), researchers have found it challenging to reconcile the theory of mental simulation with the failures to replicate some of the effects (e.g., Morey et al., 2022). In an empirical discipline like cognitive science, a theory requires the support of reproducible results.

The reliability of match advantage effects seems to vary depending on both the object properties and the languages under study. Mental simulation effects for object shape have consistently been found in English (Zwaan et al., 2017; Zwaan & Madden, 2005; Zwaan & Pecher, 2012), Chinese (Li & Shang, 2017), Dutch (De Koning et al., 2017a; Engelen et al., 2011; Pecher et al., 2009; Rommers et al., 2013), German (Koster et al., 2018), Croatian (Šetić & Domijan, 2017), and Japanese (Sato et al., 2013). Object orientation, on the other hand, has produced mixed results across languages: namely, positive evidence in English (Stanfield & Zwaan, 2001; Zwaan & Pecher, 2012) and in Chinese (Chen et al., 2020), and null evidence in Dutch (De Koning et al., 2017a; Rommers et al., 2013) and in German as second language (Koster et al., 2018). Among studies on shape and orientation, the effects of object orientation have been smaller than those of object shape (e.g., $d = 0.10$ vs. 0.17 in Zwaan & Pecher, 2012; $d = 0.07$ vs. 0.27 in De Koning et al., 2017b). To understand the causes for the discrepancies among object properties and languages, it is imperative to consider the cross-linguistic and experimental factors of the sentence-picture verification task.

Fig. 1 Procedure of the sentence-picture verification task, with an example of matching orientation



Cross-linguistic, methodological, and cognitive factors

Several factors might contribute to cross-linguistic differences in the match advantage of object orientation. First, languages differ in how they encode motion and placement events in sentences (Newman, 2002; Verkerk, 2014). Second, the potential role of mental rotation as a confound has been considered (Rommers et al., 2013). We expand on how linguistic, methodological, and cognitive factors hinder the improvement of theoretical frameworks below.

Linguistic factors The probe sentences used in object orientation studies usually contain several motion events (e.g., “*The ant walked towards the pot of honey and tried to climb in*”). Languages encode motion events in different ways, and grammatical differences between lexical encodings could explain different match advantage results. According to Verkerk (2014), Germanic languages (e.g., Dutch, English, and German) generally encode the manner of motion in the verb (e.g., “*The ant dashed*”), while conveying the path information through satellite adjuncts (e.g., “*towards the pot of honey*”). In contrast, other languages, such as the Romance family (e.g., Portuguese, Spanish), more often encode the path in the verb (e.g., “*crossing*”, “*exiting*”). Crucially, past research on the match advantage of object orientation is exclusively based on Germanic languages, and yet, there were differences across those languages, with English being the only one that consistently yielded the match advantage. As a minor difference across Germanic languages in this regard, Verkerk (2014) notes that path-only constructions (e.g., “*The ant went to the feast*”) are more common in English than in other Germanic languages.

Another topic to be considered is the lexical encoding of placement in each language, as the stimuli contain several placement events (e.g., “*Sara situated the expensive plate on its holder on the shelf*”). Chen et al. (2020) and Koster et al. (2018) noted that some Germanic languages, such as German and Dutch, often make the orientation of objects more explicit than English. In English, for example, the verb “*put*” does not convey a specific orientation in the sentences “*She put the book on the table*” and “*She put the bottle on the table*.” However, in German and Dutch, speakers preferred the verbs “*laid*” or “*stood*” in the above sentences. In this case, the verb “*lay*” encodes a horizontal orientation, whereas the verb “*stand*” encodes a vertical orientation. This distinction extends to verbs indicating existence. As Newman (2002) exemplified, an English speaker would be likely to say “*There’s a lamp in the corner*,” whereas a Dutch speaker would be more likely to say “*There ‘stands’ a lamp in the corner*.” Nonetheless, we cannot conclude that these cross-linguistic differences

are affecting the match advantage across languages because the current theories (e.g., Language and Situated Simulation, Barsalou, 2008) have not addressed the potential influence of linguistic aspects such as the lexical encoding of placement.

Methodological factors Inconsistent findings on the match advantage of object orientation may be due to variability in task design. For example, studies failing to detect the match advantage may not have required participants to verify the probe sentence after the response to the target picture (see Zwaan, 2014a). Without such a verification, participants might have paid less attention to the meaning of the probe sentences, in which they would have been less likely to form a mental representation of the objects (e.g., Zwaan & van Oostendorp, 1993). In this regard, variability originating from differences in the characteristics of experiments can substantially influence the results (Barsalou, 2019; Kaschak & Madden, 2021).

Cognitive factors Since Stanfield and Zwaan (2001) showed a match advantage of object orientation, later studies on this topic have examined the association between the match advantage and visual imagery rather than situated simulation. Visual imagery can be measured with indirect measurements such as vividness questionnaires (Zwaan & Pecher, 2012) and via more direct measurements such as mental rotation tasks (Chen et al., 2020). Indeed, studies have suggested that mental rotation tasks offer valid reflections of previous spatial experience (Frick & Möhring, 2013) and of current spatial cognition (Chu & Kita, 2008; Pouw et al., 2014). Some previous studies have drawn on mental rotation to study mental simulation. For instance, De Koning et al. (2017a) observed that the effectiveness of mental rotation increased with the size of the depicted object. Chen et al. (2020) examined the implication of this finding for the match advantage of object orientation (Stanfield & Zwaan, 2001) and implemented a picture-picture verification task using the mental rotation paradigm (Cohen & Kubovy, 1993). In each trial, two pictures appeared on opposite sides of the screen. Participants had to verify whether the pictures represented identical or different objects in which the two objects presented in congruent orientations (both are horizontal or vertical) or in incongruent orientations (one is horizontal and one is vertical).

Chen et al. (2020) not only revealed shorter verification times for congruent orientations (i.e., two identical pictures presented in horizontal or vertical orientation) but also replicated the advantage for larger objects (i.e., “*pictures of bridges versus pictures of pens*”). The results were consistent across the three languages investigated—English, Dutch and Chinese. Chen et al. (2020) converted the picture-picture verification times to the mental rotation scores that were the discrepancy of verification times between the identical and different orientations.

Their analysis showed that mental rotation affected the Dutch participants' sentence-picture verification performance. With the measurement of mental rotation scores¹, we explore the association of spatial cognition and the effect of orientation in comprehension across the investigated languages.

Purposes of this study

To scrutinize the discrepancies in findings across languages and cognitive factors, we examined the reproducibility of the object orientation effect in a multi-lab collaboration. Our pre-registered plan aimed at detecting a general match advantage of object orientation across languages and evaluated the magnitude of match advantage in each specific language. Additionally, we examined whether the match advantages were related to the mental rotation index. Thus, this study followed the original methods from Stanfield and Zwaan (2001) and addressed two primary questions: (1) How much of the match advantage of object orientation can be obtained within different languages, and (2) How are differences in the mental rotation index associated with the match advantage across languages?

Method

Hypotheses and design

The study design for the sentence-picture and picture-picture verification tasks was mixed, using between-participant (language) and within-participant (match versus mismatch object orientation) independent variables. In the sentence-picture verification task, the match condition reflects a match between the sentence and the picture, whereas in the picture-picture verification it reflects a match in orientation between two pictures. The only dependent variable for both tasks was response time. The time difference between match conditions in each task is the measurement of mental simulation effects (for the sentence-picture task) and mental rotation scores (for the picture-picture task). We did not select languages systematically but instead based on our collaboration recruitment with the Psychological Science Accelerator (PSA, Moshontz et al., 2018).

We pre-registered the following hypotheses:

- (1) In the sentence-picture verification task, we expected response times to be shorter for matching compared to mismatching orientations within each language. In the picture-picture verification task, we expected shorter

response time for identical orientation compared to different orientations. We did not have any specific hypotheses about the relative size of the object orientation match advantage in different languages.

- (2) Based on the assumption that “*mental rotation is a general cognitive function*”, we expect equal mental rotation scores across languages, but no association between mental rotation scores and mental simulation effects (see Chen et al., 2020).

Participants

We performed a pre-registered power analysis, which sought to achieve a power of 80% in a directional one-sample *t*-test. When an effect size of $d = 0.20$ was hypothesized, a sample size of $N = 156$ was required. Instead, for a hypothetical effect size $d = 0.10$, a sample size of $N = 620$ was required. In addition, a power analysis tailored to mixed-effects models was performed. The effect size hypothesized in this analysis was equal to that observed by Zwaan and Pecher (2012), and the number of items was 100 (i.e., 24 planned items nested within at least five languages). The result revealed that a sample size of $N = 400$ would be required to achieve a power of 90%. We expected laboratories to show differences in orientation effects, and therefore, the mixed effect analysis treated the laboratories as a random variable to account for different analyses. The laboratories were allowed to follow a secondary plan: a team collected at least their pre-registered minimum sample size (suggested 100 to 160 participants, most implemented 50), and then determine whether or not to continue data collection via Bayesian sequential analysis (stopping data collection if $BF_{10} = 10$ or $\frac{1}{10}$)².

We collected data in 18 languages from 50 laboratories. Each laboratory chose a maximal sample size and an incremental n for sequential analysis before their data collection. Because the pre-registered power analysis did not align with the final analysis plan, we conducted a sensitivity analysis (See Appendix 1) to determine the smallest effect size our study could reliably detect given our combination of degrees of freedom and standard error found in the study. The results indicated that our design was sufficiently sensitive to detect an orientation effect of 2.36 ms or larger as statistically significant. For comparison, the average difference was 44 ms in the Stanfield and Zwaan (2001), 35 ms in Zwaan and Pecher (2012), 1 ms in Rommers et al. (2013), and 7 ms in De Koning et al. (2017b).

The original sample sizes are presented in Table 1 for the teams that provided raw data to the project. Data collection

¹ In the pre-registered plan, we used the term “imagery score” but this term was confusing. Therefore, we used “mental rotation scores” instead of “imagery scores” in the final report.

² See details of power analysis in the pre-registered plan, pp. 13 - 15. https://osf.io/preprints/psyarxiv/t2ppv_v1.

Table 1 Demographic and sample size characteristics

Language	SP_{Trials}	PP_{Trials}	SP_N	PP_N	$Demo_N$	$Female_N$	$Male_N$	M_{Age}	SD_{Age}
Arabic	2544	2544	106	106	107	42	12	32.26	18.59
Brazilian Portuguese	1200	1200	50	50	50	36	13	30.80	8.73
English	45189	45312	1884	1888	2055	1360	465	21.71	3.85
German	5616	5616	234	234	248	98	26	22.34	3.40
Greek	2376	2376	99	99	109	0	0	33.86	11.30
Hebrew	3576	3571	149	149	181	0	0	24.25	9.29
Hindi	1896	1896	79	79	86	57	27	21.66	3.46
Magyar	3610	3816	151	159	168	3	1	21.50	2.82
Norwegian	3576	3576	149	149	154	13	9	25.22	6.40
Polish	1368	1368	57	57	146	0	0	23.25	7.96
Portuguese	1488	1464	62	61	55	26	23	30.74	9.09
Serbian	3120	3120	130	130	130	108	21	21.38	4.50
Simplified Chinese	2040	2016	85	84	96	0	1	21.92	4.68
Slovak	3881	3599	162	150	325	1	0	21.77	2.33
Spanish	3120	3096	130	129	146	0	0	21.73	3.83
Thai	1200	1152	50	48	50	29	9	21.54	3.81
Traditional Chinese	3600	3600	150	150	186	69	46	20.89	2.44
Turkish	6456	6432	269	268	274	36	14	21.38	4.59

SP Sentence Picture Verification, PP Picture Picture Verification. Sample sizes for demographics may be higher than the sample size for this study, as participants could have only completed the bundled experiment. Additionally, not all entries could be unambiguously matched by lab ID, and therefore, demographic sample sizes could also be less than data collected

proceeded in two broad stages: initially we collected data in the laboratory. However, when the global COVID-19 pandemic made this practice impossible to continue, we moved data collection online. In total, 4,248 unique participants completed the present study with 2,843 completing the in-person version and 1,405 completing the online version³. The in-person version included 35 research teams and the online version included 19 with 50 total teams across both data collection methods (i.e., some labs completed both in-person and online data collection). Based on recommendations from the research teams (Turkish lab: TUR_007, Taiwanese lab: TWN_002), two sets of data were excluded from all analyses due to participants being non-native speakers. Figure 2 provides a flow chart for participant exclusion and inclusion for analyses. All participating laboratories had either ethical approval or institutional evaluation before data collection. All data and analysis scripts are available on the source files (<https://codeocean.com/capsule/9287673>). Appendix 2 summarizes the average characteristics by language and laboratory.

