



Mechanisms underlying approach-avoidance instruction effects on implicit evaluation: Results of a preregistered adversarial collaboration



Pieter Van Dessel^{a,*}, Bertram Gawronski^b, Colin Tucker Smith^c, Jan De Houwer^a

^a Ghent University, Belgium

^b University of Texas at Austin, USA

^c University of Florida, USA

HIGHLIGHTS

- We examined effects of approach-avoidance (AA) instructions on implicit evaluations.
- We tested predictions of a propositional and an associative self-anchoring account.
- Both approach and avoidance instructions influenced implicit evaluations.
- Effects were partially mediated by changes in implicit self-stimulus linking.
- Results fit best with a propositional explanation of AA instruction effects.

ARTICLE INFO

Article history:

Received 22 April 2016

Revised 7 September 2016

Accepted 11 October 2016

Available online 17 October 2016

Keywords:

Approach

Avoidance

Instruction

Implicit evaluation

Self-anchoring

Propositional theory

ABSTRACT

Previous research demonstrated that mere instructions to approach one stimulus and avoid another stimulus result in an implicit preference for the to-be-approached over the to-be-avoided stimulus. To investigate the mechanisms underlying approach-avoidance (AA) instruction effects, we tested predictions of a propositional account and an associative self-anchoring account in a preregistered adversarial collaboration. Consistent with the propositional account, Experiment 1 showed that avoidance instructions had a negative effect on implicit evaluations over and above the positive effect of approach instructions. Consistent with the associative self-anchoring account, Experiment 2 showed that changes in implicit self-stimulus linking mediated AA instruction effects on implicit evaluations. However, mediation was only partial, in that AA instructions showed a significant effect on implicit evaluations after controlling for implicit self-stimulus linking. Together, the results support the contribution of propositional processes to AA instruction effects; the results remain ambiguous regarding an additional contribution of associative self-anchoring.

© 2016 Elsevier Inc. All rights reserved.

It has been recognized for decades that behavior is shaped by likes and dislikes (Allport, 1935). Hence, understanding how these preferences are acquired is an important endeavor for psychological research. Interestingly, preferences sometimes arise as the result of performing specific behaviors (Olson & Stone, 2005). For example, previous research has shown that the repeated performance of approach and avoidance actions can cause changes in stimulus evaluations. When participants repeatedly approach one stimulus and avoid another stimulus, they typically develop a preference for the approached stimulus over the avoided stimulus (Laham, Kashima, Dix, Wheeler, & Levis, 2014). These approach-avoidance (AA) training effects have been observed

for a wide variety of stimuli, such as pictures of unfamiliar faces (Woud, Maas, Becker, & Rinck, 2013), racial groups (Kawakami, Phills, Steele, & Dovidio, 2007), alcoholic beverages (Wiers, Eberl, Rinck, Becker, & Lindenmeyer, 2011), unhealthy foods (Zogmaister, Perugini, & Richetin, in press), insects and spiders (Jones, Vilensky, Vasey, & Fazio, 2013), and contamination-related objects (Amir, Kuckertz, & Najmi, 2013).

In a recent set of studies, Van Dessel, De Houwer, Gast, and Smith (2015) obtained evidence that AA effects can also be observed as a result of mere instructions in the absence of actually performed actions. When participants were instructed to approach certain stimuli and avoid other stimuli, their evaluations of the to-be-approached stimuli were more positive than their evaluations of the to-be-avoided stimuli even though participants never actually performed the AA actions. Effects of AA instructions have been observed for novel non-words,

* Corresponding author at: Ghent University, Department of Experimental-Clinical and Health Psychology, Henri Dunantlaan 2, B-9000 Ghent, Belgium.
E-mail address: Pieter.vanDessel@UGent.be (P. Van Dessel).

fictitious social groups, and unfamiliar faces (Van Dessel, De Houwer, Roets, & Gast, 2016b). Importantly, these AA instruction effects were similar to the effects involving actual AA training in that both AA instructions and AA training influenced not only explicit (i.e., non-automatic) stimulus evaluations but also implicit (i.e., automatic) stimulus evaluations (Van Dessel, De Houwer, Gast, Smith, & De Schryver, 2016a).

Effects of AA instructions on implicit evaluation pose a challenge to a particular type of associative models that assume that (a) implicit evaluations reflect the automatic activation of associations in memory and (b) these associations are formed as the result of a slow-learning process that capitalizes on repeated co-occurrences, such as recurrent pairings of AA actions and stimuli (Rydell & McConnell, 2006; Smith & DeCoster, 2000). Yet, instruction-based AA effects are consistent with propositional models, which assume that implicit evaluations reflect the activation and generation of mental propositions about the relation between objects and events (e.g., De Houwer, 2009, 2014; Mitchell, De Houwer, & Lovibond, 2009). When participants are instructed to approach or avoid a stimulus, they might generate propositions about these stimulus-action relations, and these propositions can influence their implicit evaluations of the stimuli (Van Dessel et al., 2016a). For example, participants who learn that they will approach a stimulus may infer that this stimulus is positive, and participants who learn that they will avoid a stimulus may infer that this stimulus is negative. These inferences could arise because of the knowledge that positive objects are typically approached and negative objects are avoided (Schneirla, 1959). People may have learned this rule through previous experiences during which they approached liked stimuli and avoided disliked stimuli. Although this knowledge does not logically imply that approached things are good and avoided things are bad, people are known to be prone to affirm the consequent (i.e., conclude that A is true on the basis of the fact that A implies B and B is present). Thus, when participants infer that the to-be-approached stimulus is good and the to-be-avoided stimulus is bad, the (automatic) activation of this mental proposition could impact their implicit evaluations (De Houwer, 2014).

