

THE INFLUENCE OF MORAL BEHAVIORS
ON PERSON PERCEPTION PROCESSES:
AN FMRI INVESTIGATION

By

Allison Lehner Eden

A DISSERTATION

Submitted to
Michigan State University
in partial fulfillment of the requirements
for the degree of

DOCTOR OF PHILOSOPHY

Communication

2011

ABSTRACT

THE INFLUENCE OF MORAL BEHAVIORS ON PERSON PERCEPTION PROCESSES: AN FMRI INVESTIGATION

By

Allison Lehner Eden

Explicating the underlying neural processes underlying the perception of and reaction to media characters and their behaviors is of central importance to further understanding many theories of media effects. This dissertation uses moral cognitive neuroscience to identify and examine the hypothesized connections between moral judgment and person perception processes central to media enjoyment, by using functional magnetic resonance imaging to test predictions relevant to media theory. Of central concern was the influence of moral relevance (the extent to which narratives were moral or neutral in content), moral valance (the morality or immorality of the content) and moral domain (the extent to which the content evoked theoretically distinct domains) on moral judgment. Additionally, the role individual differences in moral intuitions play in moral judgment was included to identify the role of moral salience in making moral judgments along these dimensions.

A 2(Moral Relevance) x 2(Moral Valance) x 2(Domain) x 2(Moral Salience) experiment was conducted using short statements that varied in moral relevance, moral valance, and domain specificity as stimuli. Participants who varied in the moral salience of specific moral intuitions as judged by the Moral Foundations Questionnaire judged the relative morality of the behaviors presented in these statements while undergoing functional magnetic resonance imaging scanning. Both the moral judgments and the imaging data were analyzed separately in order to gain a complete picture of the processes underlying moral judgment.

Findings suggest that moral content, compared to neutral content, activates a distinct “moral judgment” network in the brain. This moral judgment network further depends on the valence of the moral stimuli. That is, positively and negatively valenced statements elicited distinct patterns of neural activation. This activation also varied based on the domain of the moral behavior, supporting theoretical distinctions between moral domains central to recent theorizing in moral psychology. Furthermore, both moral judgments and neural activation varied based on participants’ self-reported salience of moral intuitions. That is participants with distinct self-report scores on moral salience showed distinct patterns of neural activation across all conditions. Importantly, this indicates that the MFQ can be used to detect important differences in moral judgments between participants. Additionally, it suggests that people for whom moral intuitions are differentially salient when making moral judgments rely on different neural processes when judging moral and immoral behavior.

This study highlights the importance of moral relevance, moral valence, and domain specificity of moral behaviors in moral judgment at both behavioral and neural levels. Furthermore, it distinguishes the separate processes involved in moral intuition and social information processing in a manner that is theoretically relevant for media scholars, moral psychologists, and social cognitive neuroscientists. Finally, it suggests the importance of understanding the moral intuitions that underlie moral judgments in order to gain a fuller understanding of the moral judgment process.

ACKNOWLEDGEMENTS

This dissertation would not be complete without significant contributions from others. Although I am the author, I owe my successful defense of this project to several people who have helped to make it possible, and I would be remiss if I did not acknowledge their efforts here.

First and foremost, I have to thank my advisor, Ron Tamborini, for encouraging me to complete this particular project. I would not have begun nor completed it without his unflagging encouragement and belief in my abilities. He has been a mentor of unparalleled dedication and commitment during my time at Michigan State University, and I owe my current success directly to his efforts on my behalf. I also must thank my committee, especially Issidoros Sarinopolous for his guidance in experimental design and fMRI analysis, and Rene Weber for his encouragement and support of my interest in combining media theory and fMRI research. John Sherry and Frank Biocca also deserve a large thanks for helping me see the broader implications of the work beyond the current project. I also would like to thank my research assistants and coworkers Lu Wang, Kellen Fox, Austin Lee, Matthew Grizzard, Robert Lewis, and Allison Gardner.

