

Central Lancashire Online Knowledge (CLoK)

Title	The importance of detailed context reinstatement for the production of identifiable composite faces from memory
Type	Article
URL	https://clock.uclan.ac.uk/id/eprint/36787/
DOI	https://doi.org/10.1080/13506285.2021.1890292
Date	2021
Citation	Fodarella, Cristina, Marsh, John Everett, Chu, Simon, Athwal Kooner, Palwinder, Jones, Helen orcid iconORCID: 0000-0002-2716-051X, Skelton, Faye C., Wood, Ellena Louise, Jackson, Elizabeth and Frowd, Charlie (2021) The importance of detailed context reinstatement for the production of identifiable composite faces from memory. Visual Cognition, 29 (3). pp. 180-200. ISSN 1350-6285
Creators	Fodarella, Cristina, Marsh, John Everett, Chu, Simon, Athwal Kooner, Palwinder, Jones, Helen, Skelton, Faye C., Wood, Ellena Louise, Jackson, Elizabeth and Frowd, Charlie

It is advisable to refer to the publisher's version if you intend to cite from the work.
<https://doi.org/10.1080/13506285.2021.1890292>

For information about Research at UCLan please go to <http://www.uclan.ac.uk/research/>

All outputs in CLoK are protected by Intellectual Property Rights law, including Copyright law. Copyright, IPR and Moral Rights for the works on this site are retained by the individual authors and/or other copyright owners. Terms and conditions for use of this material are defined in the <http://clock.uclan.ac.uk/policies/>

The importance of detailed context reinstatement for the production of identifiable composite faces from memory

Cristina Fodarella, John E. Marsh, Simon Chu, Palwinder Athwal-Kooner, Helen S. Jones, Faye C. Skelton, Ellena Wood, Elizabeth Jackson & Charlie D. Frowd

To cite this article: Cristina Fodarella, John E. Marsh, Simon Chu, Palwinder Athwal-Kooner, Helen S. Jones, Faye C. Skelton, Ellena Wood, Elizabeth Jackson & Charlie D. Frowd (2021): The importance of detailed context reinstatement for the production of identifiable composite faces from memory, *Visual Cognition*, DOI: [10.1080/13506285.2021.1890292](https://doi.org/10.1080/13506285.2021.1890292)

To link to this article: <https://doi.org/10.1080/13506285.2021.1890292>



© 2021 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group



Published online: 02 Mar 2021.



Submit your article to this journal [↗](#)



Article views: 59



View related articles [↗](#)



View Crossmark data [↗](#)

The importance of detailed context reinstatement for the production of identifiable composite faces from memory

Cristina Fodarella ^a, John E. Marsh^{a,b}, Simon Chu^a, Palwinder Athwal-Kooner^c, Helen S. Jones^a, Faye C. Skelton^d, Ellena Wood^a, Elizabeth Jackson^a and Charlie D. Frowd^a

^aSchool of Psychology and Computer Science, University of Central Lancashire, Preston, UK; ^bDepartment of Business Administration, Technology and Social Sciences, Humans and Technology, School of Psychology and Computer Science, Gävle, Sweden; ^cDepartment of Psychology, Nottingham Trent University, Nottingham, UK; ^dSchool of Applied Sciences, Edinburgh Napier University, Edinburgh, UK

ABSTRACT

Memory is facilitated by reflecting upon, or revisiting, the environment in which information was encoded. We investigated these “context reinstatement” (CR) techniques to improve the effectiveness of facial composites – visual likenesses of a perpetrator’s face constructed by eyewitnesses. Participant-constructors viewed a face and, after a one-day-delay, revisited (Physical CR) or recalled the environmental context (Mental/Detailed CR) before recalling the face and constructing an EvoFIT or a PRO-fit composite. Detailed CR increased correct naming of ensuing composites, but only when participant-constructors suitably encoded the environment. Detailed CR was also effective when combined with another interviewing technique (Holistic-Cognitive Interview), with focus on a target’s character; it was no more effective prompting constructors to engage in greater environmental recall. Analyses indicate that the Detailed CR advantage was mediated by an increase in face recall. Results are applicable by forensic practitioners to aid eyewitness memory, thereby potentially increasing suspect identification and subsequent arrest rates.

ARTICLE HISTORY

Received 9 October 2019
Accepted 9 February 2021

KEYWORDS

Context reinstatement; facial composites; EvoFIT; PRO-fit; Holistic-Cognitive Interview

Hearing an old song or returning to a place after a long time can unexpectedly revive memories thought to have been forgotten. Any aspect of the environment may trigger a memory and, often, the trigger can be peripheral to the retrieved episode. The phenomenon can be explained by the Encoding Specificity principle (Tulving & Thomson, 1973) which proposes that encoding a memory involves not only the central aspects of the episode but also information related to the context of the event. Contextual information includes an observer’s emotions at the time of encoding and their perception of the environment, which can involve a range of senses (e.g., visual, auditory, olfactory). For the visual modality in particular, there is usually a large amount of information that a person perceives about objects in the visual scene in terms of their shape, size and colour as well as their spatial arrangement and dynamics (i.e., whether and how they move). Such contextual information may later potentially act as retrieval cues, facilitating

access to the desired (or “target”) memory (e.g., returning to a childhood playground, a long-forgotten conversation, a person’s facial appearance). Recalling these ancillary retrieval cues should promote recall of the target memory.

The benefit of contextual reinstatement (CR) has been repeatedly demonstrated in the literature (e.g., see Smith, 2013 for a review; Smith & Vela, 2001 for meta-analysis), and in a variety of different areas including verbal memory (e.g., Campeanu et al., 2014; Godden & Baddeley, 1975; Smith, 1979), object memory (Barak et al., 2013; Koen et al., 2013; Levy et al., 2008), eyewitness memory (e.g., Dando, 2013; Wagstaff et al., 2011; Wong & Read, 2011) and facial memory (e.g., Davies & Milne, 1985; Shapiro & Penrod, 1986). Arguably one of the most well-known verbal learning studies was conducted by Godden and Baddeley (1975) in which divers learned word lists either on land or under water. Considerably more words (46%) were retrieved from

CONTACT Cristina Fodarella  cfodarella3@uclan.ac.uk  School of Psychology and Computer Science, University of Central Lancashire, Preston PR1 2HE, UK

© 2021 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group

This is an Open Access article distributed under the terms of the Creative Commons Attribution-NonCommercial-NoDerivatives License (<http://creativecommons.org/licenses/by-nc-nd/4.0/>), which permits non-commercial re-use, distribution, and reproduction in any medium, provided the original work is properly cited, and is not altered, transformed, or built upon in any way.

memory if recalled in the same environment as encoding than in the alternate environment, suggesting that reinstating the environmental context during retrieval facilitates memory recall.

Rather than reinstating the environmental context by physically returning to the location of encoding, comparable benefits to memory may be achieved by mentally visualising and recalling the environmental context in which the to-be-remembered target was encoded prior to retrieval (e.g., Smith, 1979), a cognitive or “mental” type of context reinstatement (MCR). In addition, the psychological state of the observer forms retrieval cues that can also facilitate memory recall. When a specific mood (Bower, 1981; Bower & Cohen, 1982; Eich & Metcalfe, 1989) or level of arousal (Clark et al., 1983) is congruent at learning and test, memory recall can be facilitated. Similarly, reproducing the pharmacological state of participants at learning (e.g., sleep deprivation, or influence of drugs or alcohol) can help trigger memory during recall, an effect termed state-dependent learning (see Eich, 1980; Eich et al., 1975). This indicates that besides physical factors, other associative and cognitive elements also form retrieval cues (Anderson & Bower, 1972).

While research assessing the effect of CR on the recall of verbal memory is theoretically interesting and may be applicable to a real-world setting, methodologies used in the literature often lack ecological validity. For example, remembering individual words through reinstating the context in form of a background image (e.g., Smith & Manzano, 2010) or the position on a screen (e.g., Macken, 2002; Murnane & Phelps, 1993) is not necessarily indicative of how memory might be improved in everyday situations. However, some research has also been *applied* in nature to produce findings that are applicable to real-life settings, most prominently on eyewitness memory. Results in this area also favour use of contextual cues for increasing recall (Dando et al., 2009; Hammond et al., 2006; Wagstaff et al., 2007; Wong & Read, 2011), with little or no increase in information that is inaccurate (e.g., Davis et al., 2005; Emmett et al., 2003; Memon & Bruce, 1995; Milne & Bull, 2002). One of the best-known methods for facilitating recall is the Cognitive Interview (CI).

The CI is an interviewing procedure usually administered by police officers for eyewitnesses to recall information about a crime (e.g., Geiselman et al.,

1985; Wells et al., 2007). It consists of a series of techniques or *mnemonics*, such as to report everything (to try to prevent witnesses from holding back information) and to attempt recall in a different temporal order (to encourage use of different retrieval paths). While the CI has been extensively assessed and improved (e.g., for a review, see Frowd, 2011), one of the mnemonics that has been consistently included since the original interview is to reinstate the context. Using Mental Context Reinstatement (MCR), observers are asked to mentally reinstate the environmental context at the time of encoding, taking into account other physical senses (e.g., smells, sounds) as well as their own psychological states (e.g., reactions, mood). MCR has been consistently shown to lead to the retrieval of a greater amount of information than a standard (“question and answer type”) police interview with only minor (or no) increase in inaccurate information recalled (Memon et al., 2010). In fact, as well as the “report everything” mnemonic, MCR is considered to be the most effective mnemonic for triggering retrieval (e.g., Davis et al., 2005; Milne & Bull, 2002).

Context can also facilitate memory for faces (e.g., Brown, 2003; Rainis, 2001; Shapiro & Penrod, 1986; Thomson et al., 1982). Studies using a line-up (or an identity parade) show that participants tend to recognize a target face significantly more accurately when tested in the same room where facial encoding initially occurred rather than in a different room (Evans et al., 2009; Wagstaff, 1982; Wong & Read, 2011). MCR is also effective in increasing correct recognition (Cutler et al., 1987; Malpass & Devine, 1981).

Another practical aspect of facilitating memory for faces, also likely to benefit from context reinstatement, be it physical or mental, is the construction of facial composites. Composites are facial likenesses usually produced from witnesses’ and victims’ memory of a perpetrator of crime. These visual representations of a face are used by law enforcement to identify potential suspects. There are many documented cases where facial composites have led to a serious criminal (e.g., a rapist, murderer, confidence) being identified and later – following further compelling evidence – convicted (e.g., Frowd, Pitchford, et al., 2012), and thus methods to improve their effectiveness are both theoretically interesting and valuable to security.

Sketch artists have traditionally worked with eyewitnesses to sketch a face, but production systems have since been developed, initially from mechanical

“feature” types (e.g., Photofit and Identikit), which are now archaic, to computerized “feature” systems (e.g., PRO-fit, E-FIT, FACES), and, more recently, “holistic” systems (e.g., EvoFIT, EFIT-V [now called EFIT-6], ID). A detailed review of the systems can be found in Frowd (2017) or Frowd ([in press](#)). In essence, feature systems allow an eyewitness to select individual facial features (e.g., eyes, nose, mouth, hair) to construct a facial composite, while holistic systems involve witnesses repeatedly selecting from screens of whole faces (or whole-face regions), with choices combined, to “evolve” a facial likeness. Composite recognition then occurs based on both featural information, that is, individual facial features, and configural information, that is, the spacing between these features (see e.g., Fodarella & Frowd, 2013; Frowd et al., 2014). Research suggests that holistic systems are more effective at producing identifiable composites presumably because they are based on face recognition processes (which are more stable over time; Davies, 1983) rather than recall processes (e.g., Frowd, 2017; Frowd et al., 2010, 2015).

