

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Investigating the Composite Effect in Prototype-Defined Checkerboards vs. Faces

Permalink

<https://escholarship.org/uc/item/0nx8p75p>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 44(44)

Authors

Waguri, Emika

McLaren, IPL

Civile, Ciro

Publication Date

2022

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Investigating the Composite Effect in Prototype-Defined Checkerboards vs. Faces

Emika Waguri (ew518@exeter.ac.uk)

School of Psychology, College of Life and Environmental Sciences,
University of Exeter, UK

R. McLaren (r.p.mclaren@exeter.ac.uk)

School of Psychology, College of Life and Environmental Sciences,
University of Exeter, UK

I.P.L. McLaren (i.p.l.mclaren@exeter.ac.uk)

School of Psychology, College of Life and Environmental Sciences,
University of Exeter, UK

Ciro Civile (c.civile@exeter.ac.uk)

School of Psychology, College of Life and Environmental Sciences,
University of Exeter, UK

Abstract

The study reported here examined the role of expertise in the composite face effect which constitutes better recognition of the top half of a face when in composite with a *congruent* vs. an *incongruent* (in terms of response required) bottom half. **Experiment 1a (n=96)** used prototype-defined artificial stimuli (checkerboards) to investigate the composite effect. The advantage of using these stimuli is that expertise can be fully controlled. **Experiment 1b (n=96)** aimed to replicate the composite effect in face stimuli which served as a control and provided a direct comparison of the composite effect between face and checkerboard stimuli. A full experimental design including congruent/incongruent aligned and misaligned composites was used in both experiments to measure the composite effect. Experiment 1a revealed that the composite effect could not be obtained in checkerboard composites. Experiment 1b confirmed the robust composite face effect. We interpret our results as suggesting that expertise/perceptual learning does not contribute to the composite effect for faces.

Keywords: Composite Face Effect, Congruency Effect, Perceptual Learning, Face recognition

Introduction

The debate over whether face recognition mechanisms are “special” or due to “expertise” was amplified by the discovery of a phenomenon named the face inversion effect (Yin, 1969; Civile et al., 2011; Civile et al., 2014; Civile et al., 2016; McCourt et al., 2021). This refers to impaired performance at recognizing upside-down faces, as opposed to their usual upright orientation. It was initially interpreted as a marker of “face-specificity”, as the effect was found to be larger in faces than objects (Yovel & Kanwisher, 2005). However, Diamond and Carey (1986)’s finding of an inversion effect for dog images by dog breeders/experts, followed by Gauthier and Tarr’s (1997) finding of an inversion effect for mono-orientated artificial objects named Greebles after training, and McLaren’s (1997) finding of an inversion effect for

non mono-orientated artificial stimuli after pre-exposure, showed how “expertise” contributes to the inversion effect. Importantly, first McLaren and Civile (2011) and then Civile, Zhao et al., (2014) extended McLaren’s (1997) work, using prototype-defined checkerboards in the old/new recognition task typically employed to study the inversion effect. Participants were first engaged in a categorization task (the pre-exposure phase) where they sorted a set of checkerboards created from two prototype-defined categories. Participants were then asked to memorize new checkerboards drawn from one of the two familiar categories seen, and from a novel category not seen previously. Half of the checkerboards were presented in the orientation familiarized during the categorization task (i.e., upright/inverted by rotating stimuli 180°). The checkerboards are non-mono-orientated (no predefined orientation), so participants had no sense of an upright or inverted orientation for checkerboards drawn from the novel category which served as a baseline. In the final recognition task, “old” exemplars seen in the study phase were intermixed with “new” ones split by the same conditions: familiar upright and inverted, novel upright and inverted. Participants indicated if they had or had not seen the exemplars previously in the study phase. The results showed a robust inversion effect for checkerboards in the familiar category compared to novel, mostly due to increased performance in upright, familiar category checkerboards.

Evidence that the inversion effect for checkerboards and faces share at least some of the same causal mechanisms is supported by a recent line of research using a particular transcranial Direct Current Stimulation (tDCS) procedure. Civile, Verbruggen et al (2016) showed that anodal tDCS (for 10 mins at 1.5mA) delivered on the left dorsolateral prefrontal cortex (Fp3)

during the same old/new recognition task used by Civile, Zhao et al (2014) abolished the checkerboard inversion effect. This was due to reduced recognition of upright familiar checkerboards compared to sham. The same tDCS procedure applied to faces resulted in a significant reduction, but not elimination, of the inversion effect (compared to sham), also due to impaired recognition for upright faces (Civile, McLaren et al., 2018). This result has been replicated across multiple studies and it is now an established finding (Civile, Obhi et al., 2019; Civile, Cooke et al., 2020; Civile, McLaren et al., 2020; Civile, Waguri et al., 2020; Civile, Quaglia et al., 2021; Civile, McLaren et al., 2021).