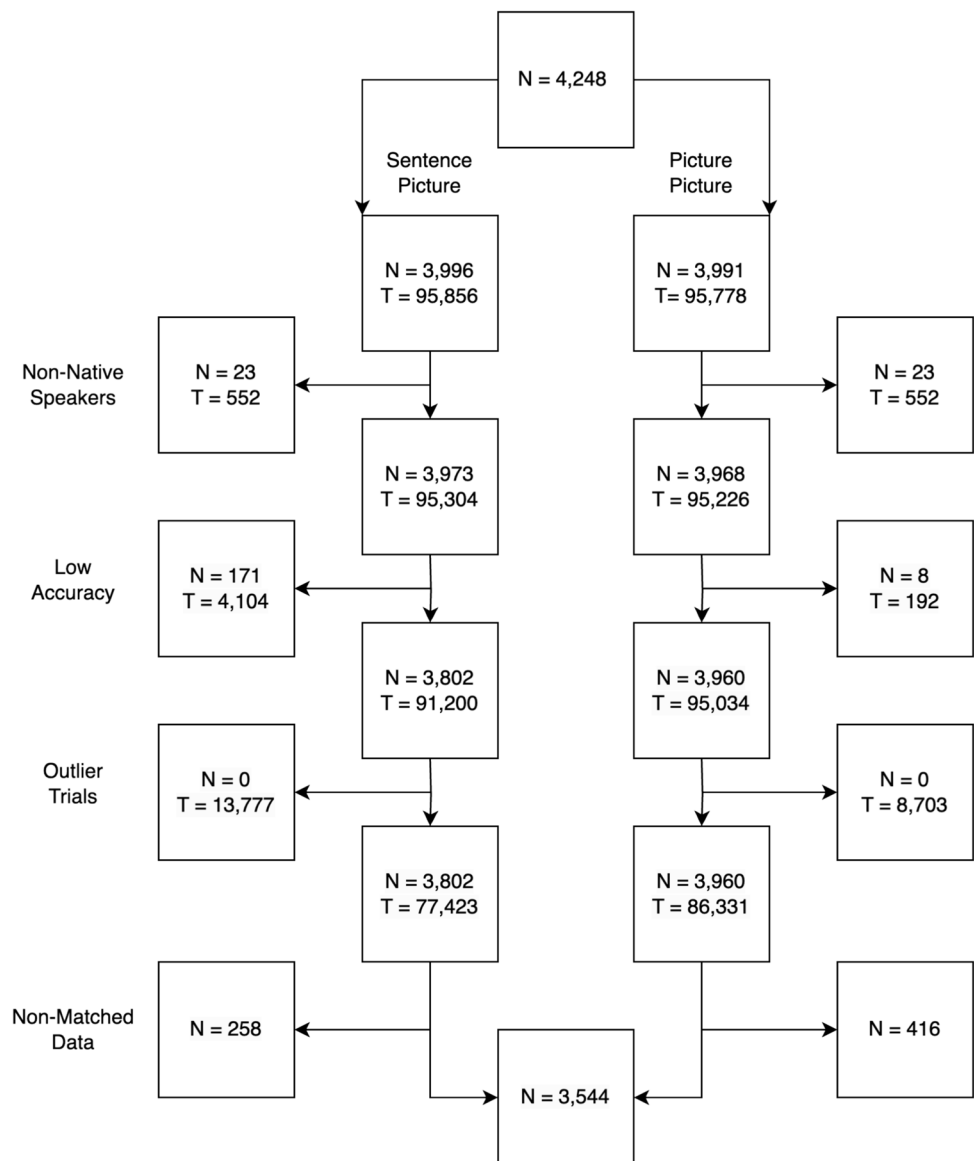
³ Data for this study was collected together with another unrelated study during the same data collection session, with the two studies using different data collection platforms. The demographic data was collected within the platform of the other study during the in-person sessions. Some participants only completed the Phillips et al. study and dropped out without completing the present study, and there were also some data entry errors in the demographic data. Thus, the demographic data of some participants who took the present study are missing or unidentifiable ($n = 39$ cannot be matched to a lab, $n = 2,053$ were missing gender information, and $n = 332$ were missing age information). Importantly, this does not affect the integrity of the experimental research data.

Materials

Sentences 24 critical sentence pairs (48 total sentences) were included in this study following Stanfield and Zwaan (2001). Each pair consisted of versions that differed in their implied orientation of the object embedded in the sentence. For instance, the sentence “*The librarian put the book back on the table*” — which implies a horizontal orientation — had a counterpart in the sentence “*The librarian put the book back on the shelf*” — which implies a vertical orientation. Another two sets of 24 sentences were included as filler sentences for the task demand. These sentences were not matched to any particular orientation but included a potential object for depiction. For example, “*After a week the painting arrived by mail*”, and “*The flowers that were planted last week had survived the storm*” were included as filler sentences. Each participant was shown 24 critical sentences and 24 filler sentences in the study. The filler sentences were included to counterbalance the number of yes-no answers to create an even 50% ratio.

Pictures The study included 24 critical matched pictures that only varied in their orientation (vertical/horizontal) for a total of 48 critical pictures (from Zwaan & Pecher, 2012). These pictures were matched to their respective sentences for implied orientation. “*The librarian put the book back on the table*” was matched with a horizontally oriented book, while “*The librarian put the book back on the shelf*” was matched with a vertically oriented book. For

Fig. 2 Sample size and exclusions. N = number of unique participants, T = number of trials. The final combined sample was summarized to a median score for each match/mismatch condition, and therefore, includes one summary score per person



counterbalancing, the mismatch between picture orientation and sentence was created, and the book would be shown in the respective opposite orientation (see orientation pairs at <https://osf.io/utqxb>). Another 48 pictures from Zwaan and Pecher (2012) were included for the fillers which were unrelated to the corresponding sentence. Therefore, the answer to critical pairs was always “yes”, while the filler sentence-picture combinations answer was always “no”.

Picture-picture trials The picture-picture verification task used the same object pictures as the above task. The 24 critical picture pairs were included as match trials and were counterbalanced such as half the time they appeared with the same object and orientation (i.e., the same picture), and half the time with the opposite orientation (i.e., horizontal and vertical). The filler pictures were randomly

paired to create mismatch trials. Table 2 shows the counterbalancing and combinations for trials.

Procedure

Sentence-picture task The sentence-picture verification task was always administered first. This task began with six practice trials. Each trial started with a left-aligned vertically centered fixation point displayed for 1,000 ms, immediately followed by the probe sentence. The sentence remained on the screen until the participant pressed the space key, acknowledging that they had read the sentence. Then, the object picture was presented in the center of the screen until the participant responded, or it disappeared after two seconds. Participants were instructed to verify, as quickly and accurately as possible,

Table 2 Trial conditions for the sentence-picture and picture-picture verification task

Condition	Item 1	Item 2	Answer	Number
Sentence-Picture Critical Match	Critical Sentence: Horizontal	Critical Picture: Horizontal	Yes	6
Sentence-Picture Critical Match	Critical Sentence: Vertical	Critical Picture: Vertical	Yes	6
Sentence-Picture Critical Mismatch	Critical Sentence: Horizontal	Critical Picture: Vertical	Yes	6
Sentence-Picture Critical Mismatch	Critical Sentence: Vertical	Critical Picture: Horizontal	Yes	6
Sentence-Picture Filler	Sentence	Picture	No	24
Picture-Picture Critical Match	Critical Picture: Horizontal	Critical Picture: Horizontal	Yes	6
Picture-Picture Critical Match	Critical Picture: Vertical	Critical Picture: Vertical	Yes	6
Picture-Picture Critical Mismatch	Critical Picture: Horizontal	Critical Picture: Vertical	Yes	6
Picture-Picture Critical Mismatch	Critical Picture: Vertical	Critical Picture: Horizontal	Yes	6
Picture-Picture Filler	Picture	Picture	No	24

whether the object on screen had been mentioned in the probe sentence. Following Stanfield and Zwaan (2001), a memory check test was carried out after every three to eight trials to ensure that participants had read each sentence carefully.

As shown in Table 2, the trials for the sentence-picture task were created by counterbalancing the sentence implied orientation (vertical, horizontal) by the pictured object orientation creating a fully crossed combination between matching sentences and objects. Therefore, each participant only saw one of the four possible combinations (sentence orientation 2 X object orientation 2). For the filler items, sentences and unrelated target pictures were randomly assigned in two separate lists. These fillers were included with the critical pairs but were excluded from the analysis. Stimuli lists were created in Excel, and this information can be found at <https://osf.io/utqxb>.

Translation of sentences The translation of probe sentences followed our pre-registered plan. Every non-English language coordinator was required to recruit at least four translators who were fluent in both English and the target language. Every language coordinator supervised the translators using the Psychological Science Accelerator

guidelines. In addition, the coordinator and participating laboratories consulted about each of the following points:

- (1) Four translators could denote the items that are unfamiliar to a particular language based on object familiarity ratings. The two forward translators would suggest alternative probe sentences and object pictures to replace the unfamiliar objects. The two backward translators would evaluate the suggested items.
- (2) Some objects in a particular language have different spellings or pronunciations among countries and geographical zones due to dialect. For example, American speakers tend to write “*tire*” whereas British speakers tend to write “*tyre*”. Every coordinator would mark these local translations in the final version of translated materials. Participating laboratories could replace the names to match the local dialect.

Picture-picture task Next, the picture-picture verification task was administered. In each trial, two objects appeared on either side of the central fixation point until either the participant indicated that the pictures displayed the same object or two different objects, or until two seconds elapsed. As shown in Table 2, four possible combinations of critical orientations could be shown with the picture (same, different) by orientation (same, different). Each participant only saw one of the critical combinations, and filler items were randomly paired in two combinations to match. The stimuli lists can be found at <https://osf.io/utqxb>.

Software implementation The study was executed using OpenSesame software for millisecond timing (Mathôt & March, 2022). After data collection moved online, to minimize the differences between on-site and web-based studies, we converted the original Python code to Javascript and collected the data using OpenSesame through a JATOS server. We proceeded with the online study from February to June 2021 after the changes in the procedure were approved by the journal editor and reviewers. Following the literature, we did not anticipate any theoretically important differences between the two data collection methods. The instructions and experimental scripts are available at the public OSF folder (<https://osf.io/e428p/> “Materials” in Files).

Analysis plan

To test Hypothesis 1, our first planned analysis⁴ used a random-effects meta-analysis model that estimated the match advantage across laboratories and languages. Analyses were

⁴ See the analysis plan in the pre-registered plan, p. 19 - 20. https://osf.io/preprints/psyarxiv/t2pju_v1. This plan was changed to a random-effects model to ensure that we did not assume the exact same effect size for each language and lab.

restricted to correct responses on critical trials. Filler trials were excluded because they were not designed to test the match advantage, and incorrect responses were excluded because accuracy is a prerequisite for interpreting potential match-advantage effects. Other data considerations including participant level accuracy and outliers are discussed in the results section. The meta-analysis summarized the median reaction times by match condition to determine the effect size by laboratory. The following formula was used:

$$d = \frac{Mdn_{Mismatch} - Mdn_{Match}}{\sqrt{MAD_{Mismatch}^2 + MAD_{Match}^2 - 2 \times r \times MAD_{Mismatch} \times MAD_{Match}}} \times \sqrt{2 \times (1 - r)}$$

where d is Cohen's d , Mdn is Median, MAD is median absolute deviation, and r is correlation between match and mismatch condition. Meta-analytic effect sizes were computed for those languages that had data from more than one team.

Continuing to test Hypothesis 1, next, we ran planned mixed-effects models using correct response times from the critical trials of sentence-picture verification task as the dependent variable. In each analysis, we first built a simple linear regression model with a fixed intercept only. Then, we systematically added random intercepts and fixed effects, arriving at the final model. First, the random intercepts were added to the model one-by-one in the following order: participant ID, target, laboratory ID, and finally language. See below section for decision criteria for determining the final random-effect structure. Then, the fixed effect of matching condition (match vs. mismatch) was added to the model. Language-specific mixed-effects models were conducted in the same way if the meta-analysis showed a significant orientation effect.

According to the pre-registration, we planned to test Hypothesis 2 by first evaluating the equality of mental rotation scores across languages using an ANOVA. However, this plan was updated to use mixed models instead due to the nested structure of the data. The same analysis plan was used for model building and selection as described above for the sentence-picture verification task. To further assess Hypothesis 2, the last planned analysis was to use mental rotation scores for the prediction of mental stimulation with an interaction between language and mental rotation scores computed from the picture-picture task to determine if there were differences in prediction of match advantage in the sentence-picture task. Here, we used a mixed-effects model as well to control for the random effect of the data collection lab, and with language, mental rotation score, and their interaction as fixed effect predictors.