Davies and Milne (1985) appear to be the first published study to investigate the influence of context on facial-composite construction. In their work, one of four target individuals was seen entering a room and searching for a calculator. After a one-week delay, observers constructed a composite face using the now archaic Photofit while in the same room or a different room, and following a guided memory procedure for recalling the environmental context or without such a procedure. The ensuing composites were given to other people to match to a recent photograph of the targets. Matching scores indicated that significantly better-quality composites were produced (i) under guided memory (cf. spontaneous memory recall) and, to a lesser extent, (ii) in the same (cf. different) room. In other work, Ness and Bruce (2006) investigated a novel procedure for reinstating physical context for face construction. Using the modern PRO-fit system, constructors who were given the opportunity to review video footage of the encoding environment (for a target identity just seen) created composites with better likenesses, faces that were matched more accurately to the target identity.

The current project investigates if cues available in the environment would enable participants to create a more accurate face using modern composite systems after a realistic delay. In Experiment 1,

participants viewed an unfamiliar target face and underwent one of three procedures the following day to reinstate context. Participants were met by an experimenter either in the original environment where target encoding had taken place (Physical CR), or in a different environment for a mental CR. For the latter, participants underwent either (i) Minimal CR, where they were instructed to “think back” to the environment, or (ii) Detailed CR, as (i) but then recalling both the environment and the person’s mood and feelings at the time.¹ Afterwards, participants freely described the face (using further mnemonics of the CI) and constructed a single composite of it using either a typical feature system (PRO-fit) or a typical holistic system (EvoFIT). The resulting composites were given to another group of participants who attempted to name these composites as measure of effectiveness.

Based on the aforementioned research, it was predicted that Detailed and Physical CR would improve a constructor’s face-recognition ability during composite construction for both systems. Therefore, more identifiable composites would be constructed (i) under Detailed CR than Minimal CR, and (ii) under Physical CR than Minimal CR. The literature also suggests that composites should be constructed more effectively from holistic than from feature systems, and we expected to observe the same result.

Following on from Experiment 1, subsequent experiments in this paper focus on combining the CR technique with other interviewing mnemonics to potentially increase composite effectiveness further. Additionally, participant-witnesses’ attention to the environment at encoding will be manipulated in an attempt to decipher whether the CR technique may be stronger when attention is specifically directed to the environment and thus is intentionally encoded. Finally, the underlying mechanism driving CR effects will be explored by investigating whether its effectiveness is due to increased face recall or face recognition, or a combination of both.

EXPERIMENT 1: Using context reinstatement to facilitate face construction using modern composite systems

Method

The most effective facial composite research mirrors, to the greatest extent possible, the real-life situation

in which a witness or victim observes a (usually unknown) perpetrator during a crime, and is required to create a visual likeness of the face after a minimum of 24 h (Frowd et al., 2005). When the composite is subsequently shown to police officers or the public, anyone who is familiar with the individual may recognize the composite and be able to provide investigators with a possible identity. In this experiment, we model this situation in two stages: (i) composite construction, where a composite of a target face is created from memory 20–28 h later, by someone unfamiliar with the target, and (ii) composite naming, where someone familiar with the target is asked to identify the composite.

In the experiments presented here, we chose to present target faces as static images; this is often the case in composite research (e.g., Fodarella et al., 2017; Frowd, Skelton, et al., 2012; Gawrylowicz et al., 2012; Hasel & Wells, 2007; Kehn et al., 2014). Although it could be argued that a staged event or use of video stimuli would be more realistic, composite identifiability changes little following a target presented in video or as a static image (Frowd et al., 2015). Therefore, use of static images seems to produce realistic findings, generalisable to viewing a face in motion.

Participants viewed a target face in the knowledge that they would produce a composite on the following day. This tends to promote an intentional type of encoding likely to increase memory (Shapiro & Penrod, 1986); however, this design choice is not different to how real eyewitnesses may remember faces. In many cases, witnesses and victims make a deliberate attempt to remember the appearance of an offender's face, having the intuition that they will be asked about the face at a later date (Frowd et al., 2015), so we copy this method of encoding here.

The aim was for the experiments to have sufficient power to be able to detect a medium-to-large effect size, should one exist. While dependent on variability of data, this effect size usually leads to around (at least) 50% difference in mean correct naming – for example, a mean of 20% correct in one condition and 30% in another $[(30-20) / 20 = 50\%]$. This aim was achieved using ten target identities and at least ten participants per group for composite naming (Frowd, 2015). Such an increase should translate into a useful benefit for policing; in the paper, we report Cohen's d for composite naming, where a

value of 0.5 is considered a “medium” effect and 0.8 as “large” (Cohen, 1988).

Stage 1: Composite construction

Design

A 2 (System: EvoFIT, PRO-fit) $\times$ 3 (CR: Minimal, Physical, Detailed) between-participants design was used for composite construction. Context reinstatement was manipulated over three conditions. The Minimal CR (“control”) condition consisted of “thinking back” to the encoding environment, with face recall and composite construction taking place in a different environment to that in which the face had been seen (encoded). The Physical CR condition was the same as the first except that face recall and composite construction were conducted in the same environment (room) as encoding. The third condition was Detailed CR. As in the first condition, a different environment was used to that in which the face had been encoded, and, prior to face recall and composite construction, participants were asked to recall both the environment and their moods and feelings at that time. In all cases, after the relevant CR manipulation and face recall, each participant created a single composite using either the holistic system EvoFIT or the feature system PRO-fit.

Participants

An opportunity sample of 60 (24 males, 36 females; $M_{age} = 30.3$, $SD_{age} = 11.3$ years) participants took part either on a voluntary basis or in return for course credits. They were staff and students from the University of Central Lancashire (UCLan). To simulate the usual situation for real eyewitnesses, participants were recruited on the basis of being unfamiliar both with the testing environment (a student café) and with the target faces. Participants' race was not recorded as there is no evidence for an own-race bias in composite construction involving Caucasian faces (Bhardwaj & Hole, 2020; McQuiston-Surrett & Topp, 2008).

Materials

Target faces were 10 current characters from a popular UK soap opera, Coronation Street, sourced from the Internet (Ken Barlow, Leanne Battersby, Peter Barlow, Michelle Connor, Jason Grimshaw, Tracey McDonald, David Platt, Kirk Sutherland, Sally

Webster and Sophie Webster). These pictures were of good quality, shown in full-frontal pose with minimal facial expression; male actors had little or no facial hair. Stimuli were printed in colour to dimensions of 8 cm (width) × 10 cm (high).

Procedure

After providing consent, participants were tested individually, and the procedure was self-paced. To allow good control of exposure to the testing environment (described below), each person was met at a convenient meeting point and taken to the room used for target encoding. Participants were randomly assigned to one of the six experimental conditions, defined above, with equal sampling. Each person was shown a target picture, randomly selected, and asked whether the identity was familiar; if it was, another picture was similarly shown. For the first face reported to be unknown, participants were given 60 sec to remember the face. For this part of the procedure, the participant was aware that a composite would be constructed of this face the following day. The experimenter was blind to the identities included in the experiment as well as the face seen by each participant. Participants viewed the face in a student café located in an unfamiliar building on the UCLan (Preston) campus, an environment selected to be unfamiliar to participants as well as rich in environmental recall cues: it included tables, chairs, a television, plants, a vending machine and a small counter selling refreshments and confectionary. Participants' unfamiliarity with the environment was established by asking whether they had visited the café prior to the experiment. If anyone had reported previously visiting the café, they would have been excluded (there were no such occurrences). The experimenter made no reference to the importance of the café (to allow environmental context to be encoded incidentally).

Following a delay of 20–28 h, according to assignment, participants met the researcher either in the student café for Physical CR, or in a neutral office space for Minimal and Detailed CR. For Minimal and Physical CR, participants were first interviewed via Cognitive Interview² in which they were asked to mentally visualize the target face and then freely recall it in as much detail as possible. In the Detailed CR condition, prior to face recall, participants were also asked to mentally visualize the environment in

which they saw the target face, and then to reflect silently on their mood and psychological state at the time of viewing. Following this, participants were asked to freely recall the environment as well as their mood and feelings at the time. As elsewhere, participants then freely recalled the face.

Once face recall had been completed, each participant constructed a single composite of the target on a laptop computer using EvoFIT or PRO-fit. The experimenter controlled the relevant software programme and took the participant through the procedure to construct the face, the aim of which was to construct the best likeness possible. A detailed description of the relevant procedure for each system can be found in Fodarella et al. (2015). In brief, for PRO-fit, participants were asked to select the best matching facial features (e.g., eyes, brows, nose, mouth, hair, ears) for their given target, and then resize and position each feature to give the best likeness possible. For EvoFIT, participants were asked to repeatedly select overall best matches from arrays of internal features to evolve a composite, use software tools to enhance the overall likeness and facial features, and then add external features (hair, ears, neck).

The procedure to construct a facial composite took between 20 and 45 min per person including debriefing.

Stage 2: Composite naming

Design

This time, participants who were familiar with Coronation Street characters were recruited. They were given a set of composites to name that had been constructed using one of the six individual procedures in Stage 1. Thus, the design was a 2 (System: EvoFIT, PRO-fit) × 3 (CR: Minimal, Physical, Detailed) between-participants. It is worth mentioning that this design may lead to an elimination strategy, with participants deciding between possible identities when attempting to name a face. In real life, an elimination strategy may also be involved to some extent, as other information (e.g., offender's build, age and accent) is usually published alongside the composite. In addition, given the nature of the design, one would imagine that any such strategy would apply equally to each experimental condition. Overall, the naming procedure used here has been found to lead to consistent results (e.g., Frowd et al., 2015), and to be a

good indicator of identification of composites in the real world (Frowd, Pitchford, et al., 2012).

Participants

Forty-eight (42 males, 6 females; $M_{age} = 41.0$, $SD_{age} = 15.3$ years) volunteer Coronation Street fans were recruited outside the “Coronation Street Tour Set” in Manchester. These participants were familiar with the relevant identities, and the testing procedure (as detailed in *Procedure* below) ensured that they would recognize a minimum of 80% of the targets.

Materials

Sixty composites and the 10 target photographs were standardized to dimensions of 8 cm (width) $\times$ 10 cm (high) and printed individually in greyscale; see [Figure 1](#) for example composites.