These findings were explained in terms of the McLaren, Kaye and Mackintosh (MKM) (McLaren et al., 1989; McLaren & Mackintosh; McLaren et al., 2012; McLaren et al., 2016) theory of perceptual learning which applies only to upright stimuli because we have little or no experience in seeing inverted stimuli and so performance on these is not aided by any significant amount of perceptual learning. On this account, when tDCS is applied, the reduced inversion effect is due to impaired recognition for upright stimuli because of the disruption of perceptual learning. The tDCS affects the ability to discriminate between and recognise different upright stimuli, making them look more “similar”.

Civile, Quaglia et al (2021) extended the tDCS-induced effects on the inversion effect using a matching task paradigm to ensure a comparable level of performance between checkerboard and face stimuli. The results confirmed that tDCS stimulation is able to fully reduce the inversion effect for checkerboard exemplars compared to sham. Hence, the same procedure significantly reduced the face inversion effect compared to sham. Critically, when compared directly, the reduced face inversion effect was found to be significantly larger than any residual checkerboard inversion effect. Thus, the authors proposed an explanation based on which face recognition mechanisms would be partly based on expertise manifesting through perceptual learning. The tDCS procedure would eliminate this perceptual learning component causing a full reduction of the checkerboard inversion effect and a partial reduction of the face inversion effect. However, the remaining face inversion effect would be due to a specificity component that would not be affected by the tDCS procedure.

Civile, McLaren et al (2021) made a first step toward the investigation of the potential specificity component involved in face recognition. Specifically, by applying the same tDCS procedure to another robust phenomenon named the *composite face effect*, the authors hoped to reveal to what extent holistic processing may be the mechanism specific to face stimuli. When using the complete/full experimental design the composite effect refers to people having lower accuracy at identifying the top half of one face presented in composite with the bottom half of another face that supports the alternative response when upright and aligned, than when the two halves are laterally misaligned. Importantly, while the inversion effect has been used as an index of disrupted configural processing, the composite effect has been used

to index holistic processing which is the reliance on the small spatial differences between the facial components in the context of the whole face (for a review see Maurer et al., 2002). Hence, upright composite faces are perceived as a “new” face because the internal features are strongly integrated and becomes difficult to isolate them (Murphy et al., 2017) and this makes it difficult to ignore the bottom half of the face. Civile, McLaren et al (2021) showed that the tDCS procedure used to affect the inversion effect did not influence the composite face effect despite reducing overall performance (composite stimuli are still upright faces) compared to sham. Thus, the authors suggested that the holistic processing indexed by the composite face effect could be face specific and not a result of expertise.

In the current study we aimed to investigate directly whether the composite effect would be found for prototype-defined artificial stimuli (checkerboards) that participants had never seen before entering the lab and for which expertise can be fully controlled. Showing that a composite effect cannot be obtained for checkerboards but can be obtained for faces, would provide additional evidence in support of holistic processing being a face specific mechanism and perhaps the factor unaffected by the tDCS procedure which does impact the face inversion effect (Civile, Quaglia et al., 2021).

The only two studies reporting a composite effect for lab trained, non-face artificial stimuli were by Gauthier & Tarr (2002) using Greebles, and Wong et al., (2009) using Ziggerins. Both studies used a full design. However, initially Gauthier et al (1998) did not find a composite effect for Greebles when using a partial/original design. Studies have highlighted that composite effects calculated using the complete and partial/original designs would not seem to always correlate (Richler & Gauthier, 2014). A key difference between the complete and partial/original design is that the complete design would include congruent and incongruent composites presented aligned and misaligned. In the full design when extracting the composite effect, a component named the *congruency effect*, higher performance for congruent face halves vs incongruent composites (in terms of the response required by each face half) is crucial as the former is calculated by subtracting the reduced congruency effect in response to misaligned composites from the robust congruency effect typically obtained in aligned composites (Civile, McLaren et al., 2021).

Recently, one study demonstrated a significant congruency effect for prototype-defined categories of checkerboards (Waguri et al., 2021). However, the design adopted did not include misaligned composites but only aligned ones, and so did not investigate the composite effect as in the full design used in the literature (Gauthier & Tarr, 2002; Wong et al., 2009; Civile, McLaren et al., 2021). The current study aims to extend these findings by applying the full composite effect design (including misaligned trials) in **Experiment 1a**. **Experiment 1b** aimed to find a composite effect for faces allowing a comparison between the two types of stimuli. This is the first study looking at the composite effect for checkerboard stimuli.