Finally I must thank my family, in particular my husband, Jason Weller, who deserves a special mention for his continued love, encouragement, and support of my endeavors throughout graduate school. Also my daughter, Hazel Weller, who gave me an unarguable end date for data collection and analysis, and my mother for demonstrating in her own way the value of a PhD. The rest of my family was unfailingly supportive throughout this process, and I thank them for their love and support as well.

TABLE OF CONTENTS

LIST OF TABLES.....	vi
LIST OF FIGURES.....	vii
KEY TO ABBREVIATIONS.....	ix
INTRODUCTION.....	1
APPENDICES.....	63
BIBLIOGRAPHY.....	117

LIST OF TABLES

Table 1 <i>Moral Statements</i>	22
Table 2 <i>Descriptive statistics for all statements</i>	25
Table 3 <i>Moral Judgment and Reaction Time Data</i>	30
Table 4 <i>Means of Post-Scan Emotion Ratings</i>	33
Table 5 <i>Brain Regions Showing a Significant Main Effect of Morality (Moral vs. Morally Irrelevant)</i>	38
Table 6 <i>Brain Regions Showing a Significant Main Effect of Trait Progressivism</i>	39
Table 7 <i>Brain Regions Showing a Significant Morality x Trait Progressivism Interaction</i>	40
Table 8 <i>Brain Regions Showing a Significant Main Effect of Moral Valence (Negative vs. Positive)</i>	41
Table 9 <i>Brain Regions Showing a Significant Moral Valence x Trait Progressivism Interaction</i>	42
Table 10 <i>Brain Regions Showing a Significant Main Effect of Broad Moral Domain</i>	42
Table 11 <i>Brain Regions Showing a Significant Trait Progressivism x Broad Moral Domain Interaction</i>	43
Table 12 <i>Brain Regions Showing a Significant Main Effect of Moral Contrasts (Negative-Positive Autonomy > Negative-Positive Community)</i>	44

Table 13 <i>Brain Regions Showing a Significant Contrast</i> <i>(Negative-Positive Autonomy > Negative-Positive Community) x</i> <i>Trait Progressivism Interaction</i>	44
--	----

LIST OF FIGURES

Figure 1 <i>Pilot Ratings of Moral Emotions, Visual Representation</i>	26
Figure 2 <i>Pilot Ratings of Moral Emotions x Moral Valence</i>	26
Figure 3 <i>Trial Structure</i>	29
Figure 4 <i>Faces Used in fMRI Experiment</i>	85
Figure 5 <i>Post-scan Survey</i>	95
Figure 6 <i>Pilot Survey Emotion Ratings</i>	111

KEY TO ABBREVIATIONS

Direction

R Right

L Left

Region of Interest

ADLPFC Anterior Dorsolateral Prefrontal Cortex

AINS Anterior Insula

AMYG Amygdala

ASTS Anterior Superior Temporal Sulcus

CERBL Cerebral Cortex

HPThal Hypothalamus

IPL Inferior Parietal Lobule

FUSIG Fusiform Gyrus

LCLC Left Calcarine Cortex

LNGG Lingual Gyrus

MCC Middle Cingulate Cortex

MOGG Middle Occipital Gyrus

OCG Occipital Gyrus

OPRC Operculum

PDLPFC Posterior Dorsolateral Prefrontal Cortex

PINS Posterior Insula

POCG Posterior Occipital Gyrus

POTZ	Parietooccipital Zone
PRCG	Precentral Gyrus
PRCN	Precuneus
PSTS	Posterior Superior Temporal Sulcus
RPDLPFC	Right Dorsolateral Prefrontal Cortex
SMA	Supplementary Motor Area
SMDLPFC	Superior Dorsolateral Prefrontal Cortex
SMPFC	Superior Medial Prefrontal Cortex
SPL	Superior Parietal Lobule
SPFC	Superior Prefrontal Cortex
STS	Superior Temporal Sulcus
THA	Thalamus
TPJ	Temporal Junction
TPL	Temporal Lobe