Procedure

Participants were randomly allocated with equal sampling to the two experimental factors (i.e., one of six individual conditions) described above. During briefing, they were informed that they would view and should attempt to identify a set of composites depicting Coronation Street characters. After providing consent, participants were shown 10 facial composites (based on assignment) sequentially to

name. Participants were encouraged to guess, and it was explained that it was also acceptable to give identifying semantic information if they were unable to remember a name, or not to give a name at all. Once all 10 composites had been presented, the target photographs were shown likewise for naming. We applied an *a priori* rule, to ensure participants were suitably familiar with the relevant identities, such that each person was required to correctly name a minimum of eight of the ten targets (or another person was to be recruited as replacement); as it turned out, all participants met this rule. Participants each received a different random order of presentation of composites and target photographs. Testing sessions lasted for approximately 10–15 min, including debriefing.

Results

Composite naming

Participant responses to target photographs and facial composites were scored for accuracy: a value of 1 was assigned when a given response was correct and 0 otherwise (incorrect name or no name given). Participants correctly named all of the target photographs and accordingly familiarity with the relevant identities was at 100%. This result suggests that all of the composites had the potential to be correctly named by all of the participants. As can be seen in [Table 1](#), mean correct naming of composites was considerably less than 100%, which is the usual case for this type of error-prone facial stimuli.

The number of correct responses per participant was analysed using Independent Samples ANOVA for CR type (Minimal, Physical, Detailed) and System (EvoFIT, PRO-fit). There was a significant main effect of CR [$F(2,42) = 3.27$, $p = .048$, $\eta_p^2 = .14$] and two-



Figure 1. Example EvoFITs (top row) and PRO-fits (bottom row) of Coronation Street character Leanne Battersby constructed following (a) Minimal, (b) Physical, and (c) Detailed CR. Each composite was constructed from memory by a different person. Due to copyright issues, we are unable to reproduce the target face used in the experiment, but an internet search should readily reveal the appearance of this actress, Jane Danson.

Table 1. Percentage of EvoFIT and PRO-fit composites correctly named by Context Reinstatement and System.

System*	Context Reinstatement (CR)†			Mean
	Minimal	Physical	Detailed	
EvoFIT	23.8 (17.7)	22.5 (14.9)	33.8 (9.2)	26.7 (14.6)
PRO-fit	6.3 (5.2)	7.5 (4.6)	13.8 (9.2)	9.2 (4.0)
Mean	15.0 (15.5)	15.0 (13.2)	23.8 (13.6)	17.9 (14.4)

Note: † Significant main effect of CR, $p < .05$; * Significant main effect of facial-composite System, $p < .001$. In parentheses are (by-participant) SD values.

tailed Simple Contrasts indicated that Detailed promoted significantly higher naming rates than both Minimal ($p = .032$, Cohen's $d = 0.66$) and Physical ($p = .032$, $d = 0.66$) CR; an additional t-test revealed that there was no significant difference between Minimal and Physical [$t(30) < 0.001$, $p = 1.00$, $d < 0.01$] CR. The main effect of System was also significant [$F(1,42) = 29.40$, $p < .001$, $\eta_p^2 = .41$], with EvoFIT composites named significantly higher than PRO-fit composites. The interaction between CR and System was not significant [$F(2,42) = 0.20$, $p = .82$, $\eta_p^2 = .01$].

Additional analyses

There are other methods of assessment that can be carried out on data produced from the experiment. First, incorrect (mistaken) naming of composites can be analysed to give a measure of composite misidentification and indication of response bias (guessing). Second, ratings of similarity (likeness) between each composite and its target face can be collected by another group of participants as a supplementary measure of composite utility. Third, information recalled about the face from each participant who constructed a composite can also be assessed, hypothesized to be greater under Physical and Detailed CR (cf. Minimal). In order to maintain brevity, these additional assessments are not reported in detail in the current paper, but have been conducted. In brief, across all experiments presented in the paper (i) there was no reliable difference by CR in terms of incorrect names given, (ii) likeness ratings generally mirrored the pattern of results from correct composite naming, and (iii) significantly greater information was recalled overall about the target face by participant constructors following Detailed compared to Minimal CR (these measures are discussed in more detail in the Discussion following Experiment 3; for Method and Results of face recall analysis, see Appendix 1).

Discussion

Our results support Davies and Milne's (1985) findings: composites were more identifiable in the Detailed CR condition compared to both Physical and Minimal CR (Control). This was shown using two types of composite system: a modern feature system, PRO-fit, and one of the newer holistic systems, EvoFIT, with the latter system outperforming the former. Overall, the data suggest that detailed recall of the physical and

psychological context is advantageous for reproducing faces from memory using two contrasting methods of face production. Our working hypothesis is that Detailed CR would be effective as it improved a constructor's memory of the face, allowing them to more-effectively process the face as a whole – that is, leading to a better end result whether that is achieved through selection of individual features (PRO-fit) or from face arrays (EvoFIT).

Experiment 2 attempted to replicate the effect of the Detailed CR procedure and ascertain whether participants would be able to construct even more effective composites using an enhanced type of interview. In Experiment 1, prior to face construction, witnesses were interviewed using a Cognitive Interview (CI) to help them recall the appearance of the target face. An enhancement of this interview, termed the Holistic-Cognitive Interview (H-CI), involves focus on the personality or character of the face (e.g., Frowd et al., 2008; Frowd, Nelson, et al., 2012). Specifically, after face recall, witnesses are asked to reflect silently on the personality of the face for 60 sec and then make seven personality judgements (e.g., intelligence, friendliness, kindness) on a three-point Likert scale. The procedure is believed to improve a constructor's face recognition by encouraging holistic processing of the face. The resulting composite is more identifiable than that produced following the more standard face recall CI (see Frowd et al., 2015 for a meta-analysis).

Therefore, Experiment 2 involved three factors: CR (Minimal CR, Detailed CR), method of witness interview (CI, H-CI) and system (EvoFIT, PRO-fit). If Detailed CR promotes a better memory of the face, then using H-CI should improve composite effectiveness even further. Based on the Experiment 1 composite naming data, Physical CR was not considered any further since there was no evidence that it helped participants to construct a more identifiable composite; we consider potential explanations for this (null) effect in the General Discussion.

EXPERIMENT 2: Combining Context Reinstatement and Holistic-Cognitive Interview

Stage 1: Composite construction

Design

A 2 (System: EvoFIT; PRO-fit) $\times$ 2 (CR: Minimal, Detailed) $\times$ 2 (Interview: CI, H-CI) between-

participants design was used. Thus, nominally 24 h after encoding an unfamiliar target face in an unfamiliar environment (café), participants underwent Minimal or Detailed CR, described the face using CI or H-CI and constructed a single composite using EvoFIT or PRO-fit.

There was also a change in the procedure used with EvoFIT. Recent findings reveal that composites are produced more identifiably if constructors are requested to select a match for the *upper half* rather than for the overall match of the presented internal features arrays when evolving the face (Fodarella et al., 2017). This enhanced procedure tends to produce composites with a more accurate upper facial region, an area known to be important to recognition of both facial photographs (Goldstein & Mackenberger, 1966; Pellicano et al., 2006) and facial composites (Laughery et al., 1986). This instruction is now used regularly with witnesses and victims of crime (Frowd et al., 2019), and including it here allows results to reflect current forensic practice.

Participants

Sixty-four (26 males, 38 females; $M_{age} = 29.1$, $SD_{age} = 8.6$ years) UCLan staff and students took part voluntarily. As in Experiment 1, participants were recruited on the basis of being unfamiliar with the target identities.

Materials

To have greater confidence in the generalisability of the CR advantage, target faces were drawn from a different pool of identities: current football players who play at international level in the UK. Recruitment of participants followed the same criteria as before (i.e., participants were recruited to be unfamiliar with targets for the face construction stage, but familiar for the naming stage). Target faces were eight photographs (Ross Barkley, Gary Cahill, Michael Carrick, Joe Hart, Harry Kane, Adam Lallana, James Milner and Jack Wilshere) sourced from the Internet and of the same standard as Experiment 1, printed in colour (8 × 10 cm).

Procedure

Each participant viewed a target face in the same café as in Experiment 1, and met with the experimenter the following day in a different room to construct a composite of this face. The procedure was the same as in Experiment 1, except for the following

differences. At the start of the second session, as Physical CR was no longer used, participants engaged in either Minimal CR or a Detailed CR. After this CR manipulation, participants were interviewed using either (i) the CI, in which participants were asked to mentally visualize and then freely recall the target face in as much detail as possible, or (ii) the H-CI, as (i) but then were asked to reflect silently on the personality of the face for 60 sec and make seven personality attributions, rating each on a three-point Likert scale. After the interview, participants created a single composite using either PRO-fit, as described in Experiment 1, or EvoFIT. The EvoFIT procedure differed from that used in Experiment 1, in so far as participants were instructed to select best matches in the presented arrays for the upper facial region; after evolving the face, participants were requested (as before) to focus on all aspects of the face (not just the upper region), in order to enhance the facial appearance using the software tools.

Stage 2: Composite naming

Design, material and procedure

Composites were named by a different group of participants using the same three-factor design as described in Stage 1. Materials were 64 composites and eight target photographs, printed individually in greyscale as in Experiment 1. The procedure used to name the composites was also the same as in Experiment 1, except that there were now 64 composites and eight targets, and participants were randomly allocated, with equal sampling, to the eight cells of the design.

Participants

Eighty (77 males, 3 females; $M_{age} = 40.6$, $SD_{age} = 14.5$ years) participants took part on a voluntary basis. They were recruited opportunistically from Manchester Football Museum on the basis of being familiar with the target identities.

Results

Composite naming

Participant responses to targets and composites were scored for accuracy. Familiarity with the target identities was at 100%; a summary of composite naming is shown in Table 2.

Table 2. Percentage of EvoFIT and PRO-fit composites correctly named by Context Reinstatement, System and Interview.

System†	Context Reinstatement*				Mean
	Minimal		Detailed		
	CI	H-CI‡	CI	H-CI‡	
EvoFIT	10.0 (9.9)	22.5 (17.5)	26.3 (16.1)	31.3 (18.9)	22.5 (17.3)
PRO-fit	2.5 (5.3)	6.3 (6.6)	5.0 (6.5)	6.3 (8.8)	5.0 (6.8)
Mean	10.3 (12.9)		17.2 (17.6)		13.8 (15.7)

Note: † Significant main effect of System, $p < .01$; * Significant main effect of CR, $p < .02$; ‡ Significant main effect of Interview, $p < .05$; Significant CR $\times$ System interaction, $p < .05$. In parentheses are (by-participant) SD values.

Independent Samples ANOVA revealed a significant main effect of (i) CR [$F(1,72) = 6.26$, $p = .015$, $\eta_p^2 = .08$], indicating Detailed CR led to better-named composites than Minimal CR, (ii) Interview [$F(1,72) = 4.19$, $p = .044$, $\eta_p^2 = .06$], with H-CI ($M = 16.6$, $SD = 17.3$) leading to better-named composites than CI ($M = 10.9$, $SD = 13.6$), and (iii) System [$F(1,72) = 40.55$, $p < .001$, $\eta_p^2 = .36$], as EvoFIT composites were named more accurately than PRO-fit composites. The interaction between CR and System was also significant [$F(1,72) = 4.19$, $p = .044$, $\eta_p^2 = .06$], as Detailed CR significantly increased naming (cf. Minimal CR) for EvoFIT [$t(38) = 2.43$, $p = .020$, $d = 0.66$] but not for PRO-fit ($p = .57$, $d = 0.33$). All other interactions were non-significant ($F < 1.30$, $p > .25$).