Method

Participants

Experiment 1a. 96 naïve participants (mean age = 25.39, age range = 18-40) were recruited via Prolific. They had an approval rating of at least 90% from participation in other studies and received monetary compensation adhering to the fair pay policies of Prolific Academic. The sample size was in part determined by the counterbalance of the stimuli and previous studies using a similar experimental paradigm (Civile, Quaglia et al., 2021; Civile, McLaren et al., 2021; Waguri et al., 2021).

Experiment 1b. 96 naïve participants (mean age = 23.8, age range = 18-38) were recruited via Prolific with the same inclusion criterion and compensation as in Experiment 1a.

Materials

Experiment 1a used the same 4 prototype-defined categories of checkerboards (A, B, C, D) from Civile, Zhao et al (2014, Experiment 1a). Category prototypes (16 x 16) were randomly generated with the constraint that they shared 50% of their squares with each of the other prototypes (50% black squares and 50% white types). Exemplars were generated from these prototypes by randomly changing forty-eight squares thus, on average, 24 squares would be expected to alter from black to white/white to black. Composite checkerboards were presented at the resolution of 256 x 256 pixels on a grey background. The composites consisted of top and bottom halves of different checkerboards (each containing 16 x 16 squares) drawn from the same prototype-defined category (e.g., A65 Top, A73 Bottom). 64 of the composites were aligned, while the other 64 were modified into misaligned checkerboards by shifting the top half to the left (total of 128 composite checkerboards).

Experiment 1b used a total of 256 face images (174 x 225 pixels), all standardized to greyscale on a black background. The original images were selected from the Psychological Image Collection at Stirling open database, (<https://pics.stir.ac.uk>). All the images were cropped to a standardized oval shape, removing distracting features such as the hairline, and adjusted to standardize image luminance (Civile, Quaglia et al., 2021). The set of faces was then used to create the composite faces as in Civile, McLaren et al (2021). Both experiments were programmed and run on the online platform *Gorilla*.

The Behavioral Task

Experiment 1a. Categorization phase commenced after participants provided consent and were first shown instructions. Participants were shown exemplar checkerboards from categories A&C or B&D depending on the counterbalance group they were assigned to (64 from each category; 128 in total). Each exemplar was shown one at a time at random order. They were instructed to sort these exemplars into two categories (A-C or B-D) through trial-and-error, by pressing one of the two indicated keys on the keyboard (counterbalanced). They were given immediate feedback on whether their response was correct or incorrect. If they did not respond

within 4 seconds, they were timed out. A fixation cross preceded each stimulus in the center of the screen (1 s).

Training phase. The purpose of this task was for the participants to associate the response keys “x” and “.” with words SAME and DIFFERENT. They were instructed to press “x” or “.” as quickly as possible when classifying them as SAME or DIFFERENT (counterbalanced). 48 trials (24 SAME, 24 DIFFERENT) were presented randomly one at a time for <1 s after a fixation cross (1s). They received feedback on each response as correct or incorrect.

Composite Checkerboard Matching-task. This phase involved a matching-task following the full design procedure to measure the composite effect used in Civile, McLaren et al (2021) however this time with composite checkerboards (128 trials) instead of composite faces. Overall, participants saw 32 trials of “same” aligned, 32 “different” aligned, 32 “same” misaligned and 32 “different” misaligned composites split by the 8 stimulus conditions (each 16 aligned, 16 misaligned trials): familiar and novel congruent aligned/misaligned, familiar and novel incongruent aligned/misaligned. Each trial commenced with a fixation cross (1s), followed by a TARGET composite checkerboard stimulus (1s), an interstimulus interval (1.5s), and a TEST composite checkerboard stimulus (≤ 2 s). Participants were to press either the ‘x’ key or ‘.’ key (same as previous training phase) when identifying the top halves of the TARGET and TEST stimulus as same or different. In-line with Civile, McLaren et al (2021) and Waguri et al (2021), congruent and incongruent trials were presented in a counterbalanced fashion across participants with aligned and misaligned stimuli randomly intermixed.