However, AA instruction effects on implicit evaluation are not necessarily incompatible with the view that implicit evaluations reflect the automatic activation of associations in memory (Gawronski & Bodenhausen, 2011). Some dual-process models, such as the associative-propositional evaluation (APE) model (Gawronski & Bodenhausen, 2006), postulate that mental associations can be formed as the result of propositional inferences. According to the APE model, any information that allows participants to entertain the proposition that a stimulus is positive or negative may instigate the proactive construction of new evaluative associations, which in turn may influence implicit evaluations. In line with this idea, changes in implicit evaluations have been observed when participants are provided with verbal information about the evaluative properties of a stimulus (Castelli, Zogmaister, Smith, & Arcuri, 2004; Cone & Ferguson, 2015; Gawronski, Walther, & Blank, 2005; Gregg, Seibt, & Banaji, 2006). Importantly, these models predict a specific pattern of mediation such that instruction effects on explicit evaluation should mediate effects on implicit evaluation (e.g., Gawronski & Walther, 2008; Peters & Gawronski, 2011a; Whitfield & Jordan, 2009; see Gawronski & Bodenhausen, 2006; Case 4).

Van Dessel et al. (2016a) recently performed two experiments that tested the mediating role of explicit evaluations in the effect of AA instructions on implicit evaluations. In these experiments, participants first received information about the evaluative traits of members of two fictitious social groups and were then given instructions to approach or avoid the names of members of these groups. The results showed that trait information eliminated the effects of AA instructions on explicit, but not implicit, evaluations. Statistical mediation analyses further showed that AA instructions had a direct effect on implicit evaluations that was not mediated by changes in explicit evaluations. These

findings contradict the idea that AA instructions influence implicit evaluations only if these instructions are considered a valid basis for evaluation and, hence, are incorporated in explicit evaluations (see Gawronski & Bodenhausen, 2006). Yet, the results are consistent with a propositional explanation of AA instruction effects and support the propositional model of evaluation which postulates that mental propositions, rather than associations, underlie implicit evaluation (De Houwer, 2014). Specifically, AA instructions might allow participants to consider the proposition that a to-be-approached stimulus is positive and a to-be-avoided stimulus is negative. A dissociation between implicit and explicit evaluation will arise when this proposition is judged to be invalid (and thus dismissed when making an explicit evaluation) but still automatically retrieved when the stimuli are implicitly evaluated.

Nevertheless, there is an important alternative explanation of AA instruction effects on implicit evaluation that is compatible with associative theories of implicit evaluation. Effects of AA instructions on implicit evaluation could arise as the result of associative self-anchoring, which involves the transfer of positive valence from the self to a stimulus associated with the self as the result of a newly formed association between the representation of the stimulus and the representation of the self (see Gawronski, Bodenhausen, & Becker, 2007). It is often assumed that approach behaviors are fundamentally related to pulling objects closer to the self (Förster, 2001), which may result in accentuated psychological closeness between approached stimuli and the self (Nussinson, Seibt, Häfner, & Strack, 2010). In line with this idea, it has been argued that the repeated performance of approach behavior in response to a stimulus allows for the formation of a mental association between the representation of the approached stimulus and the positive representation of the self (Kawakami, Steele, Cifa, Phills, & Dovidio, 2008; Phills, Kawakami, Tabi, Nadolny, & Inzlicht, 2011). Once such an association has been established, the positive valence of the self may spread to the approached stimulus, and thereby influence implicit evaluations of the stimulus (Gawronski et al., 2007). This associative transfer of valence is assumed to be driven by processes of spreading activation without requiring any kind of higher-order propositional processes (Gawronski, Strack, & Bodenhausen, 2009). Although many theories assume that the formation of new associations in memory is a slow, gradual process that requires repeated co-occurrences (e.g., Fazio, Sanbonmatsu, Powell, & Kardes, 1986; Smith & DeCoster, 2000), some researchers have rejected this idea and argued that sufficiently strong associations can be formed as the result of mere instructions (e.g., Field, 2006; Gawronski & Bodenhausen, 2007). From this perspective, mere instructions to approach a given stimulus might allow for the formation of self-stimulus associations, which may lead to more favorable implicit evaluations of the to-be-approached stimulus.

In the current research, we engaged in a preregistered adversarial collaboration to test predictions of a propositional account and an associative self-anchoring account of AA instruction effects in two experiments. Experiment 1 investigated whether both approach instructions and avoidance instructions can cause changes in implicit stimulus evaluations. From the perspective of the associative self-anchoring account, AA instruction effects should occur due to the formation of self-stimulus associations as the result of approach instructions. Processing the semantic meaning of approach instructions should lead to the co-activation of a representation of the self-connected approach action and the to-be-approached stimulus, thereby instigating the automatic formation of an association between the to-be-approached stimulus and the self. Given that most people's implicit self-evaluation is highly positive (Yamaguchi et al., 2007), the subsequent associative transfer of valence should result in a more positive implicit evaluation of the to-be-approached stimulus. In its original formulation, the associative self-anchoring hypothesis does not imply any additional effect of avoidance instructions. Associative self-anchoring is assumed to involve a projection of characteristics of the self to stimuli that are connected to the self but it does not involve a projection of self-characteristics to stimuli that are

Download English Version:

<https://daneshyari.com/en/article/5045711>

Download Persian Version:

<https://daneshyari.com/article/5045711>

[Daneshyari.com](https://daneshyari.com)