Discussion

Experiment 2 revealed an overall benefit for the Detailed CR manipulation improving correct naming of composites for both types of system. Also, composites produced via the H-CI interview procedure were overall significantly better named than those following the CI technique. Finally, EvoFIT composites were named significantly better than PRO-fit composites.

In Experiment 3, we were interested in manipulating participants' attention to the environment in attempt to explore whether the contextual effect observed thus far would be stronger if participants are specifically asked to encode the environment. Therefore, attention of half the face constructors was explicitly directed to the environment during encoding. The approach is in line with past research indicating that an MCR benefit is more likely to occur when experimental instructions emphasize a so-called interactive encoding between study items (in this case, the target face) and the environmental context (see

Hanczakowski et al., 2014; Hockley, 2008). Interactive encoding between items and context would ensure a stronger association between the two, which would presumably in turn ensure that they act as retrieval cues for one another during recall. While we are not directly manipulating the face to interact with the environment, the aforementioned evidence implies that increased attention to the environment should promote stronger context effects.

An attempt was also made in this experiment to facilitate the potential benefit of the environment on face construction (for one of the groups of participants) in another way. The approach was based on the theory that memory for information can be facilitated by cued recall. As well as MCR, the extensive recall technique is an effective interviewing mnemonic of the Cognitive Interview, used to facilitate eyewitness recall (Fisher & Geiselman, 1992). Put simply, memory recall should be greater when multiple retrieval attempts are made rather than stopping after an initial memory search (Fisher & Geiselman, 1992). In Experiment 3, therefore, once participants had provided free recall of the environmental and internal context, they were asked to try to remember further information; specifically, participants were prompted (or "cued") by open-ended questions based on the information they had recalled. For example, having mentioned the presence of chairs and tables in the room, participants would then be asked if they could recall further information about these items. It was hypothesized that this "cued" technique would lead to greater recall of the environment, which in turn should facilitate memory of the face as well as production of a composite.

The PRO-fit feature system was not used in this or in the following experiments, principally due to low naming rates produced from its composites and the ensuing difficulty of then making sensible conclusions. The following experiments therefore focussed on the EvoFIT system.

EXPERIMENT 3: Increasing the focus of attention on the environmental context

Stage 1: Composite construction

Design

The design was between-participants: 2 (Context Attention: Incidental, Intentional) $\times$ 3 (CR: Minimal,

Detailed, Extensive). In the Intentional condition, participants' attention was directed to the environment by asking them to inspect the environment closely prior to viewing the target face; participants in the Incidental condition did not receive these instructions, with encoding of the environment carried out in the same (incidental) way as before. Prior to composite construction, CR was manipulated on three levels: (i) Minimal CR, (ii) Detailed CR, and (iii) Extensive CR as (ii) with participants then invited to try to remember further information as part of a cued-recall format.

Participants

Sixty (20 males, 40 females; $M_{age} = 30.4$ years, $SD_{age} = 9.4$ years) UCLan staff and students volunteered or received course credits in return. They were recruited opportunistically on the basis of being unfamiliar with the target faces.

Materials

Ten photographs of characters from the TV soap East-Enders were target images (Ian Beale, Jane Beale, Jack Branning, Lauren Branning, Max Branning, Stacey Branning, Shirley Carter, Martin Fowler, Billy Mitchell and Jean Slater). They were of the same standard as in the previous experiments and were printed likewise.

Procedure

The construction procedure was the same as the two previous experiments except for the following differences. Prior to face encoding, participants were either asked to pay close attention to the café environment in which the face was to be subsequently shown (intentional encoding of context), or did not receive such an instruction (incidental encoding). If participants assigned to the former condition did not seem to study the room (which occurred about a third of the time, $N = 11 / 30$), the experimenter gave a prompt: "I will give you a little more time to look at the environment". The following day, participants were interviewed in one of three conditions prior to constructing the face via EvoFIT (incl. the instruction to select for the upper facial half in face arrays). Minimal and Detailed CR conditions were administered as before. Extensive CR followed the Detailed CR procedure, after which participants were asked questions about objects recalled in the

environment. For instance, "You remembered tables and chairs. Can you say anything more about these?". The researcher prompted for further information in the order in which objects had been initially recalled.

Stage 2: Composite naming

Design, materials and procedure

The design was the same as in Stage 1. Materials were 60 composites and 10 target photographs, printed as before. The naming procedure was the same as in Experiments 1 and 2.

Participants

Sixty (15 males, 45 females; $M_{age} = 41.6$, $SD_{age} = 12.7$ years) staff and student volunteers were recruited opportunistically on the UCLan campus. Participants were recruited on the basis of being familiar with the target identities.

Results

Composite naming

Composites and targets were scored in the same way as in Experiments 1 and 2. Familiarity with the target identities was once again high, at 98.8%; mean composite naming is shown in Table 3.

Independent Samples ANOVA revealed a non-significant main effect of CR [$F(2,54) = 3.09$, $p = .05$, $\eta_p^2 = .10$], but a significant main effect of Attention [$F(1,54) = 11.91$, $p = .001$, $\eta_p^2 = .18$], with composites named significantly better when attention was directed to the environment (cf. Incidental). These two factors also interacted with each other [$F(2,54) = 3.34$, $p = .043$, $\eta_p^2 = .11$]. When attention was Incidental, there were no significant differences between CR conditions ($p > .05$). However, when attention was

Table 3. Percentage of EvoFIT composites correctly named by Context Reinstatement and Attention

Attention*	Context Reinstatement			Mean
	Minimal	Detailed	Extensive	
Incidental	23.0 (15.0)	23.0 (18.3)	22.0 (19.9)	22.7 (17.2)
Intentional	22.0 (14.8)	49.0 (22.8)	50.0 (25.8)	40.3 (24.7)
Mean	22.5 (14.5)	36.0 (24.1)	36.0 (26.6)	31.5 (22.9)

Note: * Significant main effect of Attention, $p < .01$; Significant Attention $\times$ CR interaction, $p < .05$. In parentheses are (by-participant) SD values.

Intentional, composites were named significantly better if constructed following Detailed [$p = .006$, $d = 1.40$] and Extensive [$p = .008$, $d = 1.33$] compared to Minimal CR; there was no significant difference between Detailed and Extensive CR [$t(18) = 0.09$, $p = .93$, $d = 0.05$]. Also, Intentional was significantly better than Incidental, for Detailed [$p = .012$, $d = 1.26$] and Extensive CR [$p = .014$, $d = 1.21$], but there was no reliable difference between the two Minimal CR conditions [$t(18) = 0.15$, $p = .88$, $d = 0.07$].

Discussion

When participants' attention was directed to the environment, composite naming increased significantly compared to incidental encoding. However, this effect was driven by higher naming of composites created under Detailed and Extensive CR. Consequently, the advantage of Detailed CR (cf. Minimal CR) emerged when the environment had been encoded *intentionally*. When attention to the environment was incidental, Detailed and Extensive CR were not effective in increasing composite naming. Overall, these results imply that, in two out of three experiments, participants had encoded the environment to some extent despite not having been specifically asked to do so, and this influence positively affected ensuing likenesses. Participants in Experiment 3, however, were not benefitting from context cues when their attention was not directed to the environment. This seems to insinuate that incidental encoding can lead to inconsistent findings regarding the effectiveness of Detailed CR as a result from insufficient environmental encoding.

Results further revealed that Minimal CR did not differ reliably between incidental and intentional attention. This finding would appear to be sensible since Minimal CR following intentional encoding did not then make active use of the environment. This implies that trying to remember the environment at the point of encoding is not sufficient in itself to later provide access to the target memory: what is necessary is to actively visualize and recall the environment. A final noteworthy result relates to Extensive CR. It was predicted that this condition would promote more effective composites than Detailed CR. This was not the case in either incidental or intentional encoding. Thus, attempts to recall more information about the environment, even if there

should have been more information available to recall (esp. following intentional encoding), does not seem to have influenced the ensuing composites. The implications of these findings are discussed in greater depth in the General Discussion.

We argued earlier that an absent or weak advantage for Detailed CR might stem from insufficient encoding of the encoding environment. In the next experiment, participants again encoded the environment incidentally (as was the design of Experiments 1 and 2, and for components of Experiment 3) but this time we attempted to improve memory for the environment more-naturally. Previously, participants followed the experimenter into the room in which the target face had been seen. However, this procedure may not promote good encoding of the environment: participants generally engaged in informal conversation with the researcher and may have followed "blindly" to the table where face encoding was carried out. Thus, they may have paid little attention to the environment, preoccupied in conversation and the subsequent testing task. This suggestion would seem to be in line with the theory of environmental suppression (Glenberg, 1997), whereby participants were focussed on the conceptual processing of face encoding, leading to poor encoding of the environment.

Therefore, participants in Experiment 5 were invited to enter the room in front of the researcher (who waited outside of the room until seated) and to navigate to a specific table as requested. In doing so, it was hypothesized that participants should naturally take in details of the relevant context in order to navigate the room, without engaging in conversations during this process. This should also more-closely reflect how eyewitnesses encode an environment in the real world (as opposed to the intentional encoding we used in the previous experiment).

We were also interested in assessing a simpler version of the memory enhancement technique, one that does not require participants to recall their emotional state at the time of encoding. Crimes for which composites are usually constructed involve considerable stress for an observer (esp. in cases of assault). Recalling such evocative information may be traumatic, potentially leading to anxiety and the subsequent effect of inhibiting composite construction (Davies, 2009); conversely, attempts to reduce anxiety also seem to promote more-identifiable

composites (Martin et al., 2017). We reasoned that recall of the environment on its own should be sufficient to facilitate face construction (i.e., without the need for recall of psychological/emotional states).

We also thought it could be beneficial to law enforcement to assess the effectiveness of CR for a novel method of face construction. Recent technological advances have enabled witnesses to construct a composite themselves in their own time (e.g., Martin et al., 2018). This development allows composites to be used in investigations of less serious crime (e.g., minor theft, vandalism and anti-social behaviour), cases where police practitioners may not have the time (e.g., Alison et al., 2013) to interview witnesses to create a composite. As such, the production team behind EvoFIT created a “self-administered” version (www.EvoFIT.co.uk). This novel approach was designed to be functionally equivalent to the “face-to-face” method used by police practitioners (as followed in the current experiments), with witnesses taken through the same procedure by following written instructions on-screen in their own home.

In light of this development, the following experiment assessed whether the advantage of Detailed CR would be apparent for participants who constructed the face in the normal manner via face-to-face interview, or by themselves. The experiment recruited residents from a small town in the UK (cf. university staff and students) to extend findings to other participant pools, with target encoding taking place in a small office (unfamiliar to participants) and face construction (for both face-to-face and self-administered) elsewhere in a relaxed home environment.