THE INFLUENCE OF MORAL BEHAVIORS ON PERSON PERCEPTION PROCESSES: AN FMRI INVESTIGATION

The perception of and reaction to media characters and their behaviors is central to many theories of media effects, such as social cognitive theory (Bandura, 2001) and disposition theory (Zillmann, 2000). As such, explicating the underlying neural processes that govern these perceptions and reactions is central to understanding audience reactions to media. Recent research into the neuroscience of person perception suggests that moral information can alter perceptions of others' character (Delgado, Frank, & Phelps, 2005; Todorov, Gobbini, Evans, & Haxby, 2007). Past research in this area, however, has not explicated theoretical mechanisms for the influence of moral information on character perception, nor the contexts under which this influence occurs (Ames, Fiske, & Todorov, 2011). On the other hand, the theoretical mechanisms linking character morality to character perception has been a central focus of media research (Raney, 2004), although this research has not yet investigated the neural networks underlying these processes. A key linkage between these areas can be found in recent work on moral processing networks in moral cognitive neuroscience. The central goal of the present study, therefore, is to use moral cognitive neuroscience to identify and examine the hypothesized connections between moral judgment and person perception processes, using functional magnetic resonance imaging to test predictions relevant to media theory.

Specifically, this project is concerned with the following questions: What effect does moral information have on moral judgment processes? What neural processes and networks are affected or involved? Can the moral judgment process be affected by characteristics of the moral information itself such as valence or contextual domain? Is this process moderated by

individual differences, such as the salience of moral intuitions? These questions will be addressed by assessing the neural activations evoked by exposure to moral messages about people in the news, as well as explicit moral judgments about these people. Furthermore, by using stimuli designed to evoke intuitive moral judgments within specific content domains suggested by recent theorizing in moral psychology, this study will explore the domain specificity of moral judgment and the importance of moral intuitions in moral judgment. This will distinguish the separate processes involved in moral intuition and social information processing as well as the possible domain-specificity of moral information in a manner that is theoretically relevant for media scholars, moral psychologists, and social cognitive neuroscientists.

The following paper describes past research informing these questions and the proposed methodology for examining them. First, the role of moral information on character perceptions in media is introduced. Next, research on moral cognitive neuroscience is reviewed, with a focus on intuitive models of morality shown to be relevant in media appreciation, such as the networks involved in person perception processes. Next, the paper introduces evidence for the domain-specificity of morality and the role that moral intuitions play in understanding individual differences salient for media appreciation. Finally, the current study is described in detail.

Character perception processes

The role of moral judgment in media entertainment is important for understanding how individuals relate to and empathize with characters. Zillmann (2000) proposed that perceptions of character morality are critical in forming dispositions towards characters (Zillmann & Bryant,

1975; Zillmann & Cantor, 1976; Raney & Bryant, 2002) in that viewers are unable to feel empathy for characters unless viewers can positively appraise the morality of characters' behavior. Zillmann suggested that viewers first feel an automatic, (or intuitive, implicit – that is, not consciously deliberated upon), empathic response to a characters behavior. Viewers then reappraise their initial emotional reaction by cognitively assessing the behavior and social information associated with it in terms of their own morality. The initial process takes place almost instantaneously, and is then constantly revised as viewers receive more information about a character.