In the *congruent familiar trials*, participants first saw a TARGET composite checkerboard created by selecting the top and bottom halves of two different new (not seen in the categorization task) exemplars selected from familiar categories (A-C or B-D) as seen in the categorization phase (e.g., top-half of exemplar A65 and bottom-half of A73 or top-half of exemplar C65 and bottom-half of C73). In the TEST trial, they would either see the “same” or “different” composite, where in the latter case the top and bottom halves are different exemplars from the same categories (e.g., top-half of A89 and bottom-half of A81 or top-half of exemplar C89 and bottom-half of C81). Overall, 32 A or B and 32 C or D composites were presented (16 same, 16 different) randomly. An A-TARGET composite would correspond to an A-TEST composite, and a C-TARGET composite would correspond to a C-TEST composite. The same applied to B- and D- TARGET/TEST. The *congruent novel trials* TARGET and TEST “same” or “different” composites were created by selecting the top and bottom halves of exemplars drawn from the other categories (32 each, 16 same and 16 different) not seen during the categorization task. The novel composites were also created from exemplars drawn from the same novel category, and the TARGET/TEST would correspond to the same category. For *incongruent familiar and novel trials*, the TARGET/TEST would be considered ‘same’ if the top halves of the composites were the same, but both

would have different bottom halves (e.g., TARGET: A65/81; TEST: A65/A73). The converse was true for ‘different’, wherein the top halves of the TARGET and TEST are different, but have the same bottom halves (e.g., TARGET: A89/A73; TEST: A65/A73).

In congruent and incongruent misaligned trials, the top and bottom halves of each composite were shifted horizontally relative to one another so that they overlapped across half their length (Figure 1a).

Experiment 1b. With the aim of matching Experiment 1a, in Experiment 1b participants were presented first with a **categorization phase** during which they were asked to sort a set of regular faces presented one at a time in random order. Hence, they pressed one of the two indicated keys (counterbalanced) if the image presented was a male face or the other key if it was a female face. If they did not respond within 4 seconds, they were timed out. They were given immediate feedback on whether their response was correct or incorrect. A fixation cross preceded each stimulus in the center of the screen (1 s). In total 128 (64 male and 64 female) faces were presented.

Training phase. The same as for Experiment 1a.

Composite Face Matching-task. This followed the full design procedure used in Civile, McLaren et al (2021). Each trial began with a fixation cue presented in the center of the screen (1s), followed by a TARGET face stimulus (1s), an interstimulus interval (1.5s), and a TEST face stimulus (≤ 2 s). Participants pressed either ‘x’ key or ‘.’ key to identify the top half of the test face as “same” or “different” to the top half of the target face. All the composite faces were presented upright and were split by four conditions (Congruent Aligned, Incongruent Aligned, Congruent Misaligned and Incongruent Misaligned). No ‘novel’ condition was used in this case hence an overall of 64 trials were presented. Congruent and incongruent trials were presented in a counterbalanced fashion across participants with aligned and misaligned stimuli randomly intermixed. In the *congruent aligned* trials, participants first saw a TARGET face composite created by selecting the top and bottom halves of two different faces (e.g., A-B, where A is the top half and B the bottom half). In the TEST face trial, they would either see the same TARGET face or a new face composite created by selecting the top and bottom halves of two different faces (e.g., C-D). The *Incongruent aligned* trials were presented either with the same top halves as for the TARGET faces but with different bottom halves (A-D), or with different top halves from the TARGET faces but the same bottom halves (C-B). In the congruent and incongruent misaligned trials the top and bottom halves of each composite were shifted horizontally to overlap across half their length (Figure 1b).

Results

Accuracy from the matching task in both experiments is the primary measure which has been used to extract a *d-prime* (d') sensitivity measure (Stanislaw & Todorov, 1999) for ‘same’ and ‘different’ trials. To calculate d' , we used participants’ hit rate (H), the proportion of SAME trials to which the participant responded SAME, and false alarm rate (F), the proportion of DIFFERENT trials to

which the participant responded SAME. However, d' is not simply $H - F$; rather it is the difference between z transforms of the two rates: $d' = z(H) - z(F)$. A d' of 0 indicates chance-level performance. We assessed performance against chance to show that the stimulus’ conditions were recognized significantly above chance (for all conditions in both experiments we found $p < .001$). Reaction time data was inspected to confirm that no effect of speed-accuracy trade-off was found.

Experiment 1a. We conducted a 3-way ANOVA using within-subjects factors *Congruency* (congruent or incongruent), *Familiarity* (familiar or novel), and *Alignment* (aligned or misaligned). Analysis of Variance (ANOVA) showed a significant main effect of *Congruency* $F(1, 95) = 10.08, p = .002, \eta^2_p = .09$, but no significant main effect of *Familiarity* $F(1, 95) = .026, p = .87, \eta^2_p < .01$, nor of *Alignment* $F(1, 95) = .06, p = .806, \eta^2_p < .01$. Critically, the interaction between *Congruency* $\times$ *Alignment* was not significant, $F(1, 95) = 2.07, p = .153, \eta^2_p = .02$, giving no evidence of a composite effect for checkerboard exemplars though we note that the numerical effect is in the correct direction for that effect. None of the other interactions revealed a significant effect including the overall three-way interaction (*Congruency* $\times$ *Familiarity* $\times$ *Alignment*), $F(1, 95) = .24, p = .620, \eta^2_p < .01$. We conducted two additional paired sample t-tests to further investigate the *congruency effect* (i.e., better performance for congruent vs incongruent trials) for aligned and misaligned composites. A significant congruency effect was found in aligned trials with congruent composites ($M = 1.83, SE = .08$) being better identified than incongruent ones ($M = 1.46, SE = .12$), $t(95) = 3.22, p < .001, \eta^2_p = .09$. There was also a significant congruency effect on misaligned trials with congruent composites ($M = 1.77, SE = .08$) being better identified than incongruent ones ($M = 1.53, SE = .10$), $t(95) = 2.40, p = .018, \eta^2_p = .06$ (Figure 2a).