EXPERIMENT 4: Facilitating more-naturalistic encoding of the environment

Stage 1: Composite construction

Design

The design was between-subjects with two experimental factors: 2 (CR: Minimal, Detailed) $\times$ 2 (Face construction: Face-to-face, Self-administered). It was the same as the previous experiments by CR (Minimal vs. Detailed), except that Detailed CR did not involve recall of the participant’s emotional state at encoding, and half of the composites were constructed using a self-administered procedure

(with the other half produced with the assistance of the experimenter, as before). The other change relates to the room used for target encoding. It was a small office, previously unseen by the recruited participants, with potentially useful recall cues: computer, chair, desk, bookcase, stationery, etc. To facilitate encoding of the environment, participants were invited to enter the room in front of the experimenter and navigate to where they were instructed to sit.

Participants

Thirty-two (10 males, 22 females; $M_{age} = 24.0$ years, $SD_{age} = 7.8$ years) local residents of a small town in the UK (Whitchurch, Shropshire) volunteered. They were recruited opportunistically, on the basis that they were not familiar with the target faces.

Materials

Eight photographs of current EastEnders characters were target images (Ian Beale, Jack Branning, Lauren Branning, Max Branning, Stacey Branning, Shirley Carter, Billy Mitchell and Jean Slater), sourced and printed as before to the same standard.

Procedure

Face construction was the same as Experiments 1–3, except for the following differences. As before, participants were randomly assigned to CR (Minimal, Detailed), and now to Face construction (Face-to-face, Self-administered). Face-to-face construction proceeded as described previously. Those assigned to self-administered construction wrote down on a piece of plain paper what they could remember of the environmental context and then, on the reverse side, provided a free-recall description of the previously seen face. When complete, participants followed the on-screen instructions on a laptop computer as described above to construct an EvoFIT composite.

Stage 2: Composite naming

Design, materials and procedure

The two-factor design was the same as in Stage 1. Thirty-two composites and eight target photographs were printed in greyscale (8 $\times$ 10 cm). Except for variation in the two experimental factors (CR and

Table 4. Percentage of EvoFIT composites correctly named by Context Reinstatement and by Face Construction

Face Construction*	Context Reinstatement†		
	Minimal	Detailed	Mean
Face-to-face	54.7 (9.3)	71.9 (12.9)	63.3 (14.0)
Self-administered	42.2 (9.3)	50.0 (11.6)	46.1 (10.9)
Mean	48.4 (11.1)	60.9 (16.4)	54.7 (15.1)

Note: † Significant main effect of CR, $p < .005$; * Significant main effect of Face Construction, $p < .001$. In parentheses are (by-participant) SD values.

Face construction), the naming procedure was the same as that described previously.

Participants

Thirty-two (13 males, 19 females; $M_{age} = 20.0$, $SD_{age} = 1.5$ years) students were recruited opportunistically on the UCLan campus. Participants took part on a voluntary basis or in return for course credits, and were recruited to be familiar with the target identities.

Results

Responses to composites and target photographs were scored for accuracy. All participants correctly named all eight target photographs, except for one participant who named seven, and consequently familiarity with the target set was high, at 99.7%. Table 4 provides a summary of composite naming, again suggesting benefit for the Detailed CR procedure.

Independent Samples ANOVA revealed a significant main effect of CR [$F(1,28) = 10.54$, $p = .003$, $\eta_p^2 = .27$], with composites named higher when constructed following Detailed than Minimal CR ($d = 0.92$). Face construction was also significant [$F(1,28) = 20.00$, $p < .001$, $\eta_p^2 = .42$], with composites constructed face-to-face named higher than when self-administered ($d = 1.37$). The interaction between CR and Face construction was not reliable [$F(1,28) = 1.48$, $p = .23$, $\eta_p^2 = .05$].

Discussion

Experiment 4 revealed that Detailed CR was still effective (cf. Minimal CR) when participants were asked to recall only the environmental context (i.e., they omitted recall of their psychological state at encoding), thus indicating a persistent advantage of Detailed CR with recall of physical environmental

cues only. Although overall results indicate that self-administered composites were less identifiable than those produced face-to-face with the interviewer (experimenter), the positive effect of Detailed CR was consistent for both methods of face production. Thus, our effort to ensure (as far as possible) that participants encoded the environment sufficiently seems to have been successful.

We were also interested in identifying the likely engendering mechanism which drives the effectiveness of Detailed CR, to thereby better-understand its theoretical underpinning. Davies and Milne (1985) noted that there appear to be two main mechanisms by which face construction could be rendered more effective using reinstatement techniques. One mechanism could be that CR increases witnesses' face *recognition*, leading to more accurate selection of individual facial features or, in the case of a holistic system, whole-face regions. The other mechanism, which does not preclude the former, could be that CR promotes a better memory of the face. This explanation can be evidenced by an increase in witness *recall* of the face and may allow witnesses to construct composites with more accurate detail – for example, to create faces with more accurate feature shapes. Davies and Milne (1985) note that the effect of both processes in their work may have been hampered by task difficulty and low experimental power. Therefore, we attempted to overcome these issues by combining face recall over a number of studies using meta-analysis. To maintain brevity of the paper this analysis is not included here but can be found in Appendix 1.

In brief, participants' free face recall elicited prior to composite construction was coded and analysed using a Microsoft Excel template provided by Neyeloff et al. (2012). Results showed that in those cases when Detailed CR was successful in increasing composite identifiability (cf. Minimal CR) an overall effect was found for face recall, thereby providing evidence for the idea that the CR effect is driven by an increase in memory recall of the target face.

Thus, detailed recall of the environment allows constructors to achieve a better memory of a target face, the knock-on effect of which is for them to achieve a more accurate composite. However, an alternative (and not necessarily mutually-exclusive) explanation is that Detailed CR facilitates face recognition, such as observed when context is reinstated

using cues associated with the target (e.g., Shapiro & Penrod, 1986). A simple test of this theory would be to reverse the order of interviewing mnemonics: with free recall of the face carried out first followed by Detailed CR. If Detailed CR improves face recognition, then correct naming of composites would be expected to increase (cf. Minimal CR); yet, if the mechanism is mediated mainly by improvement of face recall, then no such benefit should be observed. We investigated this proposal in the following experiment using the same design as Experiment 4 but with the order of these two mnemonics reversed.

EXPERIMENT 5: Timing of detailed recall of the environmental context

Stage 1: Composite construction

Design

The design was between-subjects with two experimental factors: 2 (CR: Minimal, Detailed) $\times$ 2 (Face construction: Face-to-face, Self-administered). For Detailed CR, participants freely recalled the target face first and then the environmental context; for other participants, only free recall of the face was requested. As Experiment 4, participants entered the unfamiliar room used for encoding prior to the experimenter, to provide suitable opportunity for encoding of context; also, for Detailed CR, participants again recalled the physical, environmental (but not psychological) context.

Participants

Thirty-two (14 males, 18 females; $M_{age} = 24.9$ years, $SD_{age} = 9.5$ years) students volunteered or took part in return for course credits. Participants were recruited on the basis that they were unfamiliar with the target faces.

Materials

Targets were eight photographs of characters from the ITV Coronation Street soap (Peter Barlow, Carla Connor, Tyrone Dobbs, Tracey McDonald, David Platt, Kirk Sutherland, Leanne Tilsley and Sally Webster). Photographs were of the same standard as in the previous experiments and were printed likewise.

Procedure

Face construction was the same as in Experiment 4, except for the following differences. To keep the construction procedure closely aligned with the first three experiments, participants (randomly) assigned to self-administered construction were asked to verbally describe the target face as well as the environmental context, with the experimenter writing down recall (cf., Experiment 4, with participants writing down this information). The encoding environment was a Multi-Faith Centre on the UCLan (Preston) campus, for approximately half of the participants, and a café in a local town; both rooms were rich in environmental cues and unfamiliar to participants. Face construction requested all participants to “think back” to when the face had been seen the previous day. For Detailed CR, the order of instructions was reversed: here, free recall of face was requested first followed by free recall of environment; Minimal CR, as before, only involved free recall of the face.

Stage 2: Composite naming

Design, materials and procedure

The two factorial design was the same as Stage 1, and materials were 32 composites and eight target photographs, printed as before. The procedure was also the same as before, composite naming and then target naming, but was extended in one way. It was anticipated that if the benefit of Detailed CR was driven in part by improvement in face recall, the observed effect would be weaker. As such, statistical power was increased by asking participants to name their assigned set of composites for a second time, after having seen the target photographs; this is another commonly used procedure in composite research (e.g., Frowd et al., 2007). For this second presentation of composites, which we refer to as “cued” naming (cf. “spontaneous” naming, first presentation), images were presented again in the same (random) order.

Participants

Thirty-six (14 males and 22 females; $M_{age} = 28.1$ years, $SD_{age} = 13.3$ years) participants were recruited via opportunity sampling on the UCLan campus on the basis of being familiar with the targets.

Table 5. Percentage of EvoFITs correctly named by Context Reinstatement, Face Construction and Method of composite naming.

	Context Reinstatement	
	Minimal	Detailed
<i>Spontaneous Naming</i>		
<i>Face Construction*</i>		
Face-to-face	43.1 (14.1)	37.5 (8.8)
Self-administered	26.4 (9.8)	33.3 (24.2)
<i>Cued Naming</i>		
<i>Face Construction*</i>		
Face-to-face	94.4 (11.0)	90.3 (13.7)
Self-administered	72.2 (16.3)	66.7 (25.8)

Note: * Significant main effect of Face Construction, $p < .005$. In parentheses are (by-participant) SD values.

Results

All participants correctly named all eight target photographs, and so familiarity with the target set was at 100%. Descriptive statistics (Table 5) revealed that, for the initial Spontaneous naming task, correct naming scores reduced slightly under Detailed (cf. Minimal) CR for face-to-face construction, but the reverse trend was observed for self-administered construction. There was a similar outcome by CR, albeit with a higher level of naming, for Cued naming.

MANOVA, incorporating both Spontaneous and Cued naming, revealed a significant main effect of Face construction [$F(2, 31) = 7.88$, $p = .002$, $\eta_p^2 = .34$], as composites were correctly named higher using face-to-face ($M = 66.3\%$, $SD = 12.0$) than self-administered construction ($M = 51.4\%$, $SD = 19.4$); here, Pillai's Trace is reported, based on Levene's Test of Equality of Error Variances ($p < .05$). There was no significant effect of either CR [$F(2, 31) = 0.61$, $p = .55$, $\eta_p^2 = .04$] or the interaction between CR and Face construction [$F(2, 31) = 0.39$, $p = .68$, $\eta_p^2 = .02$].

Discussion

Experiment 5 revealed that Detailed CR was not effective at increasing correct naming of composites when used after participants had described the target face, and that this null effect emerged irrespective of whether the face was constructed using the conventional face-to-face procedure, or when self-administered. The result suggests that correct composite naming improves for Detailed (cf. Minimal) CR due to increases in face recall rather than face recognition.

Mirroring Experiment 4, it was also found that composites were constructed more effectively when participants worked with the experimenter than alone.

General discussion

The current work involved five experiments that examined environmental context techniques as retrieval cues to potentially improve a person's ability to create a facial composite from memory. Asking constructors to recall both the visual environment and their psychological context where a target face had been seen was successful in increasing the effectiveness of facial composites produced from a typical holistic system (EvoFIT) and a typical feature system (PRO-fit). When omitting recall of psychological context (Experiment 4), CR techniques remained effective. Environmental context, however, only seemed to be valuable in aiding memory if participants ("witnesses") paid sufficient attention to it at encoding (Experiment 3). Whilst CR procedures were effective in combination with a further interviewing technique, H-CI (Experiment 2), this was not the case for an extensive recall technique, one where additional retrieval of the environment using cued-recall questions was encouraged (Experiment 3).