Understanding how these person-perception processes can influence disposition formation has been found to be especially important in narrative media presentations (Raney, 2004). For example, writers may play upon our natural automatic evaluation mechanisms by portraying characters having specific appearances or behaviors designed to elicit rapid evaluation of their morality. Past media research has focused on creating taxonomies of these 'good' and 'bad' behaviors and appearances, focusing on clearly defined variables such as actions and consequences that help viewers distinguish and attribute intentionality to media characters (Hoffner & Cantor, 1991; Liss, Reinhardt, & Fredrickson, 1983; Himmelweit, Oppenheim, & Vince, 1958; Konijn & Hoorn, 2009), as well as clearly observable physical traits that lead to quick person perception and categorization (Hoffner & Cantor, 1991; Raney, 2004). Research also suggests that the processes responsible for reactions to narrative media presentations are critical for understanding reactions to news media as well (Zillmann & Knobloch, 2001). Affective reactions to hearing or watching good or bad news, for example, are

mediated by the perceived deservingness of the victim. This concept of perceived deservingness is a direct result of a viewer's moral judgment.

In psychological research, this type of person perception process falls under a large body of research known as attribution theories (Heider, 1958). These theories suggest that people understand other people based on the dispositions or personalities they ascribe to others and the situations they see others in. Copious research in social psychology suggests that people ignore the situational factors when judging others' behavior in favor of making dispositional attributions (e.g. Gilbert & Malone, 1995). Dispositional attributions offer an easy way for people to explain behavior: "He *did* x because he *was* y." This type of evaluation is generally considered an automatic process (Winter & Uleman, 1984) in which people form an initial evaluation of others quickly. This evaluation may encompass moral information such as the trustworthiness or untrustworthiness of others (Todorov, Mandisodza, Goren & Hall, 2005; Luo et al., 2006), that may be processed without conscious moral deliberations regarding the person or his/her behavior.

There is substantial information in the social cognitive literature regarding the neural circuitry involved in face processing (Haxby, Hoffman, & Gobbini, 2006), which is gradually being encompassed into a larger understanding of person perception processes. Face perception is the most highly developed social skill for most people, with face viewing persisting for the duration of most social interactions. Face perception involves specific and distributed neural systems involving the fusiform face area (FFA, Kanwisher, McDermott, & Chun, 1997), the posterior superior temporal sulcus, and the amygdala, insula, TPJ, ATC, and MPFC (Haxby, Hoffman, & Gobbini, 2000; see Vuilleumier & Pourtois, 2007 for review). Due to

the wealth of literature on neural networks involved in face processing, studying the processing of faces is one way to closely examine more nebulous concepts of person perception and impression formation.

Person perception and impression formation, separate from face processing, utilizes social cognitive networks involving the MPFC, STS, OFC, amygdala, and anterior insula (Frith & Frith, 2006). For example Mitchell, Macrae, & Banaji (2004) found that the DMPFC preferentially activated in impression formation versus sequencing information. Similar to character perceptions, person perception process can be moderated by dispositional information regarding individuals' behavior. This type of social information, such as impressions of a person's trustworthiness, future reciprocity, or affective associations, can alter person perception processes.

Several studies have examined the role of morally valenced information, such as the "goodness" or "badness" of people, on face perception. These studies show involvement of the DMPFC and MPFC regions involved in social information processing. For example, Singer et al. (2004) asked participants to judge the gender of opponents after playing an ultimatum game. Participants viewing faces ostensibly belonging to the unfair opponents showed greater activation of the amygdala, orbitofrontal cortex (OFC), posterior superior temporal sulcus (STS), and anterior insula than when viewing faces belonging to fair opponents. In Delgado et al., (2005) participants used information about an individual's moral character (good, bad, neutral) to determine impressions during a trust game. Reward system responses involving the caudate nucleus were significantly less active when playing the game with a "bad" player than when playing with a "good" or "neutral" player, regardless of the actual play style of the player.