Decision criterion for model selection and hypothesis testing The inclusion of both random and fixed effects in models was assessed using model comparison based on the

Akaike information criterion (AIC). While this method is less conservative than alternatives such as the likelihood ratio test, the AIC was deemed appropriate due to the modest effect sizes that tend to be produced by mental simulation effects, and the limited sample sizes in the present study (albeit larger samples than those of most previous studies). Models with lower AIC were preferred over models with higher AIC, and in cases where the difference in AIC did not reach 2, the model with fewer parameters was preferred.

p -values for each effect were calculated using the Satterthwaite approximation for degrees of freedom for individual predictor coefficients and meta-analysis. p -values were interpreted using the pre-registered α level of .05.

Intra-lab analysis during data collection Before data collection, each lab decided whether they wanted to apply a sequential analysis or whether they wanted to settle for a fixed sample size. The pre-registered protocol for labs applying sequential analysis established that they could stop data collection upon reaching the pre-registered criterion ($BF_{10} = 10$ or .10), or the maximal sample size. Each laboratory chose a fixed sample size and an incremental n for sequential analysis before their data collection. Two laboratories (Hungarian lab: HUN_001, Taiwanese lab: TWN_001) stopped data collection at the pre-registered criterion, $BF_{10} = .10$. Fourteen laboratories did not finish the sequential analysis because (1) twelve laboratories were interrupted by the pandemic outbreak; (2) two laboratories (Turkish lab: TUR_007E, Taiwanese lab: TWN_002E) recruited English-speaking participants to comply with institutional policies. Lab-based records were reported on a public website as each laboratory completed data collection (details are available in Appendix 3).

Results

Data screening

As shown in Figure 2, we recruited 4,248 participants. After removing the non-native speakers ($n_{total} = 23$), 4,225 were then examined for accuracy. Following our pre-registered plan, participants' data were excluded from analyses of the sentence-picture and picture-picture tasks if their accuracy was below 70%. This threshold is common in research of this kind and helps ensure that the data reflect meaningful cognitive processing rather than random or inattentive responses. Overall, $n_{total} = 8$ participants were removed for low accuracy across both tasks, leaving 4,217 participants. Figure 2 indicates the sample sizes for each task. Please note that each task was examined separately, and participants

were excluded when their accuracy was low for each task separately.

Next, the data were screened for outliers. Our pre-registered plan excluded outliers based on a linear mixed-model analysis for participants in the third quantile of the grand intercept (i.e., participants with the longest average response times). After examining the data from both online and in-person data collection, it became clear that both a minimum response latency and maximum response latency should be employed, as improbable times existed at both ends of the distribution. The minimum response time was set to 160 ms based on Hick's Law (Kvålseth, 2021). The maximum response latency was calculated as two times the mean absolute deviation plus the median calculated separately for each participant. Exclusions were performed at the trial level for these outlier response times, and thus, $n_{total} = 4,217$ participants were included in these analyses across both tasks (see Figure 2 for separation n s for each task). After removal of non-native speakers and low-accuracy participants, 3,544 participants were included in analyses that required both sentence-picture and picture-picture response times.

To ensure equivalence between data collection methods, we evaluated the response times predicted by the fixed effects of the interaction between match (match vs. mismatch) and data collection source (in-person vs. online). We included random intercepts for participants, lab, language, and random slopes for source by lab and source by language. This analysis showed no difference between data sources: $b = 2.41$, $SE = 2.77$, $t(73729.28) = 0.87$, $p = .385$. Therefore, the following analyses did not separate in-person and online data. Table 3 provides a summary of the match advantage by language for the sentence-picture verification task.

Although we combined the two data sets in the final data analysis, it is worth considering that online participants' attention may be easily distracted given the lack of environmental control and experimenter overview. However, this secondary task revealed that online participants had a higher percent correct than in-person participants, $t(3,214.86) = 35.77$, $p < .001$, $M_{online} = 85.46$ ($SD = 14.20$) and $M_{in-person} = 67.71$ ($SD = 16.26$).

Hypothesis 1: Meta-analysis of the orientation effect

The planned meta-analysis examined the effect overall and within languages wherein at least two laboratories had collected data (Arabic, English, German, Norway, Simplified Chinese, Traditional Chinese, Slovakian, and Turkish). Figure 3 showed a significant positive orientation effect across German laboratories ($b = 16.68$, 95% CI [7.75, 25.62]) but did not reveal a significant overall effect ($b = 2.05$, 95% CI [-2.71, 6.82]). Also, a significant negative orientation effect

Table 3 Descriptive summary of sentence-picture verification task by language

Language	Accuracy Percent	Mismatching	Matching	Match Advantage
Arabic	90.65	580.25 (167.53)	581.00 (200.89)	-0.75
Brazilian Portuguese	94.87	641.00 (136.40)	654.50 (146.78)	-13.50
English	95.04	576.75 (124.17)	578.75 (127.87)	-2.00
German	96.53	593.00 (106.75)	576.00 (107.12)	17.00
Greek	92.35	753.50 (225.36)	728.50 (230.91)	25.00
Hebrew	96.73	569.50 (98.59)	574.50 (110.45)	-5.00
Hindi	91.32	638.50 (207.19)	662.00 (228.32)	-23.50
Hungarian	96.47	623.00 (111.94)	643.00 (129.73)	-20.00
Norwegian	96.93	592.50 (126.39)	612.00 (136.03)	-19.50
Polish	96.11	601.00 (139.36)	586.00 (108.23)	15.00
Portuguese	95.01	616.50 (144.55)	607.00 (145.29)	9.50
Serbian	94.78	617.75 (158.64)	635.00 (168.28)	-17.25
Simplified Chinese	92.39	655.00 (170.50)	642.50 (158.64)	12.50
Slovak	96.45	610.50 (125.28)	607.25 (117.87)	3.25
Spanish	94.32	663.00 (147.52)	676.00 (154.19)	-13.00
Thai	93.92	652.50 (177.91)	637.75 (130.10)	14.75
Traditional Chinese	94.41	625.00 (139.36)	620.00 (123.06)	5.00
Turkish	95.38	654.50 (146.04)	637.00 (126.02)	17.50

Average accuracy percentage, Median response times and median absolute deviations (in parentheses) per match condition (Mismatching, Matching); Match advantage (difference in response times)

was found in the Hungarian ($b = -20.00$, 95% CI [-29.60, -10.40]) and the Serbian laboratory ($b = -17.25$, 95% CI [-32.26, -2.24]), although in these languages only a single laboratory participated, so no language-specific meta-analysis was conducted.

Hypothesis 1: Mixed-linear modeling of the orientation effect

First, an intercept-only model of response times with no random intercepts was computed for comparison purposes AIC = 1008828.79. The model with the random intercept by participants was an improvement over this model, AIC = 971783.32. The addition of a target random intercept

improved model fit over the participant intercept-only model, $AIC = 969506.32$. Data collection lab was then added to the model as a random intercept, also showing model improvement, $AIC = 969265.28$, and the random intercept of language was added last, $AIC = 969263.66$ which did not show model improvement at least 2 points change. Last, the fixed effect of match advantage was added with approximately the same fit as the three random-intercept model, $AIC = 969265.06$. This model did not reveal a significant effect of match advantage: $b = -0.17$, $SE = 1.20$, $t(69830.14) = -0.14$, $p = .887$.

We conducted an exploratory mixed-effects model on German data as this was the only language where the confidence interval did not overlap with a zero effect in the meta-analysis. An intercept-only model with random effects for participants, target, and lab was used as a comparison, $AIC = 55828.57$. The addition of the fixed effect of match showed a small improvement over this random-intercept model, $AIC = 55824.52$. Whereas the AIC values indicated a better fit for the model including matching orientation as a predictor, the advantage provided by matching orientation was small: $b = 4.84$, $SE = 4.12$, $t(4085.71) = 1.17$, $p = .241$.

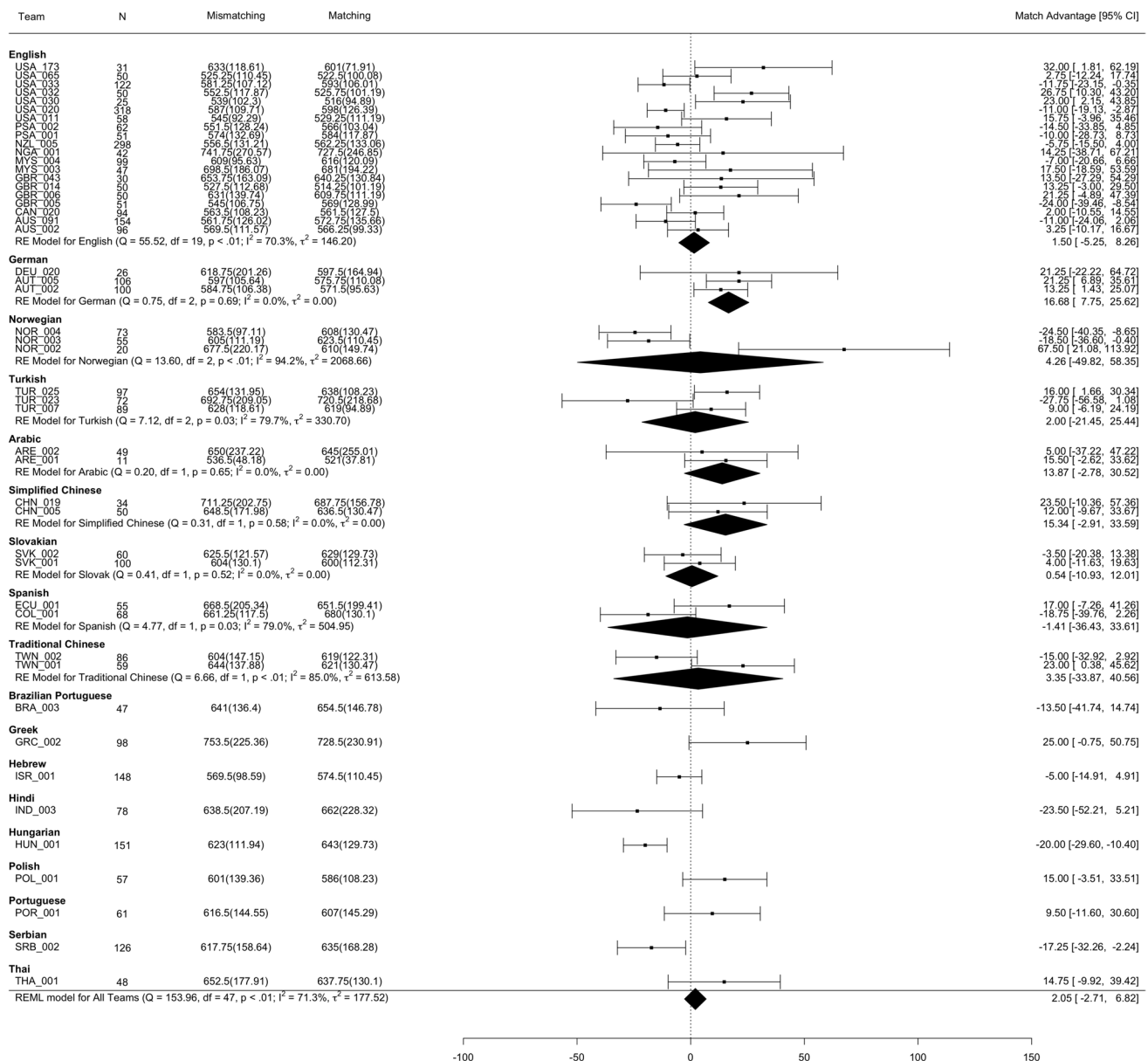


Fig. 3 Meta-analysis on match advantage of object orientation for all languages. Diamonds indicate summary estimates, the midpoint of the diamond indicating the point estimate, and the left and right endpoints

indicating the lower and upper bounds of the confidence interval of the estimated effect size. The lowermost diamond represents the estimate derived from the whole dataset

All the details of the above fixed effects and random intercepts are summarized in Appendix 4.

Hypothesis 2: Mental rotation scores

Using the same steps as described for the sentence-picture verification mixed model, we first started with an intercept-only model with no random effects for comparison, $AIC = 1029362.78$. The addition of random intercepts by subject, $AIC = 979873.47$, by item, $AIC = 977037.64$, by lab, $AIC = 976721.45$, and by language, $AIC = 976717.46$, all subsequently improved model fit. Next, the match effect for object orientation was entered as the fixed effect for mental rotation score, $AIC = 973054.93$, which showed improvement over the random intercepts model. This model showed a significant effect of object orientation, $b = 32.30$, $SE = 0.53$, $t(79585.24) = 61.23$, $p < .001$, such that identical orientations were processed faster than rotated orientations. The point estimates of the orientation effect varied between 23.79–40.24, revealing a range of 14 ms across languages. The coefficients of all mixed-effects models are reported in Appendix 5, along with all estimates presented by language.