The first two experiments involved a holistic system (EvoFIT) and a feature system (PRO-fit). The main reason for including both types of system was to investigate whether CR interviewing techniques could be successful for both, rather than comparing the systems *per se*. Other research has focussed on the effectiveness of systems (e.g., see Frowd et al., 2015 for meta-analysis), and here we find the same overall outcome, that a holistic system is more effective than a feature system.

To our knowledge, limited research (Davies & Milne, 1985; Ness & Bruce, 2006) has applied CR techniques for the purpose of enhancing composites, and our results mirror, as well as extend, previous findings. In line with the Encoding Specificity theory (Tulving & Thomson, 1973), as participants in Detailed CR were first involved in recalling the environmental and internal context, it seems as though this information acted as associative retrieval cues, facilitating access to facial memory, thereby facilitating construction of the target face. More-specifically, the underlying mechanism for the effectiveness of Detailed CR seems to be an increase in memory recall of facial

features of the target rather than improved face recognition (Experiment 5; meta-analyses of the face recall data, Appendix 1).

Physical CR did not have a facilitative effect. The interviewing procedure prior to composite construction did not differ between Minimal and Physical conditions – that is, participants only actively recalled the *target face*, but did not engage in intensive memory recall of the environment (cf. Detailed CR). Therefore, the environmental context as a physical retrieval cue in itself does not appear to be strong enough to facilitate memory. Similarly, Davies and Milne's (1985) results also indicate that Physical CR was not as effective as Mental (Detailed) CR.³ Since Physical CR would pose practical and ethical problems for policing – even though there are potential ways to overcome this issue (see Ness & Bruce, 2006) – we propose that Detailed CR is a more convenient procedure to implement, and may be less traumatic for victims (i.e., as there would not be the need to return to the scene of crime).

The effectiveness of Detailed CR, however, is dependent upon the extent of the visual encoding of the environment; in other words, it is only effective in increasing composite naming when participants paid sufficient attention to the environment (Experiment 3). As the environment in the laboratory is perhaps more “disconnected” from the target face than it would be in real life, one would theorize that the issue regarding lack of attention to the environment would not occur for real eyewitnesses. Also, efforts to ensure a more natural encoding of the environment (in Experiment 4) showed that participants had, without prompting, encoded the environment to a greater extent, a situation that is closer to real life. In a real-life crime, it is thought that a witness or victim is unlikely to intentionally encode the environment around them, and it is encouraging to observe that incidental encoding has subsequent benefit to face construction. Perhaps more importantly, even when Detailed CR was not effective, it did not reliably reduce correct naming of composites, nor alter incorrect naming, and hence there is no apparent disadvantage in using this technique in the real world.

Results from Experiment 4 also indicated that constructor's psychological state does not need to be recalled (cf. Experiments 1–3) for Detailed CR to be efficacious: recall of the visual environment alone

was sufficient to facilitate face-to-face and self-administered EvoFIT face construction. This has important practical implications since witnesses and victims (esp. of serious crimes such as assault) are likely to find it uncomfortable to recall their psychological state, which in turn is likely to inhibit face construction (Davies, 2009). It is unclear as to why recall of the internal state is not necessary to facilitate access to the memory of the face. In the current work, participants generally spent more time recalling the exterior environmental rather than their internal context, and so this may be indicative that the environment is more helpful in aiding memory access. It may also be the case that, although participants were not asked to verbalize their internal context at the time of face encoding, the act of recalling the environment may have automatically triggered memory of mood and feelings at the time, as would be implied with the Encoding Specificity principle (Tulving & Thomson, 1973). We acknowledge that we have not considered the effect of recalling the constructor's psychological state alone, and it would seem sensible from a theoretical perspective to investigate this possibility in future work (although based on the above result from Detailed CR, it is unlikely that recall of internal state would be effective).

Detailed CR can also be effective in conjunction with H-CI (Experiment 2), leading to the best-named composites compared to other conditions. Both of these procedures may be effective in combination by guiding witness attention to different aspects of the face: the H-CI shifts focus towards the central part of the face, whilst Detailed CR provides better memory for facial features. If memory is improved for both the whole face and individual features, then it would seem reasonable that the two techniques combined could improve face construction. It is now established that shifting a witness's focus of attention to the central part of the face aids composite construction (e.g., Frowd et al., 2008). On the other hand, improved memory for facial features is likely to aid discrimination between features – such as, eyes and brows. As the region around the eyes plays a central role for familiar face recognition (O'Donnell & Bruce, 2001), relevant here to composite naming, it seems plausible that better memory for this area should assist in constructing a better likeness, which in turn would increase subsequent identification (Ellis et al., 1980).

In Experiment 3, Detailed CR was considered along with another interviewing mnemonic to potentially facilitate memory further, extensive retrieval (Fisher & Geiselman, 1992), in this case extensive recall of the environment. Rather than stopping after one, initial memory search, it is thought that multiple retrieval attempts elicit greater overall recall. Our data, however, suggest that further recall was no more effective for improving composite effectiveness compared to Detailed CR, and it is not entirely clear as to why. Constructors in this condition each recalled further detail when probed, showing that the technique elicited more information compared to the initial, free recall. It is possible that this additional information simply did not facilitate facial memory any further, maybe because the information was less relevant than free recall, or perhaps participants had already visualized this information during free recall, but did not verbalize it to the experimenter, despite being asked to recall as much information as possible. Our findings are somewhat in line with results from Campos and Alonso-Quecuty (1999) who found that four multiple retrieval attempts (“try again”) were not as effective in increasing correct recall as other retrieval strategies – CR, change perspective, recall in a different order, and other techniques. Although participants in our Experiment 3 were not specifically asked to “try again” (as in Campos & Alonso-Quecuty, 1999), our cued-based questions may not be as effective as other retrieval mnemonic strategies. If this is the case, one straightforward way to trigger new information may be to request recall in a different spatial order, perhaps starting from the location of last item remembered. Future work could investigate if such a retrieval strategy (combined with CR), or others such as recalling from another person’s perspective, might facilitate face construction further.

Conclusion

This paper is the first demonstration of an advantage of recalling contextual cues for forensic face construction using modern composite systems in a realistic experimental design (incl. target face unfamiliar at encoding, nominal 24 hr delay to construction, and naming of ensuing composites). This mnemonic of the Cognitive Interview was also effective in combination with an H-CI, with focus on the perceived

character of the target face. These two procedures are straightforward, take little time to administer, and if used with eyewitnesses should increase visual identification of offenders. In fact, forensic practitioners in the UK and overseas now regularly use Detailed CR (for recalling scene of crime but not psychological state) with witnesses and victims when constructing a composite with EvoFIT (Frowd et al., 2019). On a final note, it is worth mentioning that, given the similarity of procedures with modern facial composite systems, it seems likely that our results would generalize to other feature and holistic systems such as EFIT-V (EFIT-6), ID, FACES 4.0 and Identikit 2000.

Notes

1. Note that “*Detailed CR*” is equivalent to the type of elaborative Mental Context Reinstatement as used in other studies of memory (but not facial composite construction). Since a simple “think back” type (Minimal) CR is already in use with face construction, we draw a distinction between Minimal and Detailed CR.
2. Please note that the term “Cognitive Interview” within facial composite research differs from the full “Cognitive Interview” (Fisher & Geiselman, 1992). In relation to composites, the CI is a more concise version and usually only includes mnemonics for rapport building, visualisation and free recall of the target face (see Frowd, 2011 for an in-depth review of the CI for facial-composite construction). In the current paper, we refer to CI (and in Experiment 2 to “Holistic-Cognitive Interview”) as used within composite research.
3. Detailed CR is equivalent to Mental CR from Davies and Milne’s (1985) experiment (see current paper, Footnote¹).
4. The Supplementary experiment was conducted by the current authors. It followed the same Design and Procedure as Experiment 2, but found no significant effect of Detailed CR. This null effect is hypothesised to be due to participants not having suitably encoded the environment, in line with findings from Experiment 3. It is not included in the main manuscript to maintain brevity of the paper.

Disclosure statement

No potential conflict of interest was reported by the author(s).