Todorov et al. (2007) presented short behaviors that were aggressive, disgusting, neutral, or sad, paired with a neutral face. After a short delay, participants then viewed the faces in the fMRI scanner. Faces paired with neutral and nice behaviors elicited greater activation in the STS compared with disgusting or aggressive faces. Disgusting behaviors showed greater left anterior insula, right STS, left cingulate gyrus, right anterior cingulate gyrus, left intraparietal sulcus, left inferior frontal gyrus, and right precuneus activation than other faces. There could also be simple contextual priming effects for valence on person perception processes. For example, Mobbs et al. (2005) examined the role of context on the processing of neutral faces. After a negative valence picture (dogs barking, guns) participants showed greater bilateral ventrolateral prefrontal cortex (vlPFC), bilateral temporal pole, and right fusiform face gyrus (FFG) activation when attributing emotion to neutral faces. After positively valenced pictures (such as babies or animals) participants showed greater ventromedial prefrontal cortex (vmPFC), bilateral temporal pole, and right FFG activation.

Taken together, existing research suggests that there are dissociable and identifiable neural networks involved in person perception that are affected by social information regarding and surrounding those people. Primarily, these networks involve the dorsomedial prefrontal cortex (DMPFC), vmPFC, amygdala, and FFG. Of particular importance to the present study, these networks can be utilized to distinguish the relevance of moral information and the effect of context on person perception.

To date, most research in this area has examined the moral judgments involved in disposition formation as a result of deliberative moral processes. Recently, however, a new line of research has elaborated on the automatic, or intuitive, processes involved in judging the

morality of others. This intuitionist model of morality has large implications for understanding person perception processes central to media and psychology research.

Moral intuitions

The proposition of moral intuitions was originally proposed in order to explain a phenomenon called *moral dumbfounding*. Moral dumbfounding occurs when people have an immediate intuitive “good-or-bad” reaction to a behavior that cannot be justified rationally (For example, many people exhibit a negative intuitive response to consensual incest or homosexuality). Trying to understand where these gut reactions originate, Haidt (2001) proposed a model of morality that is based on what he called social (or moral) intuitions. Haidt describes moral intuitions as an automatic evaluative reaction experienced in response to a socially normative behavior or violation. Although people can overcome this initial evaluative response (Haidt & Graham, 2009; Amodio, Jost, Master, & Yee, 2007) it takes cognitive effort and is a slower, deliberative process.

Haidt and Joseph (2008) proposed that these moral intuitions fall into five domains: Harm/care (concerned with the suffering of others and empathy); fairness (related to reciprocity and justice); loyalty (dealing with common good and punitiveness toward outsiders); authority (negotiating dominance hierarchies); and purity (concerned with sanctity and contamination). These five primary domains are found throughout every culture, although the salience of specific domains differs between cultures (cf., Haidt & Joseph, 2008; Shweder, Much, Mahapatra, & Park, 1997). This contextualization of moral intuitions along evolutionarily determined lines has been called moral foundations theory (MFT: Haidt & Joseph, 2008).

Moral foundations theory has been researched primarily in the political arena. For example, Graham, Haidt, and Nosek (2009) studied the role of moral foundations in political identity. They found that political liberals value primarily harm and fairness when making moral decisions, whereas political conservatives value all five foundations equally. This pattern has been found across national boundaries (Bowman, Dogruel, & Joeckel, 2011) as well as within them (van Leeuwen & Park, 2009; Lewis, Grizzard, Bowman, Eden, & Tamborini, 2011). Haidt and Kesebir (2010) suggest that, fundamentally, political liberals and conservatives use different processes to make moral judgments, and that these processes are based in the salience of moral intuitions.

Moral foundations have also been the focus of a recent program of study in media entertainment (MIME: Tamborini, in press; Eden, Oliver, Tamborini, Woolley, & Limperos, 2009; Eden and Tamborini, 2010; Tamborini, Eden, Bowman, Grizzard, & Lachlan, in press; Tamborini, Eden, Bowman, Grizzard, & Weber, 2009). This program of study has investigated the relationship between moral intuitions and media, focusing on the role of moral intuitions in understanding and interpreting character behavior and narrative justification. Results support the notion that moral intuitions may shape response to characters and narrative plots in media. These findings in this line of research are most important for indicating that these intuitive moral codes may guide reflexive responses to popular media experience. They are also important in that they provide evidence that including intuitive as well as explicit judgments of morality can help us understand how audiences evaluate acts and form perceptions of media characters in terms of their own morality. Additionally, findings from Eden and Tamborini (2010) and Tamborini et al., (2009) suggest that individual differences in the salience of moral

intuitions can alter a viewer's response to media products. For example, Eden and Tamborini (2010) found that the salience of fairness to participants moderated the formation of positive dispositions towards a character who violated or upheld fairness. Tamborini et al., (2009) found that the salience of harm or fairness to participants moderated their acceptance of violent or unjustified narratives.