Hypothesis 2: Prediction of match advantage

The last analysis included a mixed effects regression model using the interaction of language and mental rotation scores to predict match advantage in the sentence-picture task. First, an intercept-only model was calculated for comparison, $AIC = 42678.66$, which was improved slightly by adding a random intercept by data collection lab, $AIC = 42677.80$. The addition of the fixed effects interaction of language and mental rotation score improved the overall model, $AIC = 42633.44$. English was used as the comparison group for all language comparisons. Neither the mental rotation score nor the interaction of mental rotation score and language were significant, and these results are detailed in Appendix 5.

Discussion

This study aimed to test a global object orientation effect and to estimate the magnitude of object orientation effect in each particular language. The findings of our study did not support the existence of the object orientation effect as an outcome of general cognitive function. Furthermore, our data failed to replicate the effects in English and Chinese, languages in which the effect has been reported previously (Chen et al., 2020; Stanfield & Zwaan, 2001; Zwaan & Pecher, 2012). The

only language in which we found an indication of the orientation effect in the predicted direction was German, but this effect was evident only in the meta-analysis and not in the mixed-effects model approach. Although tangential to our topic, an effect of mental rotation was observed, such that identical orientations were processed faster than rotated orientations. However, the mental rotation score did not predict the object orientation effects nor interact with language. Overall, the failure to replicate the previously reported object orientation effects casts doubt on the existence of the effect as a language-general phenomenon (Kaschak & Madden, 2021). Critically, our experiment has some important methodological strengths. We recruited a large, globally diverse sample of participants. Our pre-registered research protocol was developed in collaboration between experts in object orientation and methodological experts who were agnostic regarding the presence or absence of an effect. For these reasons, the current experiment can be seen as a strong test of the object orientation effect in the sentence picture verification paradigm. Below, we further delineate the lessons and limitations of the methodology and analysis, and discuss theoretical issues related to the orientation effect as an effective probe to investigate the mental simulation process.

Methodological considerations

By examining the failed replications of the object orientation effect in the English-language labs (see Figure 3), researchers can further identify the possible factors that may have contributed to the discrepancies between the results of this project and the original studies. Although our project had a larger sample of English-speaking participants compared to the original studies (i.e., Stanfield & Zwaan, 2001; Zwaan & Pecher, 2012), our English-speaking participants came from multiple countries where the participants' lexical knowledge is not completely consistent with American English. Although we prepared an alternative version of the stimuli for British English, these two versions of English stimuli did not cover all English language backgrounds, such as participants from Malaysia and Africa. Despite the overall non-significant effect in all English-language data, the meta-analysis indicated three significant positive team-based effects (American labs: USA_173, USA_030 and USA_032, see Figure 3) but also three significant negative effects (American labs: USA_33, USA_20, and Great Britain labs: GBR_005, see Figure 3). Future cross-linguistic studies should attempt to balance sample sizes across languages to allow reliable cross-linguistic comparisons.

Regarding the failed replication of Chinese orientation effects, the past study used simpler sentence content

compared to this project. Chen et al. (2020) used the probe sentences in which the target objects were the subject of sentences (e.g., “*The nail was hammered into the wall*”; bold added to mark the subject noun). The Chinese probe sentences in this project were translated from the English sentences used in Stanfield and Zwaan (2001), in which the target objects are the object of sentences (e.g., “*The carpenter hammered the nail into the wall*”; bold added to mark object noun). It is possible that the object orientation effect may be present or stronger when the target objects are the subject of the sentence, rather than the direct object, and future studies could explore this distinction.

Lastly, past studies that employed a secondary task among the experimental trials (Chen et al., 2020; Kaschak & Madden, 2021; Stanfield & Zwaan, 2001) showed a positive object orientation effect. In our study, the memory check did not increase the likelihood to detect the mental simulation effects. In addition, we did not find that mental imagery predicted match advantage, which implies that this strategy to ensure linguistic processing had limited influence in our study.

Analysis considerations

The orientation effects were analyzed using a meta-analytic approach and mixed-effects models. Neither approach revealed an overall effect of object orientation. In the exploratory language-by-language analysis, an indication of the orientation effect was found in the German language data in the meta-analysis. However, the estimate of this effect was small in the mixed-model analysis, and considered in the general context including all the other results, the present exploratory result for German could stem from measurement error or from family-wise error.

The orientation effects were analyzed using a meta-analytic approach and mixed-effects models. Neither approach revealed an overall effect of object orientation. In the language-by-language analysis, a significant orientation effect was found in the German language data in the meta-analysis. The mixed model analysis did not confirm this result because the effect in the German data was not significant according to our pre-registered test criteria. There is considerable debate in the statistical community regarding the precision of the p values computed for linear mixed models (Bolker, 2015). One alternative, less-conservative approach to testing the significance of a fixed effect predictor is assessing the difference in the AIC model fit index between a model that contains a fixed effect predictor and one that does not. Using this approach in an exploratory analysis, we found that the effect of orientation in the German language

data was not negligible, rendering this result compatible with the result obtained for German in the meta-analysis. However, considered in the general context including all the other results, the present exploratory result for German could stem from measurement error (Loken & Gelman, 2017) or from family-wise error.

When a topic area yields inconsistent or small effects, some researchers have questioned the utility of further research. However, research on embodied cognition should continue with the aim of determining the factors behind the variability of the effects. One of these factors could be the nature of the variables used — for instance, categorical versus continuous. The object orientation design is a factorial, congruency paradigm, based on congruent (matching) and incongruent (mismatching) conditions. Another paradigm of similar characteristics, namely the action sentence compatibility effect, similarly failed to replicate in a large-scale study (Morey et al., 2022). Whereas factorial paradigms require the use of categorical variables, other studies have operationalized sensorimotor information using continuous variables, and observed significant effects (Bernabeu, 2022). Since continuous variables contain more information, they may afford more statistical power. Furthermore, in addition to categorical versus continuous predictors, sensorimotor effects are likely to be moderated by factors influencing participants’ attention during experiments (Barsalou, 2019; Noah et al., 2018). Last, due to publication bias, the true size of sensorimotor effects is likely to be smaller than that observed in small-sample studies (Vasishth & Gelman, 2021). Indeed, studying these effects reliably may require samples exceeding 1,000 participants (Bernabeu, 2022). In summary, addressing the above issues may permit the analytic sensitivity needed to observe the presence and causes of object orientation effects.

Theoretical considerations

Scholars interested in mental simulation have investigated whether the human mind processes linguistic content as abstract symbols or as grounded mental representations (Barsalou, 1999, 2008; Zwaan, 2014b). Some of the tasks used to test these theories—such as the sentence-verification task—rely on priming-based logic, whereby a designed sentence generates representations along some dimension (such as orientation) that facilitates or interferes with the processing of the subsequent stimulus (Roelke et al., 2018). Furthermore, embodied cognition theories suggest that the reading of the sentence will activate perceptual experience, thus facilitating a matching object picture and causing interference for a mismatching picture (Kaschak & Madden,

2021; McNamara, 2005). To scrutinize these effects, future studies could augment the sentence-picture verification task to compare the degree of priming based on object orientation with the priming based on other semantic information. The present study constitutes the first large-scale, cross-linguistic approach to the object orientation effect. Cross-linguistic studies are rare in the present topic, and generally in the topic of conceptual/semantic processing. In future studies, the basis for cross-linguistic comparisons in conceptual processing should be expanded, for instance, by studying the lexicosemantic features of the stimuli used, how those differ across languages, and how those differences may influence psycholinguistic processing. For the development of this founding work, the field of linguistic relativity may be useful as a model.

The current study found no evidence that mental rotation scores predicted the match advantage of object orientation across languages, suggesting that mental imagery may be relatively independent of language experience. The results of the sentence-picture verification task showed significant variance across languages, with the significant cases including positive orientation effects in German and negative orientation effects in Hungarian and Serbian. These results suggest that specific language experience may have to be considered when measuring and interpreting the orientation effect.

Further research should compare the size of mental simulation effects with the size of effects that are associated with the symbolic account of conceptual processing. The symbolic account posits that conceptual processing (i.e., the comprehension of the meaning of words) depends on the abstract symbols (e.g., propositions and production rules). So far, some of these comparisons have supported both accounts. However, in some studies, the effects of the symbolic account have been larger than those of the embodied account (Bernabeu, 2022; Louwerse et al., 2015), whereas

the reverse has been true in other studies (Fernandino et al., 2022; Tong et al., 2022).

Limitations

This study reflects the challenges to assess the mental simulation of object orientation across languages, especially when dealing with effects that require large sample sizes. Our data collection deviated from the pre-registered plan because of the COVID-19 pandemic. Due to the lack of participant monitoring online, and an inspection of the data, we post-hoc used filtering on outliers in terms of participants' response times for both too fast (< 160 ms) and too slow responses (2 MAD beyond the median for each participant individually). It is possible that the additional data exclusion could alter the overall reaction time distribution. A mixed-effects model confirmed no difference of response times between in-person and online data. Future studies could evaluate how the task environments alter the magnitude of the orientation effect.

Conclusion

Based on the results of this project, we did not find evidence for a general object orientation effect in a culturally and linguistically diverse sample. Reliable measurement of orientation effects may require a comprehensive sample sizes justification carefully focusing on the interaction of language and cultural backgrounds. Theories of mental simulations should be refined to suggest which perceptual features and linguistic properties contribute to human mental simulations. Our findings on the orientation effects question the theoretical importance of mental simulation in linguistic processing, but they provide directions for further improvements on methodology, analysis, and theories.

Appendix 1

Sensitivity analyses

The R codes for the sensitivity analysis on the trial level were written by Erin M. Buchanan.

Load data and run models

The data for the sensitivity analysis shared the same exclusion criterion for the pre-registered mixed-effects models. The first step is to determine if there is a minimum number of trials required for stable results.

```
number_trials <- SP_V_correct %>%
group_by(Subject, Match) %>%
summarize(count = n())
Run the analyses across different numbers of trials.
models <- list()
for (i in 3:12){
subjects <- number_trials %>%
filter(count >= i) %>%
pull(Subject) %>% unique()
temp_data <- SP_V_correct %>%
filter(Subject %in% subjects)
# match model in the paper
models[[paste("fixed.three.model", i, sep = "_")] <-
lmer(response_time ~ Match + (1|Subject) + (1|Target) +
(1|PSA_ID),
control = lmerControl(optimizer = "bobyqa",
optCtrl = list(maxfun = 1e6)),
data = temp_data)
}
```

View the results

```
B_values <- 1:length(models)
```

```
p_values <- 1:length(models)
for (i in 1:length(models)){
b_values[i] <- round(fixef(models[[i]]["MatchN"],2)
p_values[i] <- apa_p(summary(models[[i]])$coefficients[2,5])
}
#b_values
#p_values
```

b values. These values represent the *b* values found for each run of 3 up to 12 trials.

-0.17, -0.17, -0.17, -0.17, -0.18, -0.12, 0.49, -0.14, 0.67, and 3.11

p values. These values represent the *p* values found for each run of 7 up to 12 trials.

.887, .887, .887, .890, .880, .918, .687, .913, .647, .150

As we can see, the effect is generally negative until participants were required to have 7-12 correct trials. When participants accurately answer all 12 trials the effect is approximately 3 ms. Examination of the *p*-values indicates that no coefficients would have been considered significant.

Calculate the Sensitivity

Given we used all data points, the smallest detectable effect with our standard error and degrees of freedom would have been:

```
qt(p = .05/2, # two tailed
df = summary(models[[1]])$coefficients[2, 3],
lower.tail = F) * summary(models[[1]])$coefficients[2,
2]
## [1] 2.356441
```

Appendix 2

Data collection logs

The log website was initiated since the data collection began. The public logs were updated when a laboratory updated their data for the sequential analysis. The link to access the public site is: https://scgeeker.github.io/PSA002_log_site/index.html

If you want to check the sequential analysis result of a laboratory, at first you have to identify the ID and language of this laboratory from “Overview” page. Next you will navigate to the language page under the banner “Tracking Logs”. For example, you want to see the result of “GBR_005”. Navigate “Tracking Logs ->English”. Search the figure by ID “GBR_005”.

The source files of the public logs are available in the github repository: https://github.com/SCgeeker/PSA002_log_site

All the raw data and log files are compressed in the project OSF repository: <https://osf.io/e428p/>

The R code to conduct the Bayesian sequential analysis is available at “data_seq_analysis.R”. This file can be found at: https://github.com/SCgeeker/PSA002_log_site

Note 1 USA_067, BRA_004 and POL_004 were unavailable because the teams withdrew.

Note 2 Some mistakes happened between the collaborators’ communications and required advanced data wrangling. For example, some AUS_091 participants were assigned to NZL_005. The Rmd file in NZL_005 folder were used to identify the AUS_091 participants’ data then move them to AUS_091 folder.

Datasets

Complete data can be found online with this manuscript or on each collaborators OSF page. Please see the Lab_Info.csv on <https://osf.io/e428p/>.