Experiment 1b. We conducted a 2 $\times$ 2 ANOVA using the within-subjects factors *Congruency* (congruent or incongruent) and *Alignment* (aligned or misaligned) which revealed no significant main effect of *Congruency* $F(1, 95) = .51, p = .475, \eta^2_p < .01$, nor of *Alignment* $F(1, 95) = .92, p = .338, \eta^2_p = .01$. Importantly, and in agreement with previous studies that adopted the same full design (the same design was adopted by Civile, McLaren et al 2021) the interaction *Congruency* $\times$ *Alignment* was significant, $F(1, 95) = 14.62, p < .001, \eta^2_p = .13$, indicating that there was a robust composite face effect. We conducted two additional paired sample t-tests that revealed a significant congruency effect in aligned trials with congruent composites ($M = 2.27, SE = .13$) being better identified than incongruent ones ($M = 2.02, SE = .11$), $t(95) = 2.84, p = .005, \eta^2_p = .08$. This congruency effect was actually reversed for misaligned trials, with congruent composites ($M = 2.03, SE = .12$) being numerically (but not significantly) worse identified than incongruent ones ($M = 2.16, SE = .15$), $t(95) = 1.44, p = .152, \eta^2_p = .02$ (Figure 2b).

Analysis Between the Experiments. We conducted an additional analysis with the aim of directly testing if the composite face effect found in Experiment 1b was

significantly larger than the non-significant composite effect in Experiment 1a. We extracted a composite effect index for each experiment by subtracting the congruency effect found in misaligned trials from that found in aligned trials. Following this, we conducted an independent t-test, which revealed a significant difference, $t(190) = 2.08$, $p = .038$, $\eta^2_p = .04$ indicating a larger composite effect in Experiment 1b ($M = .40$, $SE = .10$) vs that found in Experiment 1a ($M = .13$, $SE = .09$).

Discussion

The study reported here aimed to investigate whether a composite effect could be obtained in non mono-orientated, artificial stimuli (checkerboards) and to compare that with the robust composite effect typically found for faces. According to the perceptual learning literature there are two main factors that may contribute to face recognition skills. One is expertise manifesting through perceptual learning and previous studies evidenced how a specific tDCS procedure can abolish that by reducing entirely the checkerboard inversion effect (Civile, Verbruggen et al., 2016; Civile, Quaglia et al., 2021). The other factor is based on ‘face-specificity’ and previous work has provided some evidence that the same tDCS procedure cannot modulate the composite face effect which is often indicated as a robust index of holistic processing of faces (Civile, McLaren et al., 2021).

We conducted two experiments aiming to advance our understanding of the role that holistic processing may have specifically in face recognition. We directly tested this by using the same checkerboard exemplars employed in the perceptual learning literature to demonstrate the inversion effect (Civile, Zhao et al., 2014, but using a composite effect paradigm to index holistic processing. The main finding from Experiment 1a is that we do not find a significant composite effect for checkerboard composites. In line with Waguri et al (2021) we found a congruency effect for aligned checkerboard composites but did not find a significant reduction of the congruency effect in the misaligned composites. In both aligned and misaligned trials we found a significant congruency effect. This is the first evidence of this type based on non mono-orientated artificial stimuli such as checkerboards. Importantly, Experiment 1b demonstrated a robust composite face effect as typically obtained in the literature (e.g., Civile, McLaren et al., 2021). Here a significant congruency effect was found for aligned face composites which was numerically reversed when the composites were misaligned. Through an additional analysis between the composite effect index from Experiment 1a vs Experiment 1b we found that the composite face effect was significantly larger than that for checkerboard stimuli.