ORCID

Cristina Fodarella  <http://orcid.org/0000-0001-5551-3450>

References

- Alison, L., Doran, B., Long, M. L., Power, N., & Humphrey, A. (2013). The effects of subjective time pressure and individual differences on hypotheses generation and action prioritization in police investigations. *Journal of Experimental Psychology: Applied*, 19(1), 83–93. <https://doi.org/10.1037/a0032148>
- Anderson, J. A., & Bower, G. H. (1972). Recognition and retrieval processes in free recall. *Psychological Review*, 79(2), 97–123. <https://doi.org/10.1037/h0033773>
- Barak, O., Vakil, E., & Levy, D. A. (2013). Environmental context effects on episodic memory are dependent on retrieval mode and modulated by neuropsychological status. *Quarterly Journal of Experimental Psychology*, 66(10), 2008–2022. <https://doi.org/10.1080/17470218.2013.772647>
- Bhardwaj, K., & Hole, G. (2020). Effect of racial bias on composite construction. *Applied Cognitive Psychology*, 34(3), 616–627. <https://doi.org/10.1002/acp.3655>
- Bower, G. H. (1981). Mood and memory. *American Psychologist*, 36(2), 129–148. <https://doi.org/10.1037/0003-066X.36.2.129>
- Bower, G. H., & Cohen, P. R. (1982). Emotional influences in memory and thinking: Data and theory. In M. S. Clark, & S. T. Fiske (Eds.), *Affect and cognition* (pp. 263–289). Erlbaum.
- Brown, J. M. (2003). Eyewitness memory for arousing events: Putting things into context. *Applied Cognitive Psychology: The Official Journal of the Society for Applied Research in Memory and Cognition*, 17(1), 93–106. <https://doi.org/10.1002/acp.848>
- Campeanu, S., Craik, F. I., Backer, K. C., & Alain, C. (2014). Voice reinstatement modulates neural indices of continuous word recognition. *Neuropsychologia*, 62, 233–244. <https://doi.org/10.1016/j.neuropsychologia.2014.07.022>
- Campos, L., & Alonso-Quecuty, M. L. (1999). The cognitive interview: Much more than simply “try again”. *Psychology, Crime and Law*, 5(1–2), 47–59. <https://doi.org/10.1080/10683169908414993>
- Clark, M. S., Milberg, S., & Ross, J. (1983). Arousal cues arousal-related material in memory: Implications for understanding effects of mood on memory. *Journal of Memory and Language*, 22(6), 633–649. [https://doi.org/10.1016/S0022-5371\(83\)90375-4](https://doi.org/10.1016/S0022-5371(83)90375-4)
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Lawrence Erlbaum Associates.
- Cutler, B. L., Penrod, S. D., & Martens, T. K. (1987). Improving the reliability of eyewitness identifications: Putting context into context. *Journal of Psychology*, 72(2), 629–637. <https://doi.org/10.1037/0021-9010.72.4.629>
- Dando, C. J. (2013). Drawing to remember: External support of older adults’ eyewitness performance. *PloS One*, 8(7), Article e69937. <https://doi.org/10.1371/journal.pone.0069937>
- Dando, C. J., Wilcock, R., Milne, R., & Henry, L. (2009). A modified cognitive interview procedure for frontline police investigators. *Applied Cognitive Psychology*, 23(5), 698–716. <https://doi.org/10.1002/acp.1501>
- Davies, G. M. (1983). Forensic face recall: The role of visual and verbal information. In S. M. A. Lloyd-Bostock, & B. R. Clifford (Eds.), *Evaluating witness evidence* (pp. 103–123). Wiley.
- Davies, S. (2009). The effects of stress levels on the construction of facial composites. *Diffusion*, 2. <http://atp.uclan.ac.uk/buddypress/diffusion/?p=1763>
- Davies, G. M., & Milne, A. (1985). Eyewitness composite production. *Criminal Justice and Behavior*, 12(2), 209–220. <https://doi.org/10.1177/0093854885012002004>
- Davis, M. R., McMahon, M., & Greenwood, K. M. (2005). The efficacy of mnemonic components of the cognitive interview: Towards a shortened variant for time-critical investigations. *Applied Cognitive Psychology*, 19(1), 75–93. <https://doi.org/10.1002/acp.1048>
- Eich, J. E. (1980). The cue-dependent nature of state-dependent retrieval. *Memory and Cognition*, 8(2), 157–173. <https://doi.org/10.3758/BF03213419>
- Eich, E., & Metcalfe, J. (1989). Mood dependent memory for internal versus external events. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15(3), 443–455. <https://doi.org/10.1037/0278-7393.15.3.443>
- Eich, J. E., Weingartner, H., Stillman, R. C., & Gillin, J. C. (1975). State-dependent accessibility of retrieval cues in the retention of a categorized list. *Journal of Memory and Language*, 14(4), 408–417. [https://doi.org/10.1016/S0022-5371\(75\)80020-X](https://doi.org/10.1016/S0022-5371(75)80020-X)
- Ellis, H. D., Shepherd, J. W., & Davies, G. M. (1980). The deterioration of verbal descriptions of faces over different delay intervals. *Journal of Police Science and Administration*, 8(1), 101–106.
- Emmett, D., Clifford, B. R., & Gwyer, P. (2003). An investigation of the interaction between cognitive style and context reinstatement on memory performance of eyewitnesses. *Personality and Individual Differences*, 34(8), 1495–1508. [https://doi.org/10.1016/S0191-8869\(02\)00131-9](https://doi.org/10.1016/S0191-8869(02)00131-9)
- Evans, J. R., Marcon, J. L., & Meissner, C. A. (2009). Cross-racial lineup identification: Assessing the potential benefits of context reinstatement. *Psychology, Crime and Law*, 15(1), 19–28. <https://doi.org/10.1080/10683160802047030>
- Fisher, R., & Geiselman, R. (1992). *Memory enhancing techniques for investigative interviewing*. Charles Thomas Publishers.
- Fodarella, C., & Frowd, C. D. (2013, September). Accuracy of relational and featural information in facial-composite images. In 2013 Fourth International Conference on Emerging Security Technologies (pp. 16–20). IEEE.
- Fodarella, C., Frowd, C. D., Warwick, K., Hepton, G., Stone, K., Date, L., & Heard, P. (2017). Adjusting the focus of attention: Helping witnesses to evolve a more identifiable composite. *Forensic Research and Criminology International Journal*, 5(1), Article 00143. <https://doi.org/10.15406/frcij.2017.05.00143>
- Fodarella, C., Kuivaniemi-Smith, H. J., Gawrylowicz, J., & Frowd, C. D. (2015). Detailed procedures for forensic face construction. *Journal of Forensic Practice*, 17(4), 259–270. <https://doi.org/10.1108/JFP-10-2014-0033>
- Frowd, C. D. (2011). Eyewitnesses and the use and application of cognitive theory. In G. Davey (Ed.), *Introduction to applied psychology* (pp. 267–289). BPS Wiley-Blackwell.

- Frowd, C. D. (2015). Facial composites and techniques to improve image recognisability. In T. Valentine, & J. Davis (Eds.), *Forensic facial identification: Theory and practice of identification from eyewitnesses, composites and cctv* (pp. 43–70). Wiley-Blackwell. ISBN: 978-1-118-46911-8.
- Frowd, C. D. (2017). Facial composite systems: Production of an identifiable face. In M. Bindemann & A. Megreya (Eds.) *Face processing: Systems, disorders and cultural differences* (pp. 55–86). Nova Science: New York.
- Frowd, C. D. (in press). Forensic facial composites. In M. Togli et al. (Ed.), *Methods, measures and theories in eyewitness identification*. Taylor and Francis.
- Frowd, C. D., Bruce, V., Ross, D., McIntyre, A., & Hancock, P. J. B. (2007). An application of caricature: How to improve the recognition of facial composites. *Visual Cognition*, 15(8), 1–31. <https://doi.org/10.1080/13506280601058951>
- Frowd, C. D., Bruce, V., Smith, A., & Hancock, P. J. B. (2008). Improving the quality of facial composites using a holistic cognitive interview. *Journal of Experimental Psychology: Applied*, 14(3), 276–287. <https://doi.org/10.1037/1076-898X.14.3.276>
- Frowd, C. D., Carson, D., Ness, H., Richardson, J., Morrison, L., McLanaghan, S., & Hancock, P. J. B. (2005). A forensically valid comparison of facial composite systems. *Psychology, Crime and Law*, 11(1), 33–52. <https://doi.org/10.1080/10683160310001634313>
- Frowd, C. D., Erickson, W. B., Lampinen, J. M., Skelton, F. C., McIntyre, A. H., & Hancock, P. J. B. (2015). A decade of evolving composite techniques: Regression- and meta-analysis. *Journal of Forensic Practice*, 17(4), 319–334. <https://doi.org/10.1108/JFP-08-2014-0025>
- Frowd, C. D., Hancock, P. J. B., Bruce, V., McIntyre, A. H., Pitchford, M., Atkins, R., Webster, A., Pollard, J., Hunt, B., Price, E., & Morgan, S. (2010, September). Giving crime the 'evo': Catching criminals using EvoFIT facial composites. In *2010 International Conference on Emerging Security Technologies* (pp. 36–43). IEEE.
- Frowd, C. D., Jones, S., Fodarella, C., Skelton, F. C., Fields, S., Williams, A., Marsh, J., Thorley, R., Nelson, L., Greenwood, L., Date, L., Kearley, K., McIntyre, A., & Hancock, P. J. B. (2014). Configural and featural information in facial-composite images. *Science and Justice*, 54(3), 215–227. <https://doi.org/10.1016/j.scijus.2013.11.001>
- Frowd, C. D., Nelson, L., Skelton, F. C., Noyce, R., Atkins, R., Heard, P., Morgan, D., Fields, S., Henry, J., McIntyre, A., & Hancock, P. J. B. (2012). Interviewing techniques for Darwinian facial composite systems. *Applied Cognitive Psychology*, 26(4), 576–584. <https://doi.org/10.1002/acp.2829>
- Frowd, C. D., Pitchford, M., Skelton, F. C., Petkovic, A., Prosser, C., & Coates, B. (2012). Catching even more offenders with EvoFIT facial composites. In *2012 Third International Conference on Emerging Security Technologies* (pp. 20–26). IEEE. <https://doi.org/10.1109/EST.2012.26>
- Frowd, C. D., Portch, E., Killeen, A., Mullen, L., Martin, A. J., & Hancock, P. J. B. (2019, July). EvoFIT facial composite images: A detailed assessment of impact on forensic practitioners, police investigators, victims, witnesses, offenders and the media. In *2019 Eighth International Conference on Emerging Security Technologies* (pp. 1–7). IEEE.
- Frowd, C. D., Skelton, F., Atherton, C., Pitchford, M., Hepton, G., Holden, L., McIntyre, A. H., & Hancock, P. J. B. (2012). Recovering faces from memory: The distracting influence of external facial features. *Journal of Experimental Psychology: Applied*, 18(2), 224–238. <https://doi.org/10.1037/a0027393>
- Gawrylowicz, J., Gabbert, F., Carson, D., Lindsay, W. R., & Hancock, P. J. B. (2012). Holistic versus featural facial composite systems for people with mild intellectual disabilities. *Applied Cognitive Psychology*, 26(5), 716–720. <https://doi.org/10.1002/acp.2850>
- Geiselman, R. E., Fisher, R. P., MacKinnon, D. P., & Holland, H. L. (1985). Eyewitness memory enhancement in the police interview: Cognitive retrieval mnemonics versus hypnosis. *Journal of Applied Psychology*, 70(2), 401–412. <https://doi.org/10.1037/0021-9010.70.2.401>
- Glenberg, A. M. (1997). What memory is for. *Behavioral and Brain Sciences*, 20(1), 1–19. <https://doi.org/10.1017/S0140525X97000010>
- Godden, D. R., & Baddeley, A. D. (1975). Context-dependent memory in two natural environments: Land and underwater. *British Journal of Psychology*, 66(3), 325–331. <https://doi.org/10.1111/j.2044-8295.1975.tb01468.x>
- Goldstein, A. G., & Mackenberg, E. J. (1966). Recognition of human faces from isolated facial features: A developmental study. *Psychonomic Science*, 6(4), 149–150. <https://doi.org/10.3758/BF03328001>
- Hammond, L., Wagstaff, G. F., & Cole, J. (2006). Facilitating eyewitness memory in adults and children with context reinstatement and focused meditation. *Journal of Investigative Psychology and Offender Profiling*, 3(2), 117–130. <https://doi.org/10.1002/jip.47>
- Hanczakowski, M., Zawadzka, K., & Coote, L. (2014). Context reinstatement in recognition: Memory and beyond. *Journal of Memory and Language*, 72, 85–97. <https://doi.org/10.1016/j.jml.2014.01.001>
- Hasel, L. E., & Wells, G. L. (2007). Catching the bad guy: Morphing composite faces helps. *Law and Human Behavior*, 31(2), 193–207. <https://doi.org/10.1007/s10979-006-9007-2>
- Hockley, W. E. (2008). The effects of environmental context on recognition memory and claims of remembering. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 34(6), 1412–1429. <https://doi.org/10.1037/a0013016>
- Kehn, A., Renken, M. D., Gray, J. M., & Nunez, N. L. (2014). Developmental trends in the process of constructing own- and other-race facial composites. *The Journal of Psychology*, 148(3), 287–304. <https://doi.org/10.1080/00223980.2013.794122>
- Koen, J. D., Aly, M., Wang, W. C., & Yonelinas, A. P. (2013). Examining the causes of memory strength variability: Recollection, attention failure, or encoding variability? *Journal of Experimental Psychology: Learning, Memory, and*