Fluency test for the online study

At the beginning of the online study, participants will hear the verbal instruction narrated by a native speaker. The original English transcript is as below:

“In this session you will complete two tasks. The first task is called the sentence picture verification task. In this task, you will read a sentence. You will then see a picture. Your job is to verify whether the picture represents an object that was described in the sentence or not. The second task is the picture verification task. In this task you will see two pictures on the screen at the same time and determine whether they are the same or different. Once you have completed both tasks, you will receive a completion code that you can use to verify your participation in the study.”

The fluency test are three multiple choice questions. The question text and the correct answers are as below:

- How many tasks will you run in this session?
- A: 1 *B: 2 C: 3
- When will you get the completion code?
- A: Before the second task B: After the first task *C: After the second task
- What will you do in the sentence-picture verification task?
- A: Confirm two pictures for their objects
- *B: Read a sentence and verify a picture C: Judge sentences for their accuracy

Distributions of scripts

The instructions and experimental scripts are available at the public OSF folder (<https://osf.io/e428p/> “Materials/js” folder in Files). To upload to a jatos server, a script had to be converted to the compatible package. Researchers could do this conversion by “OSWEB” package in OpenSesame. We rent an remote server for the distributions during the data collection period. Any researcher would distribute the scripts on a free jatos server such as MindProbe (<https://www.mindprobe.eu/>).

Appendix 3

Demographic characteristics by language

Table 4 Demographic and sample size characteristics by Language Part 1

Language	SP_{Trials}	PP_{Trials}	SP_N	PP_N	$Demo_N$	$Female_N$	$Male_N$	M_{Age}	SD_{Age}
Arabic	1248	1248	52	52	53	0	0	38.00	NaN
Arabic	1296	1296	54	54	54	42	12	26.51	18.59
Brazilian Portuguese	1200	1200	50	50	50	36	13	30.80	8.73
English	2376	2376	99	99	103	46	37	20.14	3.32
English	3840	3840	160	160	160	127	25	26.03	11.55
English	2352	2376	98	99	104	54	40	20.26	3.66
English	1272	1272	53	53	76	57	13	19.96	3.90
English	1200	1200	50	50	51	37	13	20.14	2.46
English	1200	1200	50	50	58	46	11	18.74	1.62
English	720	720	30	30	32	15	11	25.70	9.40
English	1200	1224	50	51	52	38	11	22.56	3.90
English	2400	2400	100	100	109	65	30	20.73	2.00
English	1248	1248	52	52	52	24	22	23.94	11.29
English	7680	7680	320	320	320	244	56	23.21	5.43
English	1248	1272	52	53	71	50	12	18.89	0.95

SP Sentence Picture Verification, *PP* Picture Picture Verification. Sample sizes for demographics may be higher than the sample size for the this study, as participants could have only completed the bundled experiment. Additionally, not all entries could be unambiguously matched by lab ID, and therefore, demographic sample sizes could also be less than data collected. Each row represents a single lab

Table 5 Demographic and sample size characteristics by Lab Part 2

Language	SP_{Trials}	PP_{Trials}	SP_N	PP_N	$Demo_N$	$Female_N$	$Male_N$	M_{Age}	SD_{Age}
English	1536	1536	64	64	102	79	11	19.82	2.42
English	264	264	11	11	12	9	2	20.36	1.91
English	288	288	12	12	12	6	5	21.17	1.19
English	1512	1512	63	63	63	30	23	22.34	11.55
English	7980	8064	333	336	403	258	76	19.63	2.12
English	648	648	27	27	31	20	3	36.00	0.96
English	1209	1224	51	51	51	30	21	19.29	1.51
English	3000	3024	125	126	129	90	25	20.06	1.36
English	1200	1200	50	50	61	35	15	18.86	1.63
English	816	744	34	31	3	0	3	19.67	0.58
German	2400	2400	100	100	114	0	1	20.94	2.56
German	2592	2592	108	108	108	80	22	22.18	4.26
German	624	624	26	26	26	18	3	23.88	3.39
Greek	2376	2376	99	99	109	0	0	33.86	11.30
Hebrew	3576	3571	149	149	181	0	0	24.25	9.29
Hindi	1896	1896	79	79	86	57	27	21.66	3.46
Magyar	3610	3816	151	159	168	3	1	21.50	2.82
Norwegian	504	504	21	21	21	12	8	30.10	8.58
Norwegian	1320	1320	55	55	53	1	1	23.55	6.25
Norwegian	1752	1752	73	73	80	0	0	22.00	4.38

SP Sentence Picture Verification, *PP* Picture Picture Verification. Sample sizes for demographics may be higher than the sample size for the this study, as participants could have only completed the bundled experiment. Additionally, not all entries could be unambiguously matched by lab ID, and therefore, demographic sample sizes could also be less than data collected

Table 6 Demographic and sample size characteristics by Lab Part 3

Language	<i>SP</i> _{Trials}	<i>PP</i> _{Trials}	<i>SP</i> _N	<i>PP</i> _N	<i>Demo</i> _N	<i>Female</i> _N	<i>Male</i> _N	<i>M</i> _{Age}	<i>SD</i> _{Age}
Polish	1368	1368	57	57	146	0	0	23.25	7.96
Portuguese	1488	1464	62	61	55	26	23	30.74	9.09
Serbian	3120	3120	130	130	130	108	21	21.38	4.50
Simplified Chinese	1200	1200	50	50	57	0	0	18.66	3.92
Simplified Chinese	840	816	35	34	39	0	1	25.17	5.44
Slovak	2419	2400	101	100	103	1	0	21.59	2.51
Slovak	1462	1199	61	50	222	0	0	21.96	2.14
Spanish	1680	1656	70	69	70	0	0	21.36	3.36
Spanish	1440	1440	60	60	76	0	0	22.10	4.30
Thai	1200	1152	50	48	50	29	9	21.54	3.81
Traditional Chinese	1440	1440	60	60	70	45	14	20.73	1.21
Traditional Chinese	2160	2160	90	90	116	24	32	21.04	3.66
Turkish	2184	2184	91	91	93	0	0	20.92	2.93
Turkish	1896	1896	79	79	80	36	14	21.58	8.64
Turkish	2376	2352	99	98	101	0	0	21.63	2.19

SP Sentence Picture Verification, *PP* Picture Picture Verification. Sample sizes for demographics may be higher than the sample size for this study, as participants could have only completed the bundled experiment. Additionally, not all entries could be unambiguously matched by lab ID, and therefore, demographic sample sizes could also be less than data collected

Appendix 4

Model Estimates for Mental Simulation

All model estimates are given below for the planned mixed linear model to estimate the matching effect for object orientation in the sentence picture verification task.

Note. Fixed indicates fixed parameters in multilevel models, while “ran_pars” indicates the random intercepts included in the model.

Table 7 Intercept only object orientation results

Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>p</i>
(Intercept)	654.71	0.84	775.11	<.001

Table 8 Subject-random intercept object orientation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	654.26	2.69	243.34	3,787.12	<.001
ran_pars	Subject	sd__(Intercept)	161.40	NA	NA	NA	
ran_pars	Residual	sd__Observation	165.05	NA	NA	NA	

Table 9 Subject and item-random intercept object orientation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	655.47	5.16	126.97	84.63	<.001
ran_pars	Subject	sd__(Intercept)	161.37	NA	NA	NA	
ran_pars	Target	sd__(Intercept)	30.54	NA	NA	NA	
ran_pars	Residual	sd__Observation	162.17	NA	NA	NA	

Table 10 Subject, item, and lab-random intercept object orientation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	659.63	9.67	68.24	65.54	<.001
ran_pars	Subject	sd__(Intercept)	153.76	NA	NA	NA	
ran_pars	Target	sd__(Intercept)	30.56	NA	NA	NA	
ran_pars	PSA_ID	sd__(Intercept)	55.76	NA	NA	NA	
ran_pars	Residual	sd__Observation	162.17	NA	NA	NA	

Table 11 Subject, item, lab, and language-random intercept object orientation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	672.75	11.88	56.61	23.94	<.001
ran_pars	Subject	sd__(Intercept)	153.78	NA	NA	NA	
ran_pars	Target	sd__(Intercept)	30.56	NA	NA	NA	
ran_pars	PSA_ID	sd__(Intercept)	48.60	NA	NA	NA	
ran_pars	Language	sd__(Intercept)	25.52	NA	NA	NA	
ran_pars	Residual	sd__Observation	162.17	NA	NA	NA	

Table 12 Fixed effects object orientation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	659.71	9.68	68.12	66.05	<.001
fixed	NA	MatchN	-0.17	1.20	-0.14	69,830.14	.887
ran_pars	Subject	sd__(Intercept)	153.76	NA	NA	NA	
ran_pars	Target	sd__(Intercept)	30.56	NA	NA	NA	
ran_pars	PSA_ID	sd__(Intercept)	55.76	NA	NA	NA	
ran_pars	Residual	sd__Observation	162.17	NA	NA	NA	

Table 13 Random effects german object orientation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	631.96	19.65	32.15	1.74	.002
ran_pars	Subject	sd__(Intercept)	129.89	NA	NA	NA	
ran_pars	Target	sd__(Intercept)	33.04	NA	NA	NA	
ran_pars	PSA_ID	sd__(Intercept)	28.11	NA	NA	NA	
ran_pars	Residual	sd__Observation	134.89	NA	NA	NA	

Table 14 Fixed effects german object orientation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	629.52	19.74	31.90	1.77	.002
fixed	NA	MatchN	4.84	4.12	1.17	4,085.71	.241
ran_pars	Subject	sd__(Intercept)	129.90	NA	NA	NA	
ran_pars	Target	sd__(Intercept)	33.06	NA	NA	NA	
ran_pars	PSA_ID	sd__(Intercept)	28.05	NA	NA	NA	
ran_pars	Residual	sd__Observation	134.88	NA	NA	NA	

Appendix 5

Model estimates for mental rotation

All model estimates are given below for the mixed linear model for the prediction of mental rotation scores by

orientation, and the effects of predicting mental simulation effects (object orientation) with the mental rotation scores.

Note. Fixed indicates fixed parameters in multilevel models, while “ran_pars” indicates the random intercepts included in the model.

Table 15 Intercept only mental rotation results

Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>p</i>
(Intercept)	589.10	0.40	1,485.40	<.001

Table 16 Subject-random intercept mental rotation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	588.28	1.35	436.44	3,957.55	<.001
ran_pars	Subject	sd__(Intercept)	83.04	NA	NA	NA	
ran_pars	Residual	sd__Observation	79.04	NA	NA	NA	

Table 17 Subject and item-random intercept mental rotation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	589.21	2.70	217.83	79.98	<.001
ran_pars	Subject	sd__(Intercept)	82.94	NA	NA	NA	
ran_pars	Picture1	sd__(Intercept)	16.26	NA	NA	NA	
ran_pars	Residual	sd__Observation	77.56	NA	NA	NA	

Table 18 Subject, item, and lab-random intercept mental rotation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	590.18	5.16	114.35	69.45	<.001
ran_pars	Subject	sd__(Intercept)	78.36	NA	NA	NA	
ran_pars	Picture1	sd__(Intercept)	16.32	NA	NA	NA	
ran_pars	PSA_ID	sd__(Intercept)	30.12	NA	NA	NA	
ran_pars	Residual	sd__Observation	77.56	NA	NA	NA	

Table 19 Subject, item, lab, and language-random intercept mental rotation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	596.71	6.98	85.46	22.17	<.001
ran_pars	Subject	sd__(Intercept)	78.37	NA	NA	NA	
ran_pars	Picture1	sd__(Intercept)	16.32	NA	NA	NA	
ran_pars	PSA_ID	sd__(Intercept)	24.00	NA	NA	NA	
ran_pars	Language	sd__(Intercept)	19.29	NA	NA	NA	
ran_pars	Residual	sd__Observation	77.56	NA	NA	NA	

Table 20 Fixed effects mental rotation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	581.52	7.02	82.82	22.57	<.001
fixed	NA	IdenticalN	32.30	0.53	61.23	79,585.24	<.001
ran_pars	Subject	sd_(Intercept)	78.24	NA	NA	NA	
ran_pars	Picture1	sd_(Intercept)	16.89	NA	NA	NA	
ran_pars	PSA_ID	sd_(Intercept)	24.01	NA	NA	NA	
ran_pars	Language	sd_(Intercept)	19.33	NA	NA	NA	
ran_pars	Residual	sd__Observation	75.80	NA	NA	NA	