Overall, these results are in line with Civile, McLaren et al., (2021) which demonstrated how the same tDCS procedure able to remove the perceptual learning component of the face inversion effect is not able to influence the composite face effect which is based on holistic processing. We have now provided some evidence that the composite effect can be obtained for faces but not for non mono-orientated checkerboards

even after participants had been trained with them. All together the results from Civile, McLaren et al (2021) and ours provide support to the proposition that the composite effect is specific to faces. Furthermore, if we assume that the composite effect is a robust index of holistic processing (for a review see Maurer et al., 2022), we will then suggest that this is the type of processing that is face-specific and thus not influenced by the tDCS procedure. Coming back to the findings by Civile, Quaglia et al (2021) one may say that the remaining inversion effect for faces would be due to holistic information not being affected (or at least in part) by the tDCS. Future work should investigate this directly perhaps by looking at the tDCS-induced effects on the inversion effect for sets of face stimuli not containing the usual holistic information (e.g., a novel face outline).

One could argue that our results contradict the results of Gauthier and Tarr (2002), and Wong et al., (2009). Both these studies found a composite effect for artificial stimuli that participants had never seen before entering the lab and being trained with them. However, there are two main differences with our study. The first one regards the stimuli used and the fact that checkerboards are non mono-orientated. In this sense, these stimuli have no featural similarities with faces. Despite Gauthier and Tarr (2002), and Wong et al., (2009) demonstrating how familiarity is important in obtaining the composite effect, both Greebles and at least some categories of Ziggerins present a configuration of features that could resemble those of upright faces. Thus, this could elicit the type of holistic processing typically found for face stimuli and lead to a composite effect. The second difference regards the training phase/categorization task. A clear difference in the training task between our Experiment 1a (checkerboards) and those of Greebles/Ziggerins (Gauthier & Tarr, 2002; Wong et al., 2009), is utilizing a categorization task contrasted with individuation training. Both train participants to become experts, but this is nuanced in the sense that individuation particularly emphasizes subordinate level training as opposed to basic-level processing by categorization. While there is much debate as to what exactly it promotes and whether subordinate level training can indeed increase holistic processing, Wong et al., (2009) have demonstrated that individuation training (i.e., learning and identifying individual Ziggerins) does yield a composite effect in artificial stimuli as opposed to categorization training (class level expertise). This would indicate that there is a top-down effect akin to personification affecting the manifestation of the composite effect. Therefore, there may be an additional component other than lower-level perceptual processes (i.e., holistic) that influences face processing, which has also been suggested by Civile, McLaren et al (2021). Our results advance our understanding of the mechanisms that are the basis of face recognition skills. We have shown that expertise manifesting through perceptual learning in the case of checkerboards does not lead to a composite effect (Experiment 1a), a robust phenomenon found in the face recognition literature and confirmed in Experiment 1b.

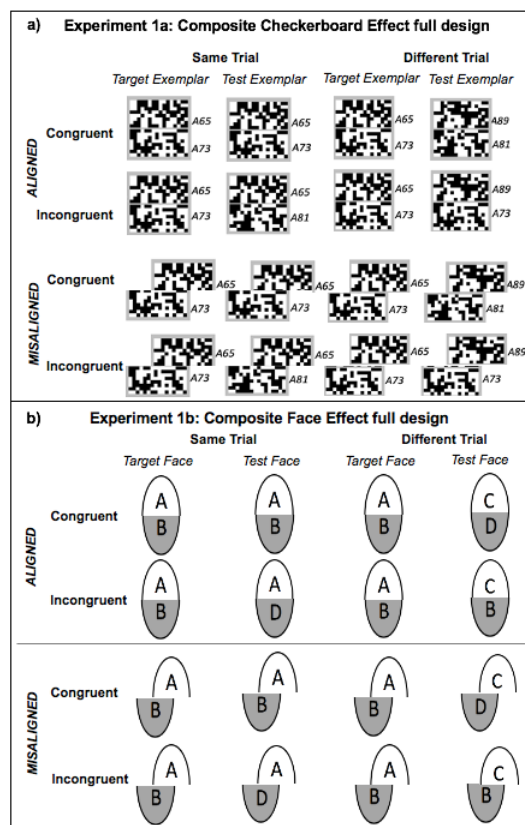


Figure 1. Panel a illustrates the full design for Experiment 1a with aligned and misaligned checkerboard composites. The same design was used for ‘familiar’ and ‘novel’ category exemplars. Panel b illustrates the full design for Experiment 1b which followed the same logic, except with composite faces instead of checkerboards.