- Cognition*, 39(6), 1726–1741. <https://doi.org/10.1037/a0033671>
- Laughery, K. R., Duval, C., & Wogalter, M. S. (1986). Dynamics of facial recall. In A. W. Young (Ed.), *Aspects of face processing* (pp. 373–387). Springer Netherlands.
- Levy, D. A., Rabinyan, E., & Vakil, E. (2008). Forgotten but not gone: Context effects on recognition do not require explicit memory for context. *The Quarterly Journal of Experimental Psychology*, 61(11), 1620–1628. <https://doi.org/10.1080/17470210802134767>
- Macken, W. J. (2002). Environmental context and recognition: The role of recollection and familiarity. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 28(1), 153–161. <https://doi.org/10.1037/0278-7393.28.1.153>
- Malpass, R. S., & Devine, P. G. (1981). Eyewitness identification: Lineup instructions and the absence of the offender. *Journal of Applied Psychology*, 66(4), 482–489. <https://doi.org/10.1037/0021-9010.66.4.482>
- Martin, A. J., Hancock, P. J. B., & Frowd, C. D. (2017, September). Breathe, relax and remember: An investigation into how focused breathing can improve identification of EvoFIT facial composites. In *2017 Seventh international Conference on Emerging security Technologies* (pp. 79–84). IEEE.
- Martin, A. J., Hancock, P. J. B., Frowd, C. D., Heard, P., Gaskin, E., Ford, C., & Hewitt, T. (2018, August). EvoFIT composite face construction via practitioner interviewing and a witness-administered protocol. In *2018 NASA / ESA Conference on Adaptive Hardware and Systems* (pp. 311–316). IEEE.
- McQuiston-Surrett, D., & Topp, L. D. (2008). Externalizing visual images: Examining the accuracy of facial descriptions vs. composites as a function of the own-race bias. *Experimental Psychology*, 55(3), 195–202. <https://doi.org/10.1027/1618-3169.55.3.195>
- Memon, A., & Bruce, V. (1995). Context effects in episodic studies of verbal and facial memory: A review. *Current Psychological Research and Reviews*, 4(4), 349–369. <https://doi.org/10.1007/BF02686589>
- Memon, A., Meissner, C. A., & Fraser, J. (2010). The cognitive interview: A meta-analytic review and study space analysis of the past 25 years. *Psychology, Public Policy, and Law*, 16(4), 340–372. <https://doi.org/10.1037/a0020518>
- Milne, R., & Bull, R. (2002). Back to basics: A componential analysis of the original cognitive interview mnemonics with three age groups. *Applied Cognitive Psychology*, 16(7), 743–753. <https://doi.org/10.1002/acp.825>
- Murnane, K., & Phelps, M. P. (1993). A global activation approach to the effect of changes in environmental context on recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 19(4), 882–894. <https://doi.org/10.1037/0278-7393.19.4.882>
- Ness, H., & Bruce, V. (2006, June). An investigation into the use of CCTV footage to improve likeness in facial composites. In *European Association of Psychology and Law Conference*, Liverpool, UK.
- Neyeloff, J. L., Fuchs, S. C., & Moreira, L. B. (2012). Meta-analyses and forest plots using a microsoft excel spreadsheet: Step-by-step guide focusing on descriptive data analysis. *BMC Research Notes*, 5(1), 1–6. <https://doi.org/10.1186/1756-0500-5-52>
- O'Donnell, C., & Bruce, V. (2001). Familiarisation with faces selectively enhances sensitivity to changes made to the eyes. *Perception*, 30(6), 755–764. <https://doi.org/10.1068/p3027>
- Pellicano, E., Rhodes, G., & Peters, M. (2006). Are pre-schoolers sensitive to configural information in faces? *Developmental Science*, 9(3), 270–277. <https://doi.org/10.1111/j.1467-7687.2006.00489.x>
- Rainis, N. (2001). Semantic contexts and face recognition. *Applied Cognitive Psychology*, 15(2), 173–186. [https://doi.org/10.1002/1099-0720\(200103/04\)15:2<173::AID-ACP695>3.0.CO;2-Q](https://doi.org/10.1002/1099-0720(200103/04)15:2<173::AID-ACP695>3.0.CO;2-Q)
- Shapiro, P. N., & Penrod, S. (1986). Meta-analysis of facial identification studies. *Psychological Bulletin*, 100(2), 139–156. <https://doi.org/10.1037/0033-2909.100.2.139>
- Smith, S. M. (1979). Remembering in and out of context. *Journal of Experimental Psychology: Human Learning and Memory*, 5(5), 460–471. <https://doi.org/10.1037/0278-7393.5.5.460>
- Smith, S. M. (2013). Effects of environmental context on human memory. In T. J. Perfect, & D. S. Lindsay (Eds.) *The SAGE handbook of applied memory* (vol. 2, pp. 162–182). SAGE Publications Ltd.
- Smith, S. M., & Manzano, I. (2010). Video context-dependent recall. *Behavior Research Methods*, 42(1), 292–301. <https://doi.org/10.3758/BRM.42.1.292>
- Smith, S. M., & Vela, E. (2001). Environmental context-dependent memory: A review and meta-analysis. *Psychonomic Bulletin and Review*, 8(2), 203–220. <https://doi.org/10.3758/BF03196157>
- Thomson, D. M., Robertson, S. L., & Vogt, R. (1982). Person recognition: The effect of context. *Human Learning*, 1(1), 37–54.
- Tulving, E., & Thomson, D. M. (1973). Encoding specificity and retrieval processes in episodic memory. *Psychological Review*, 80(5), 352–373. <https://doi.org/10.1037/h0020071>
- Wagstaff, G. F. (1982, July). Context effects in eyewitness reports [Paper presented]. *Law and Psychology Conference*, Swansea, Wales.
- Wagstaff, G. F., Cole, J., Wheatcroft, J., Marshall, M., & Barsby, I. (2007). A componential approach to hypnotic memory facilitation: Focused meditation, context reinstatement and eye movements. *Contemporary Hypnosis*, 24(3), 97–108. <https://doi.org/10.1002/ch.334>
- Wagstaff, G. F., Wheatcroft, J. M., Caddick, A. M., Kirby, L. J., & Lamont, E. (2011). Enhancing witness memory with techniques derived from hypnotic investigative interviewing: Focused meditation, eye-closure, and context reinstatement. *International Journal of Clinical and Experimental Hypnosis*, 59(2), 146–164. <https://doi.org/10.1080/00207144.2011.546180>
- Wells, G. L., Memon, A., & Penrod, S. D. (2007). Eyewitness evidence: Improving its probative value. *Psychological Science in the Public Interest*, 7(2), 45–75. <https://doi.org/10.1111/j.1529-1006.2006.00027.x>

Wong, C. K., & Read, J. D. (2011). Positive and negative effects of physical context reinstatement on eyewitness recall and identification. *Applied Cognitive Psychology*, 25(1), 2–11. <https://doi.org/10.1002/acp.1605>

Appendix 1

Meta-analyses of face recall

Method

Studies and procedure

Eleven comparisons between Minimal versus Detailed CR were available for meta-analyses of face recall, derived from Experiments 1–3, and an unpublished Supplementary experiment.⁴ Comparisons include only incidental encoding of the environment (i.e., cases where attention had been directed to the café in Experiment 3 were omitted). Cohen's *d* was used as the measure of effect size in each comparison along with its standard error (SE).

For the analysis of face recall, participants' free face recall elicited prior to composite construction were coded by assigning a value of 1 for each unit of information (UOI) recalled. For example, "small, brown eyes" would be counted as two UOI, and "eyebrows were far apart, low and quite straight" as three. Information regarding details other than the face (e.g., clothing, jewellery and shoulders) were excluded, whilst subjective information about the face, such as "pleasant, good-looking face" (two UOI) or "quite a friendly face" (one UOI), were included; information relating to either of these categories was recalled by around 1 in 4 participants. Using this scoring procedure, total face recall was calculated for each facial composite across experimental conditions. To ensure consistency, the same two experimenters coded recall. Both coders were blind to the experimental conditions under which composites had been constructed, and participant recall was presented in a random order (for both experimental conditions and target identity). After coding, scores were compared and differences resolved by discussion; this occurred only on two occasions, and thus inter-rater reliability emerged at 100%.

Three meta-analyses were conducted in Microsoft Excel using a template provided by Neyeloff et al. (2012). The initial analysis included the seven comparisons in which Detailed CR reliably increased correct naming. The second analysis included data from Experiment 3, and the third analysis also included data from the Supplementary experiment. Mainly due to low statistical power, results of face recall are presented overall (cf. by composite system). It was expected that this information would only increase under Detailed CR when the meta-analysis included experiments in which the manipulation had been successful – that is, when correct naming had increased significantly. The effect was likely to decrease when including additional data (i.e., in the second and third meta-analysis).

Table A1. Mean (and SD) face recall and Cohen's *d* by CR across experiments.

Study	Context Reinstatement		Cohen's <i>d</i>
	Detailed CR	Minimal CR	
<i>Experiment 1</i>			
EvoFIT	9.4 (2.6)	7.4 (1.8)	0.9
PRO-fit	11.3 (4.7)	11.1 (3.5)	0.0
<i>Supplementary Experiment</i>			
EvoFIT	9.0 (4.1)	10.5 (1.9)	−0.5
EvoFIT (H-Cl)	11.6 (4.9)	13.6 (4.9)	−0.4
PRO-fit	11.3 (3.9)	10.5 (2.7)	0.2
PRO-fit (H-Cl)	12.9 (5.3)	13.6 (2.6)	−0.2
<i>Experiment 2</i>			
EvoFIT	13.4 (3.0)	11.9 (4.7)	0.4
EvoFIT (H-Cl)	12.8 (5.1)	10.8 (4.1)	0.4
PRO-fit	12.3 (2.7)	10.9 (3.0)	0.5
PRO-fit (H-Cl)	13.3 (2.7)	10.1 (3.8)	1.0
<i>Experiment 3</i>			
EvoFIT	9.5 (3.1)	10.7 (2.1)	−0.5

Results

Mean (and SD) face recall for each comparison across experiments was calculated along with Cohen's *d* (Table A1). A Random-effects model involving the six comparisons under which correct naming reliably improved for Detailed CR (cf. Minimal) indicated a significant increase in face recall [$Q(5) = 5.52$, $I^2 = 9.40$, $SE(d) = 0.14$, $95\%CI(d)$ (0.22, 0.77)], with a medium ES ($d = 0.49$). Including data from Experiment 3, when no advantage was observed by correct naming, the effect is now medium in size ($d = 0.36$) and the meta-analysis was marginally significant [$Q(6) = 5.97$, $I^2 = -0.42$, $SE(d) = 0.19$, $95\%CI(d)$ (–0.01, 0.73)]. When also including data from the Supplementary experiment, ES decreases even further ($d = 0.16$) and the analysis is no longer marginally significant [$Q(10) = 9.99$, $I^2 = -0.05$, $SE(d) = 0.15$, $95\%CI(d)$ (–0.14, 0.46)].

Discussion

Meta-analyses revealed that only in those cases where Detailed CR was successful in increasing composite identifiability (cf. Minimal CR) was an overall effect found for face recall. When including "unsuccessful" comparisons, the overall effect of face recall became weaker (incl. data from Experiment 3) and weaker again (also incl. data from the Supplementary experiment), thereby supporting the idea that the CR effect, when successful, is driven by an increase in memory recall of the target face.