Table 21 Language specific mental rotation results

Language	Coefficient (<i>b</i>)	<i>SE</i>
Arabic	28.27	3.36
English	33.02	0.77
German	31.38	1.91
Brazilian Portuguese	23.79	4.85
Simplified Chinese	32.40	3.64
Spanish	40.24	3.75
Greek	30.59	3.67
Hungarian	25.43	2.57
Hindi	35.83	3.86
Hebrew	29.02	2.43
Norwegian	28.12	2.59
Polish	38.74	3.51
Portuguese	34.67	4.05
Serbian	25.93	2.75
Slovak	33.34	2.61
Thai	34.99	4.13
Turkish	37.46	2.29
Traditional Chinese	30.31	2.45

Table 22 Intercept only predicting mental simulation results

Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>p</i>
(Intercept)	-0.74	1.67	-0.44	.661

Table 23 Lab-random intercept predicting mental simulation results

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	-0.74	1.67	-0.44	3,543.00	.661
ran_pars	PSA_ID	sd_(Intercept)	0.00	NA	NA	NA	
ran_pars	Residual	sd__Observation	99.67	NA	NA	NA	

Table 24 Fixed effects interaction language and rotation predicting mental simulation results part 1

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	(Intercept)	-3.43	3.07	-1.12	3,510.00	.264
fixed	NA	LanguageArabic	16.27	16.34	1.00	3,510.00	.319
fixed	NA	LanguageBrazilian Portuguese	-17.00	16.63	-1.02	3,510.00	.307
fixed	NA	LanguageGerman	20.65	9.17	2.25	3,510.00	.024
fixed	NA	LanguageGreek	7.52	12.58	0.60	3,510.00	.550
fixed	NA	LanguageHindi	1.24	15.39	0.08	3,510.00	.936
fixed	NA	LanguageHungarian	-10.30	10.01	-1.03	3,510.00	.304
fixed	NA	LanguageNorwegian	19.66	10.31	1.91	3,510.00	.057
fixed	NA	LanguagePolish	-5.67	17.87	-0.32	3,510.00	.751
fixed	NA	LanguagePortuguese	-15.90	17.21	-0.92	3,510.00	.356
fixed	NA	LanguageSerbian	-0.48	11.20	-0.04	3,510.00	.966
fixed	NA	LanguageSimplified Chinese	22.70	14.45	1.57	3,510.00	.116
fixed	NA	LanguageSlovak	-0.34	11.58	-0.03	3,510.00	.976
fixed	NA	LanguageSpanish	5.39	11.51	0.47	3,510.00	.640
fixed	NA	LanguageThai	27.53	21.22	1.30	3,510.00	.195
fixed	NA	LanguageTraditional Chinese	-8.24	11.20	-0.74	3,510.00	.462
fixed	NA	LanguageTurkish	10.66	8.57	1.24	3,510.00	.214
fixed	NA	Imagery	0.04	0.05	0.79	3,510.00	.432

Table 25 Fixed effects interaction language and rotation predicting mental simulation results part 2

Effect	Group	Term	Estimate (<i>b</i>)	<i>SE</i>	<i>t</i>	<i>df</i>	<i>p</i>
fixed	NA	LanguageArabic:Imagery	-0.37	0.24	-1.55	3,510.00	.122
fixed	NA	LanguageBrazilian Portuguese:Imagery	0.18	0.29	0.62	3,510.00	.536
fixed	NA	LanguageGerman:Imagery	-0.37	0.19	-1.96	3,510.00	.050
fixed	NA	LanguageGreek:Imagery	-0.07	0.20	-0.36	3,510.00	.718
fixed	NA	LanguageHindi:Imagery	-0.36	0.27	-1.34	3,510.00	.181
fixed	NA	LanguageHungarian:Imagery	0.10	0.22	0.45	3,510.00	.653
fixed	NA	LanguageNorwegian:Imagery	-0.07	0.19	-0.35	3,510.00	.726
fixed	NA	LanguagePolish:Imagery	0.29	0.31	0.91	3,510.00	.363
fixed	NA	LanguagePortuguese:Imagery	-0.05	0.32	-0.15	3,510.00	.884
fixed	NA	LanguageSerbian:Imagery	-0.12	0.22	-0.56	3,510.00	.576
fixed	NA	LanguageSimplified Chinese:Imagery	-0.32	0.24	-1.32	3,510.00	.187
fixed	NA	LanguageSlovak:Imagery	0.12	0.21	0.57	3,510.00	.568
fixed	NA	LanguageSpanish:Imagery	0.05	0.17	0.28	3,510.00	.781
fixed	NA	LanguageThai:Imagery	-0.50	0.38	-1.32	3,510.00	.186
fixed	NA	LanguageTraditional Chinese:Imagery	0.13	0.23	0.59	3,510.00	.556
fixed	NA	LanguageTurkish:Imagery	-0.15	0.14	-1.09	3,510.00	.274
ran_pars	PSA_ID	sd__(Intercept)	0.00	NA	NA	NA	
ran_pars	Residual	sd__Observation	99.73	NA	NA	NA	

Supplementary Information The online version contains supplementary material available at <https://doi.org/10.1007/s12144-025-08304-x>.

Acknowledgement We thank the suggestions from the editor and two reviewers on our first and second proposals. CB would like to thank Tyler McGee for help with data collection. PB would like to thank Liam Morgillo for help with data collection. The team would like to thank Zhong Chen for data collection.

Author contributions The authors made the following contributions. Sau-Chin Chen: Conceptualization, Data curation, Formal analy-

sis, Investigation, Methodology, Resources, Software, Supervision, Validation, Visualization, Writing - original draft, Writing - review & editing; Erin Buchanan: Formal analysis, Project administration, Resources, Software, Validation, Writing - review & editing; Zoltan Kececs: Project administration, Writing - review & editing; Jeremy K. Miller: Project administration, Resources, Supervision, Writing - review & editing; Anna Szabelska: Project administration, Writing - original draft, Writing - review & editing; Balazs Aczel: Investigation, Methodology, Resources, Writing - review & editing; Pablo Bernabeu: Investigation, Methodology, Visualization, Writing - review & editing; Patrick Forscher: Methodology, Writing - review & editing; Attila Szuts: Investigation, Methodology, Resources, Writing - review

& editing; Zahir Vally: Investigation, Resources, Writing - review & editing; Ali H. Al-Hoorie: Investigation, Resources, Writing - review & editing; Mai Helmy: Investigation, Resources, Writing - review & editing; Caio Santos Alves da Silva: Investigation, Resources, Writing - review & editing; Luana Oliveira da Silva: Investigation, Resources, Writing - review & editing; Yago Luksevičius de Moraes: Investigation, Resources, Writing - review & editing; Rafael Ming Chi Santos Hsu: Investigation, Resources, Writing - review & editing; Anthonieta Looman Mafra: Investigation, Resources, Writing - review & editing; Jaroslava V. Valentova: Investigation, Resources, Writing - review & editing; Marco Antonio Correa Varella: Investigation, Resources, Writing - review & editing; Barnaby Dixon: Investigation, Writing - review & editing; Kim Peters: Investigation, Writing - review & editing; Nik Steffens: Investigation, Writing - review & editing; Omid Ghasemi: Investigation, Writing - review & editing; Andrew Roberts: Investigation, Writing - review & editing; Robert M. Ross: Investigation, Writing - review & editing; Ian D. Stephen: Investigation, Writing - review & editing; Marina Milyavskaya: Investigation, Writing - review & editing; Kelly Wang: Investigation, Writing - review & editing; Kaitlyn M. Werner: Investigation, Writing - review & editing; Dawn Liu Holford: Investigation, Writing - review & editing; Miroslav Sirota: Investigation, Writing - review & editing; Thomas Rhys Evans: Investigation, Writing - review & editing; Dermot Lynott: Investigation, Writing - review & editing; Bethany M. Lane: Investigation, Writing - review & editing; Danny Riis Sahlholdt: Investigation, Writing - review & editing; Glenn P. Williams: Investigation, Writing - review & editing; Chrystalle B. Y. Tan: Investigation, Writing - review & editing; Alicia Foo: Investigation, Writing - review & editing; Steve M. J. Janssen: Investigation, Writing - review & editing; Nwadiogo Chisom Arinze: Investigation, Writing - review & editing; Izuchukwu Lawrence Gabriel Ndukaihe: Investigation, Writing - review & editing; David Moreau: Investigation, Writing - review & editing; Brianna Jurosic: Investigation, Writing - review & editing; Brynna Leach: Investigation, Writing - review & editing; Savannah Lewis: Investigation, Writing - review & editing; Peter R. Mallik: Investigation, Writing - review & editing; Kathleen Schmidt: Investigation, Resources, Writing - review & editing; William J. Chopik: Investigation, Writing - review & editing; Leigh Ann Vaughn: Investigation, Writing - review & editing; Manyu Li: Investigation, Writing - review & editing; Carmel A. Levitan: Investigation, Writing - review & editing; Daniel Storage: Investigation, Writing - review & editing; Carlota Batres: Investigation, Writing - review & editing; Tyler McGee: Investigation, Writing - review & editing; Janina Enachescu: Investigation, Resources, Writing - review & editing; Jerome Olsen: Investigation, Resources, Writing - review & editing; Martin Voracek: Investigation, Resources, Writing - review & editing; Claus Lamm: Investigation, Resources, Writing - review & editing; Ekaterina Pronizius: Investigation, Writing - review & editing; Tilli Ripp: Investigation, Writing - review & editing; Jan Philipp Röer: Investigation, Writing - review & editing; Roxane Schnepfer: Investigation, Writing - review & editing; Marietta Papadatou-Pastou: Investigation, Resources, Writing - review & editing; Aviv Mokady: Investigation, Resources, Writing - review & editing; Niv Reggev: Investigation, Resources, Writing - review & editing; Priyanka Chandel: Investigation, Resources, Writing - review & editing; Pratibha Kujur: Investigation, Resources, Writing - review & editing; Babita Pande: Investigation, Resources, Supervision, Writing - review & editing; Arti Parganiha: Investigation, Resources, Supervision, Writing - review & editing; Noorshama Parveen: Investigation, Resources, Writing - review & editing; Sraddha Pradhan: Investigation, Resources, Writing - review & editing; Margaret Messiah Singh: Investigation, Writing - review & editing; Max Korbacher: Investigation, Writing - review & editing; Jonas R. Kunst: Investigation, Resources, Writing - review & editing; Christian K. Tammes: Investigation, Resources, Writing - review & editing; Frederike S. Woelfert: Investigation, Writing - review & editing; Kristoffer Klevjer: Investigation, Writing - review & editing; Sarah E. Martiny: Investigation,

Writing - review & editing; Gerit Pfuhl: Investigation, Resources, Writing - review & editing; Sylwia Adamus: Investigation, Resources, Writing - review & editing; Krystian Barzykowski: Investigation, Resources, Supervision, Writing - review & editing; Katarzyna Filip: Investigation, Resources, Writing - review & editing; Patrícia Arriaga: Funding acquisition, Investigation, Resources, Writing - review & editing; Vasilije Gvozdenović: Investigation, Resources, Writing - review & editing; Vanja Ković: Investigation, Resources, Writing - review & editing; Fei Gao: Investigation, Resources, Writing - review & editing; Jingxiang Li: Investigation, Resources, Writing - review & editing; Jozef Bavoľár: Investigation, Resources, Writing - review & editing; Monika Hricová: Investigation, Resources, Writing - review & editing; Pavol Kačmár: Investigation, Resources, Writing - review & editing; Matúš Adamkovič: Investigation, Writing - review & editing; Peter Babinčák: Investigation, Writing - review & editing; Gabriel Baník: Investigation, Writing - review & editing; Ivan Ropovik: Investigation, Writing - review & editing; Danilo Zambrano Ricaurte: Investigation, Resources, Writing - review & editing; Sara Álvarez-Solas: Investigation, Resources, Writing - review & editing; Harry Manley: Investigation, Resources, Writing - review & editing; Panita Suavansri: Investigation, Resources, Writing - review & editing; Chun-Chia Kung: Investigation, Resources, Writing - review & editing; Belemir Çoktok: Investigation, Resources, Writing - review & editing; Asil Ali Özdoğru: Investigation, Resources, Writing - review & editing; Çağlar Solak: Investigation, Writing - review & editing; Sinem Söylemez: Investigation, Writing - review & editing; Sami Çoksan: Investigation, Resources, Writing - review & editing; İlker Dalgat: Resources, Writing - review & editing; Mahmoud Elsherif: Writing - review & editing; Martin Vasilev: Writing - review & editing; Vinka Mlakic: Resources, Writing - review & editing; Elisabeth Oberzaucher: Resources, Writing - review & editing; Stefan Stieger: Resources, Writing - review & editing; Selina Volsa: Resources, Writing - review & editing; Erica D. Musser: Investigation, Writing - review & editing; Janis Zickfeld: Resources, Writing - review & editing; Christopher R. Chartier: Investigation, Project administration, Writing - review & editing.