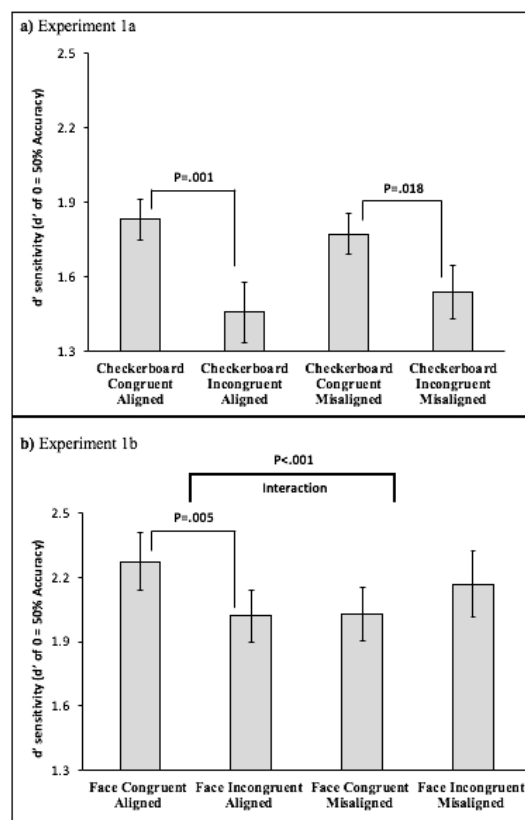


Figure 2. Panel a and Panel b reports the results from Experiment 1a and 1b respectively. In both panels, the x-axis shows the stimulus conditions, the y-axis shows d' . Error bars represent s.e.m.

Acknowledgements

This project has received funding from the Economic and Social Research Council *New Investigator Grant* (Ref.ES/R005532) awarded to Ciro Civile (PI) and I.P.L. McLaren (Co-I) and from the UKRI Covid-19 Grant Allocation awarded to Ciro Civile.

References

- Civile, C., McLaren, R., and McLaren, I.P.L. (2011). Perceptual learning and face recognition: Disruption of second-order relational information reduces the face inversion effect. In L. Carlson, C. Hoelscher, & T.F. Shipley (Eds.), *Proceedings of the 33rd Annual Conference of the Cognitive Science Society*, (pp. 2083-88). Austin, TX: Cognitive Science Society.
- Civile, C., McLaren, R., & McLaren, I.P.L. (2014). The face Inversion Effect-Parts and wholes: Individual features and their configurations. *Quarterly Journal of Experimental Psychology*, 67, 728-746.
- Civile, C., Zhao, D., Ku, Y., Elchlepp, H., Lavric, A., & McLaren, I. (2014). Perceptual learning and inversion effect: Recognition of prototype-defined familiar checkerboards. *Journal of Experimental Psychology: Animal Learning and Cognition*, 40, 144-161.
- Civile, C., McLaren, R., and McLaren, I.P.L. (2016). The face inversion effect: Roles of first and second-order relational information. *The American Journal of Psychology*, 129, 23-35.
- Civile, C., Verbruggen, F., McLaren, R., Zhao, D., Ku, Y., & McLaren, I.P.L. (2016). Switching off perceptual learning: Anodal transcranial direct current stimulation (tDCS) at Fp3 eliminates perceptual learning in humans. *Journal of Experimental Psychology: Animal Learning and Cognition*, 290-296.
- Civile, C., McLaren, R., & McLaren, I.P.L. (2018). How we can change your mind: Anodal tDCS to Fp3 alters human stimulus representation and learning. *Neuropsychologia*, 119, 241-246.
- Civile, C., Obhi, S.S., and McLaren, I.P.L. (2019). The role of experience-based perceptual learning in the Face Inversion Effect. *Vision Research*, 157, 84-88.
- Civile, C., Cooke, A., Liu, X., McLaren, R., Elchlepp, H., Lavric, A., Milton, F., and I.P.L. McLaren. (2020). The effect of tDCS on recognition depends on stimulus generalization: Neuro-stimulation can predictably enhance or reduce the face inversion effect. *Journal of Experimental Psychology: Animal Learning and Cognition*, 46, 83-98.
- Civile, C., McLaren, R., Waguri, E., and McLaren, I.P.L. (2020). Testing the immediate effects of transcranial Direct Current Stimulation (tDCS) on face recognition skills. In S. Denison, M. Mack, Y. Xu, & B.C. Armstrong (Eds.), *Proceedings of the 42st Annual Conference of the Cognitive Science Society*. Toronto, ON: Cognitive Science Society, 1141-47.