Funding Matúš Adamkovič was supported by APVV-20-0319; Ivan Ropovik was supported by APVV-23-0421 and Mediated Society (MEDIS:ON) CZ.02.01.01/00/23_025/0008713, co-financed by the European Union; Peter Babinčák was supported by APVV-23-0647; Robert M. Ross was supported by Australian Research Council (grant number: DP180102384) and the John Templeton Foundation (grant ID: 62631); Zoltan Kekecs was supported by János Bolyai Research Scholarship of the Hungarian Academy of Science; Mahmoud Elsherif was supported by Leverhulme Trust; Glenn P. Williams was supported by Leverhulme Trust Research Project Grant (RPG-2016-093); Krystian Barzykowski was supported by National Science Centre, Poland (2019/35/B/HS6/00528); Gabriel Baník was supported by PRIMUS/20/HUM/009; Patrícia Arriaga was supported by Portuguese National Foundation for Science and Technology (FCT UID/PSI/03125/2019); Monika Hricová was supported by VEGA 1/0145/23.

Data availability The project dataset and analysis scripts are available at https://osf.io/rgmky/?view_only=5210c1ec522049db8535f40bc9a07113. Materials are available at https://osf.io/74ujf/?view_only=7094d5af798047b3a2ec01ea4930e3de.

Declarations

Ethical approval Authors who collected data on site and online had the ethical approval/agreement from their local institutions. The latest status of ethical approval for all the participating authors is available at the public OSF folder (<https://osf.io/eknvw/>). “IRB approvals” in Files).

This study is the collaborative research project launched by Psycho-

logical Science Accelerator (Moshontz et al., 2018). The initial preregistered plan is available at https://osf.io/preprints/psyarxiv/t2pju_v1.

Conflict of interest The authors declare no conflicts of interest relevant to the content of the manuscript.

References

- Barsalou, L. W. (1999). Perceptual symbol systems. *Behavioral and Brain Sciences*, 22, 577–660. <https://doi.org/10.1017/S0140525X99002149>
- Barsalou, L. W. (2008). Grounded cognition. *Annual Review of Psychology*, 59, 617–645. <https://doi.org/10.1146/annurev.psych.59.103006.093639>
- Barsalou, L. W. (2019). Establishing generalizable mechanisms. *Psychological Inquiry*, 30(4), 220–230. <https://doi.org/10.1080/1047840X.2019.1693857>
- Barsalou, L. W. (2020). Challenges and opportunities for grounding cognition. *Journal of Cognition*, 3(1), 31. <https://doi.org/10.5334/joc.116>
- Bernabeu, P. (2022). Language and sensorimotor simulation in conceptual processing: Multilevel analysis and statistical power. <https://doi.org/10.17635/LANCASTER/THESIS/1795>
- Bolker, B. M. (2015). Linear and generalized linear mixed models. In G. A. Fox, S. Negrete-Yankelevich, & V. J. Sosa, (Eds.), pp. 309–333. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780199672547.003.0014>
- Chen, S.-C., de Koning, B. B., & Zwaan, R. A. (2020). Does object size matter with regard to the mental simulation of object orientation? *Experimental Psychology*, 67(1), 56–72. <https://doi.org/10.1027/1618-3169/a000468>
- Chu, M., & Kita, S. (2008). Spontaneous gestures during mental rotation tasks: Insights into the microdevelopment of the motor strategy. *Journal of Experimental Psychology: General*, 137(4), 706–723. <https://doi.org/10.1037/a0013157>
- Cohen, D., & Kubovy, M. (1993). Mental rotation, mental representation, and flat slopes. *Cognitive Psychology*, 25, 351–382. <https://doi.org/10.1006/cogp.1993.1009>
- Connell, L. (2007). Representing object colour in language comprehension. *Cognition*, 102, 476–485. <https://doi.org/10.1016/j.cognition.2006.02.009>
- De Koning, B. B., Wassenburg, S. I., Bos, L. T., & Van der Schoot, M. (2017). Mental simulation of four visual object properties: Similarities and differences as assessed by the sentence-picture verification task. *Journal of Cognitive Psychology*, 29(4), 420–432. <https://doi.org/10.1080/20445911.2017.1281283>
- De Koning, B. B., Wassenburg, S. I., Bos, L. T., & Van der Schoot, M. (2017). Size does matter: Implied object size is mentally simulated during language comprehension. *Discourse Processes*, 54(7), 493–503. <https://doi.org/10.1080/0163853X.2015.1119604>
- Engelen, J. A. A., Bouwmeester, S., de Bruin, A. B. H., & Zwaan, R. A. (2011). Perceptual simulation in developing language comprehension. *Journal of Experimental Child Psychology*, 110(4), 659–675. <https://doi.org/10.1016/j.jecp.2011.06.009>
- Fernandino, L., Tong, J.-Q., Conant, L. L., Humphries, C. J., & Binder, J. R. (2022). Decoding the information structure underlying the neural representation of concepts. *Proceedings of the National Academy of Sciences of the United States of America*, 119(6), Article e2108091119. <https://doi.org/10.1073/pnas.2108091119>
- Frick, A., & Möhring, W. (2013). Mental object rotation and motor development in 8- and 10-month-old infants. *Journal of Experimental Child Psychology*, 115(4), 708–720. <https://doi.org/10.1016/j.jecp.2013.04.001>
- Koster, D., Cadierno, T., & Chiarandini, M. (2018). Mental simulation of object orientation and size: A conceptual replication with second language learners. *Journal of the European Second Language Association*. <https://doi.org/10.22599/jesla.39>
- Kvålseth, T. O. (2021). Hick's law equivalent for reaction time to individual stimuli. *British Journal of Mathematical and Statistical Psychology*, 74(S1), 275–293. <https://doi.org/10.1111/bmsp.12232>
- Li, Y., & Shang, L. (2017). An ERP study on the mental simulation of implied object color information during Chinese sentence comprehension. *Journal of Psychological Science*, 40(1), 29–36. <http://doi.org/10.16719/j.cnki.1671-6981.20170105>
- Loken, E., & Gelman, A. (2017). Measurement error and the replication crisis. *Science*, 355(6325), 584–585. <https://doi.org/10.1126/science.aal3618>
- Louwerse, M. M., Hutchinson, S., Tillman, R., & Recchia, G. (2015). Effect size matters: the role of language statistics and perceptual simulation in conceptual processing. *Language, Cognition and Neuroscience*, 30(4), 430–447. <https://doi.org/10.1080/23273798.2014.981552>
- Mathôt, S., & March, J. (2022). Conducting linguistic experiments online with OpenSesame and OSWeb. *Language Learning*, 72(4), 1017–1048. <https://doi.org/10.1111/lang.12509>
- McNamara, T. P. (2005). Semantic priming: Perspectives from memory and word recognition. Psychology Press.
- Morey, R. D., Kaschak, M. P., Diez-Álamo, A. M., Glenberg, A. M., Zwaan, R. A., Lakens, D., Ibáñez, A., García, A., Gianelli, C., Jones, J. L., Madden, J., Alifano, F., Bergen, B., Bloxsom, N. G., Bub, D. N., Cai, Z. G., Chartier, C. R., Chatterjee, A., Conwell, E., & Ziv-Crispel, N. (2022). A pre-registered, multi-lab non-replication of the action-sentence compatibility effect (ACE). *Psychonomic Bulletin & Review*, 29(2), 613–626. <https://doi.org/10.3758/s13423-021-01927-8>
- Moshontz, H., Campbell, L., Ebersole, C. R., IJzerman, H., Urry, H. L., Forscher, P. S., Grahe, J. E., McCarthy, R. J., Musser, E. D., Antfolk, J., Castille, C. M., Evans, T. R., Fiedler, S., Flake, J. K., Forero, D. A., Janssen, S. M. J., Keene, J. R., Protzko, J., Aczel, B., ... Chartier, C. R. (2018). The psychological science accelerator: Advancing psychology through a distributed collaborative network. *Advances in Methods and Practices in Psychological Science*, 1(4), 501–515. <https://doi.org/10.1177/2515245918797607>
- Newman, J. (2002). 1. A cross-linguistic overview of the posture verbs “Sit,” “Stand,” and “Lie.” In J. Newman (Ed.), *Typological Studies in Language* (Vol. 51, pp. 1–24). John Benjamins Publishing Company. <https://doi.org/10.1075/tsl.51.02new>
- Noah, T., Schul, Y., & Mayo, R. (2018). When both the original study and its failed replication are correct: Feeling observed eliminates the facial-feedback effect. *Journal of Personality and Social Psychology*, 114(5), 657–664. <https://doi.org/10.1037/pspa0000121>
- Ostarek, M., & Huettig, F. (2019). Six challenges for embodiment research. *Current Directions in Psychological Science*, 28(6), 593–599. <https://doi.org/10.1177/0963721419866441>
- Pecher, D., van Dantzig, S., Zwaan, R. A., & Zeelenberg, R. (2009). Language comprehenders retain implied shape and orientation of objects. *The Quarterly Journal of Experimental Psychology*, 62(6), 1108–1114. <https://doi.org/10.1080/17470210802633255>
- Pouw, W. T. J. L., de Nooijer, J. A., van Gog, T., Zwaan, R. A., & Paas, F. (2014). Toward a more embedded/extended perspective on the cognitive function of gestures. *Frontiers in Psychology*, 5. <https://doi.org/10.3389/fpsyg.2014.00359>
- Roelke, A., Franke, N., Biemann, C., Radach, R., Jacobs, A. M., & Hofmann, M. J. (2018). A novel co-occurrence-based approach to predict pure associative and semantic priming. *Psychonomic Bulletin & Review*, 25(4), 1488–1493. <https://doi.org/10.3758/s13423-018-1453-6>

- Rommers, J., Meyer, A. S., & Huettig, F. (2013). Object shape and orientation do not routinely influence performance during language processing. *Psychological Science*, 24(11), 2218–2225. <https://doi.org/10.1177/0956797613490746>
- Sato, M., Schafer, A. J., & Bergen, B. K. (2013). One word at a time: Mental representations of object shape change incrementally during sentence processing. *Language and Cognition*, 5(04), 345–373. <https://doi.org/10.1515/langcog-2013-0022>
- Scorolli, C. (2014). Embodiment and language. In L. Shapiro (Ed.), *The Routledge handbook of embodied cognition* (pp. 145–156). Routledge.
- Šetić, M., & Domijan, D. (2017). Numerical congruency effect in the sentence-picture verification task. *Experimental Psychology*, 64(3), 159–169. <https://doi.org/10.1027/1618-3169/a000358>
- Stanfield, R. A., & Zwaan, R. A. (2001). The effect of implied orientation derived from verbal context on picture recognition. *Psychological Science*, 12(2), 153–156. <https://doi.org/10.1111/1467-9280.00326>
- Tong, J., Binder, J. R., Humphries, C., Mazurchuk, S., Conant, L. L., & Fernandino, L. (2022). A distributed network for multimodal experiential representation of concepts. *The Journal of Neuroscience*, 42(37), 7121–7130. <https://doi.org/10.1523/JNEUROSCI.1243-21.2022>
- Vasishth, S., & Gelman, A. (2021). How to embrace variation and accept uncertainty in linguistic and psycholinguistic data analysis. *Linguistics*, 59(5), 1311–1342. <https://doi.org/10.1515/ling-2019-0051>
- Vazire, S. (2018). Implications of the credibility revolution for productivity, creativity, and progress. *Perspectives on Psychological Science*, 13(4), 411–417. <https://doi.org/10.1177/1745691617751884>
- Verkerk, A. (2014). The correlation between motion event encoding and path verb lexicon size in the indo-european language family. *Folia Linguistica*, 35(1). <https://doi.org/10.1515/flih.2014.009>
- Wilson, M. (2002). Six views of embodied cognition. *Psychonomic Bulletin & Review*, 9(4), 625–636. <https://doi.org/10.3758/BF03196322>
- Zwaan, R. A. (2014). Embodiment and language comprehension: Reframing the discussion. *Trends in Cognitive Sciences*, 18(5), 229–234. <https://doi.org/10.1016/j.tics.2014.02.008>
- Zwaan, R. A. (2014). Replications should be performed with power and precision: A response to Rommers, Meyer, and Huettig (2013). *Psychological Science*, 25(1), 305–307. <https://doi.org/10.1177/0956797613509634>
- Zwaan, R. A., & Madden, C. J. (2005). Embodied sentence comprehension. In D. Pecher & R. A. Zwaan (Eds.), *Grounding cognition: The role of perception and action in memory, language, and thinking* (pp. 224–245). Cambridge University Press.
- Zwaan, R. A., & Pecher, D. (2012). Revisiting mental simulation in language comprehension: Six replication attempts. *PLoS One*, 7, e51382. <https://doi.org/10.1371/journal.pone.0051382>
- Zwaan, R. A., & van Oostendorp, H. (1993). Do readers construct spatial representations in naturalistic story comprehension? *Discourse Processes*, 16(1–2), 125–143. <https://doi.org/10.1080/01638539309544832>
- Zwaan, R. A., Stanfield, R. A., & Yaxley, R. H. (2002). Language comprehenders mentally represent the shapes of objects. *Psychological Science*, 13, 168–171. <https://doi.org/10.1111/1467-9280.00430>
- Zwaan, R. A., Pecher, Diane, Paolacci, G., Bouwmeester, S., Verkoeijen, P., Dijkstra, K., & Zeelenberg, R. (2017). Participant nonnaïveté and the reproducibility of cognitive psychology. *Psychonomic Bulletin & Review*. <https://doi.org/10.3758/s13423-017-1348-y>

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

Springer Nature or its licensor (e.g. a society or other partner) holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.

Authors and Affiliations

Sau-Chin Chen¹  · Erin M. Buchanan²  · Zoltan Kekecs^{3,4} · Jeremy K. Miller⁵ · Anna Szabelska⁶ · Balazs Aczel³ · Pablo Bernabeu^{7,8} · Patrick Forscher^{9,10} · Attila Szuts³ · Zahir Vally¹¹ · Ali H. Al-Hoorie¹² · Mai Helmy^{13,14} · Caio Santos Alves da Silva¹⁵ · Luana Oliveira da Silva¹⁵ · Yago Luksevicius de Moraes^{15,16} · Rafael Ming Chi Santos Hsu¹⁵ · Anthonieta Looman Mafra¹⁵ · Jaroslava V. Valentova¹⁵ · Marco Antonio Correa Varella¹⁵ · Barnaby Dixon¹⁷ · Kim Peters¹⁷ · Nik Steffens¹⁷ · Omid Ghasemi¹⁸ · Andrew Roberts¹⁸ · Robert M. Ross¹⁹ · Ian D. Stephen^{18,20} · Marina Milyavskaya²¹ · Kelly Wang²¹ · Kaitlyn M. Werner²² · Dawn Liu Holford²³ · Miroslav Sirota²³ · Thomas Rhys Evans²⁴ · Dermot Lynott²⁵ · Bethany M. Lane²⁶ · Danny Riis Sahlholdt²⁶ · Glenn P. Williams²⁷ · Chrystalle B. Y. Tan²⁸ · Alicia Foo²⁹ · Steve M. J. Janssen²⁹ · Nwadiogo Chisom Arinze³⁰ · Izuchukwu Lawrence Gabriel Ndukaihe³⁰ · David Moreau³¹ · Brianna Jurosic³² · Brynna Leach³² · Savannah Lewis³³ · Peter R. Mallik³⁴ · Kathleen Schmidt³² · William J. Chopik³⁵ · Leigh Ann Vaughn³⁶ · Manyu Li³⁷  · Carmel A. Levitan³⁸ · Daniel Storage³⁹ · Carlota Batres⁴⁰ · Tyler McGee⁴⁰ · Janina Enachescu⁴¹ · Jerome Olsen⁴² · Martin Voracek⁴³ · Claus Lamm⁴³ · Ekaterina Pronizius⁴³ · Tilli Ripp⁴⁴ · Jan Philipp Röer⁴⁴ · Roxane Schnepfer⁴⁴ · Marietta Papadatou-Pastou⁴⁵ · Aviv Mokady⁴⁶ · Niv Reggev⁴⁶ · Priyanka Chandel⁴⁷ · Pratibha Kujur⁴⁷ · Babita Pande⁴⁷ · Arti Parganiha⁴⁷ · Noorshama Parveen⁴⁷ · Sraddha Pradhan⁴⁷ · Margaret Messiah Singh⁴⁷ · Max Korbmacher⁴⁸ · Jonas R. Kunst⁴⁹ · Christian K. Tamnes⁵⁰ · Frederike S. Woelfert⁵⁰ · Kristoffer Klevjer⁵¹ · Sarah E. Martiny⁵¹ · Gerit Pfuhl^{51,52} · Sylwia Adamus⁵³ · Krystian Barzykowski⁵³ · Katarzyna Filip⁵³ · Patrícia Arriaga⁵⁴ · Vasilije Gvozdenovic⁵⁵ · Vanja Kovic⁵⁵ · Fei Gao⁵⁶ · Jingxiang Li⁵⁶ · Jozef Bavoľár⁵⁷ · Monika Hricová⁵⁷ · Pavol Kačmár⁵⁷ · Matúš Adamkovič^{58,59,60} · Peter Babinčák⁶¹ · Gabriel Baník⁶² · Ivan Ropovik^{63,64} · Danilo Zambrano Ricaurte⁶⁵ · Sara Álvarez-Solas⁶⁶ · Harry Manley^{67,68} · Panita Suavansri⁶⁷ · Chun-Chia Kung⁶⁹ · Belemir Çoktok⁷⁰ · Asil Ali Özdoğru^{70,71}  · Çağlar Solak⁷² · Sinem Söylemez⁷² · Sami Çoksarı^{73,74} · İlker Dalgıç⁷⁵ · Mahmoud Elsherif⁷⁶ · Martin Vasilev⁷⁷ · Vinka Mlakic⁷⁸ · Elisabeth Oberzaucher⁷⁹ · Stefan Stieger⁷⁸ · Selina Volsa⁷⁸ · Erica D. Musser⁸⁰ · Janis Zickfeld⁸¹ · Christopher R. Chartier³²

✉ Sau-Chin Chen
csc2009@mail.tcu.edu.tw; csc2009@mail2.tcu.edu.tw

Manyu Li
manyu.li@louisiana.edu

Asil Ali Özdoğru
asil.ozdogru@marmara.edu.tr

¹ Department of Human Development and Psychology, Tzu-Chi University, No. 67, Jei-Ren St., Hualien, Taiwan

² Harrisburg University of Science and Technology, Harrisburg, PA, USA

³ Institute of Psychology, ELTE, Eotvos Lorand University, Budapest, Hungary

⁴ Department of Psychology, Lund University, Lund, Sweden

⁵ Department of Psychology, Willamette University, Salem, OR, USA

⁶ Institute of Cognition and Culture, Queen's University Belfast, Belfast, UK

⁷ Department of Psychology, Lancaster University, Lancaster, UK

⁸ Department of Education, University of Oxford, Oxford, UK

⁹ LIP/PC2S, Université Grenoble Alpes, Grenoble, France

¹⁰ Busara Center for Behavioral Economics, Nairobi, Kenya

¹¹ Department of Clinical Psychology, United Arab Emirates University, Al Ain, United Arab Emirates

¹² Jubail English Language and Preparatory Year Institute, Royal Commission for Jubail and Yanbu, Al Jubail, Saudi Arabia

¹³ Psychology Department, College of Education, Sultan Qaboos University, Muscat, Oman

¹⁴ Psychology Department, Faculty of Arts, Menoufia University, Shebin El-Kom, Egypt

¹⁵ Department of Experimental Psychology, Institute of Psychology, University of Sao Paulo, Sao Paulo, Brazil

¹⁶ Faculty of Philosophy, University Center Foundation Santo André, Santo André, Brazil

¹⁷ School of Psychology, University of Queensland, Brisbane, Australia

¹⁸ Department of Psychology, Macquarie University, Sydney, Australia

¹⁹ School of Psychological Sciences and School of Philosophy, Macquarie University, Sydney, Australia

²⁰ Department of Psychology, Bournemouth University, Bournemouth, UK

²¹ Department of Psychology, Carleton University, Ottawa, Canada

²² Department of Psychology, University of Oregon, Eugene, OR, USA

²³ Department of Psychology, University of Essex, Colchester, UK

²⁴ School of Human Sciences and Institute for Lifecourse Development, University of Greenwich, London, UK

²⁵ Department of Psychology, Maynooth University, Maynooth, Ireland

- 26 Division of Psychology, School of Social and Health Sciences, Abertay University, Dundee, UK
- 27 School of Psychology, Faculty of Health Sciences and Wellbeing, University of Sunderland, Sunderland, UK
- 28 School of Psychology and Vision Sciences, University of Leicester, Leicester, UK
- 29 School of Psychology, University of Nottingham Malaysia, Semenyih, Malaysia
- 30 Department of Psychology, Alex Ekwueme Federal University, Ndufu-Alike, Nigeria
- 31 School of Psychology, University of Auckland, Auckland, New Zealand
- 32 Department of Psychology, Ashland University, Ashland, OH, USA
- 33 Department of Psychology, University of Alabama, Tuscaloosa, AL, USA
- 34 Hubbard Decision Research, Glen Ellyn, IL, USA
- 35 Department of Psychology, Michigan State University, East Lansing, MI, USA
- 36 Department of Psychology, Ithaca College, Ithaca, NY, USA
- 37 Department of Psychology, University of Louisiana at Lafayette, Lafayette, LA, USA
- 38 Department of Cognitive Science, Occidental College, Los Angeles, CA, USA
- 39 Department of Psychology, University of Denver, Denver, CO, USA
- 40 Department of Psychology, Franklin and Marshall College, Lancaster, PA, USA
- 41 Statistics Austria, Vienna, Austria
- 42 Faculty of Psychology, University of Vienna, Wien, Austria
- 43 Department of Cognition, Emotion, and Methods in Psychology, Faculty of Psychology, University of Vienna, Vienna, Austria
- 44 Department of Psychology and Psychotherapy, Witten/Herdecke University, Witten, Germany
- 45 School of Education, National and Kapodistrian University of Athens, Athens, Greece
- 46 Department of Psychology and School of Brain Sciences and Cognition, Ben-Gurion University of the Negev, Beer Sheva, Israel
- 47 School of Studies in Life Science, Pt. Ravishankar Shukla University, Raipur, India
- 48 Department of Health and Functioning, Western Norway University of Applied Sciences, Bergen, Norway
- 49 Department of Communication and Culture, BI Norwegian Business School, Oslo, Norway
- 50 Department of Psychology, University of Oslo, Oslo, Norway
- 51 Department of Psychology, UiT - The Arctic University of Norway, Tromsø, Norway
- 52 Department of Psychology, Norwegian University of Science and Technology, Trondheim, Norway
- 53 Institute of Psychology, Faculty of Philosophy, Jagiellonian University, Krakow, Poland
- 54 Iscte-University Institute of Lisbon, CIS-IUL, Lisbon, Portugal
- 55 Laboratory for Neurocognition and Applied Cognition, Faculty of Philosophy, University of Belgrade, Belgrade, Serbia
- 56 Faculty of Arts and Humanities, University of Macau, Macau, China
- 57 Department of Psychology, Faculty of Arts, Pavol Jozef Šafarik University in Košice, Košice, Slovakia
- 58 Centre of Social and Psychological Sciences, Slovak Academy of Sciences, Bratislava, Slovakia
- 59 Faculty of Education, Charles University, Prague, Czechia
- 60 Faculty of Humanities and Social Sciences, University of Jyväskylä, Jyväskylä, Finland
- 61 Institute of Psychology, Faculty of Arts, University of Presov, Presov, Slovakia
- 62 Department of Educational Psychology and Psychology of Health, Faculty of Arts, Pavol Jozef Šafarik University, Košice, Slovakia
- 63 Institute for Research and Development of Education, Faculty of Education, Charles University, Prague, Czechia
- 64 Institute of Psychology, Czech Academy of Sciences, Prague, Czechia
- 65 Faculty of Psychology, Fundación Universitaria Konrad Lorenz, Bogotá, Colombia
- 66 Faculty of Health, International University of La Rioja (UNIR), Logroño, Spain
- 67 Faculty of Psychology, Chulalongkorn University, Bangkok, Thailand
- 68 Faculty of Behavioral Sciences, Education, & Languages, HELP University, Subang 2, Kuala Lumpur, Malaysia
- 69 Department of Psychology, National Cheng Kung University, Tainan, Taiwan

- ⁷⁰ Department of Psychology, Üsküdar University, İstanbul, Turkey
- ⁷¹ Department of Psychology, Marmara University, İstanbul, Turkey
- ⁷² Department of Psychology, Manisa Celal Bayar University, Manisa, Turkey
- ⁷³ Department of Psychology, Erzurum Technical University, Erzurum, Turkey
- ⁷⁴ Network for Economic and Social Trends, Western University, London, Canada
- ⁷⁵ Department of Psychology, Ankara Medipol University, Ankara, Turkey
- ⁷⁶ Department of Psychology and Vision Sciences, University of Leicester, Leicester, UK
- ⁷⁷ Department of Experimental Psychology, University College London, London, UK
- ⁷⁸ Department of Psychology and Psychodynamics, Karl Landsteiner University of Health Sciences, Krems an der Donau, Austria
- ⁷⁹ Department of Evolutionary Anthropology, University of Vienna, Wien, Austria
- ⁸⁰ Department of Psychology, Barnard College, New York, NY, USA
- ⁸¹ Department of Management, Aarhus University, Aarhus, Denmark