- Civile, C., Waguri, E., Quaglia, S., Wooster, B., Curtis, A., McLaren, R., Lavric, A., and McLaren, I.P.L. (2020). Testing the effects of transcranial Direct Current Stimulation (tDCS) on the Face Inversion Effect and the N170 Event-Related Potentials (ERPs) component. *Neuropsychologia*, 143, 107470.
- Civile, C., McLaren, R., Milton, F., and McLaren, I.P.L. (2021). The Effects of transcranial Direct Current Stimulation on Perceptual Learning for Upright Faces and its Role in the Composite Face Effect. *Journal of Experimental Psychology: Animal Learning and Cognition*.
- Civile, C., Quaglia, S., Waguri, E., Ward, W., McLaren, R., and McLaren, I.P.L. (2021). Using transcranial Direct Current Stimulation (tDCS) to investigate why Faces are and are Not Special. *Scientific Reports*, 11, 4380.
- Diamond, R. & Carey, S. (1986). Why faces are and are not special: An effect of expertise. *Journal of Experimental Psychology: General*, 115, 107-117.
- Gauthier, I., & Tarr, M. (1997). Becoming a “Greeble” expert: exploring mechanisms for face recognition. *Vision Research*, 37, 1673-1682.
- Gauthier, I., & Tarr, M. (2002). Unraveling mechanisms for expert object recognition: Bridging brain activity and behavior. *Journal of Experimental Psychology: Human Perception and Performance*, 28, 431-446.
- McCourt, S., McLaren, I.P.L., and Civile, C. (2021). Perceptual processes of face recognition: Single feature orientation and holistic information contribute to the face inversion effect. In T. Fitch, C. Lamm, H. Leder, & K. Teßmar-Raible (Eds.), *Proceedings of the 43rd Annual Conference of the Cognitive Science Society* (pp. 728-34). Vienna, Austria: Cognitive Science Society.
- McLaren, I.P.L., Kaye, H. & Mackintosh, N.J. (1989). An associative theory of the representation of stimuli: Applications to perceptual learning and latent inhibition. In R.G.M. Morris (Ed.) *Parallel Distributed Processing - Implications for Psychology and Neurobiology*. Oxford, Oxford University Press.
- McLaren (1997). Categorization and perceptual learning: An analogue of the face inversion effect. *The Quarterly Journal of Experimental Psychology* 50, 257-273.
- McLaren, I.P.L. & Mackintosh, N. (2000). An elemental model of associative learning: Latent inhibition and perceptual learning. *Animal Learning and Behavior*, 38, 211-246.
- McLaren, I.P.L., and Civile, C. (2011). Perceptual learning for a familiar category under inversion: An analogue of face inversion? In L. Carlson, C. Hoelscher, & T.F. Shipley (Eds.), *Proceedings of the 33rd Annual Conference of the Cognitive Science Society*, (pp. 3320-25). Austin, TX: Cognitive Science Society.
- McLaren, I.P.L., Forrest, C., & McLaren, R. (2012). Elemental representation and configural mappings: combining elemental and configural theories of associative learning. *Learning and Behavior*, 40, 320-333.
- McLaren, I.P.L., Carpenter, K., Civile, C., McLaren, R., Zhao, D., Ku, Y., Milton, F., and Verbruggen, F. (2016). Categorisation and Perceptual Learning: Why tDCS to Left DLPC enhances generalisation. *Associative Learning and Cognition*. Homage to Prof. N.J. Mackintosh. Trobalon, J.B., and Chamizo, V.D. (Eds.), University of Barcelona
- Maurer, D., Le Grand, R., and Mondloch, C. (2002). The many faces of configural processing. *Trends in Cognitive Science*, 6, 255-260.
- Murphy, J., Gray, K. L. H., and Cook, R. (2017). The composite face illusion. *Psychonomic Bulletin & Review*, 24, 245-261.
- Richler, J. J., & Gauthier, I. (2014). A meta-analysis and review of holistic face processing. *Psychological Bulletin*, 140, 1281-1302.
- Stanislaw H, & Todorov N. (1999). Calculation of signal detection theory measures. *Behavior Research Methods Instruments & Computers* 31, 137-149.
- Waguri, E., McLaren, R., McLaren, I.P.L., and Civile, C. (2021). Using prototype-defined checkerboards to investigate the mechanisms contributing to the Composite Face Effect. In T. Fitch, C. Lamm, H. Leder, & K. Teßmar-Raible (Eds.), *Proceedings of the 43rd Annual Conference of the Cognitive Science Society* (pp. 1236-42). Vienna, Austria: Cog. Sci. Soc.
- Wong, A. C., Palmeri, T. J., & Gauthier, I. (2009). Conditions for facelike expertise with objects: becoming a Ziggerin expert—but which type? *Psychological Science*, 20, 1108-1117.
- Yin, R. K. (1969). Looking at upside-down faces. *Journal of Experimental Psychology*, 81, 141-145.
- Yovel G., & Kanwisher N. (2005). The neural basis of the behavioral face-inversion effect *Current Biology*, 15, 2256- 62.