Research on moral intuitions in both the political and media areas has been based largely in experimental and survey research that used using self-report techniques to measure moral intuitions. However, self-report data may reflect conscious deliberation rather than intuitive processes, thus providing an incomplete - or worse, inaccurate - picture of moral judgment. Therefore, recent work in moral psychology has turned to behavioral and neuroscientific paradigms in order to examine moral intuitions without relying on self-report data. Using reaction time data, for example, Luo et al., (2006) and Van Leeuwen and Park (2009) examined reactions to moral and immoral statements using an Implicit Attitudes Test; Hofmann and Baumert, (2010) adopted an affect misattribution procedure to examine the effect of moral affect on guilt; and Tamborini, Lewis, Grizzard, and Eden (2011), used a similar affect misattribution procedure to test the salience of moral primes. Results, however, have not been as strong as the self-report data would suggest. Therefore other methods of measurement, such as functional magnetic resonance imaging, have emerged as an important way to isolate moral intuitions for study.

Moral cognitive neuroscience

The notion of integrating moral emotions and deliberative moral judgment has been a focus of social cognitive neuroscience research in the past decade. As a field of inquiry, moral

cognitive neuroscience “aims to elucidate the cognitive and neural mechanisms that underlie moral behavior” (Moll, Zahn, Olivira-Souza, Krueger, & Grafman, 2005). Results in this area of research indicate that there are distinct patterns of neural activation that accompanies the processing of moral stimuli. These patterns are consistent enough across studies that some researchers have described these areas as a “moral judgment network” (Moll et al., 2005b; Prehn et al., 2010). Brain regions indicated in these studies include the medial prefrontal cortex (PFC) and ventral PFC, orbitofrontal cortex (OFC), gyrus rectus (Moll et al, 2002a; de Oliveira-Souza & Moll, 2000; Moll et al., 2005a; Harenski & Hamann, 2006; Schiach Borg, Hynes, Van Horn, Grafton, & Sinnott-Armstrong, 2006), posterior cingulate cortex (PCC: Heekeren et al., 2003, Heekeren et al., 2005; Moll et al. 2002b), amygdala (Moll et al., 2002a, Berthoz, Armony, Blair, & Dolan, 2006, Harenski & Hamann 2006, Luo et al., 2006), as well as the temporal pole and temporal junction (TPJ: de Oliveira-Souza & Moll, 2000; Moll et al 2002b, Heekeren et al., 2003, Heekeren et al., 2005, Schiach Borg et al., 2006) (see Raine & Yang, 2006, for overview). These areas are active when participants view or read moral violations (either in pictures or in short statements), as compared to neutral statements and pictures (although methodological differences between studies make further comparisons difficult).

The consistent findings in this area suggest that behaviors with a moral component elicit differential neural activation than morally irrelevant behaviors. Whether this activation is consistent for moral behaviors versus morally irrelevant behaviors, or is simply indicative of immoral behaviors, is less clear. There has been some research on moral virtue behaviors suggesting that moral behavior elicits neural activation consistent with reward processing. For example, cooperation in social games increases activation in the nucleus accumbens (nACC).

(Rilling, Sanfey, Aronson, Nystrom, & Cohen, 2002; King-Casas et al., 2005). Moll et al., (2006) found that the ventral tegmental area, dorsal striatum, ventral striatum, and subgenual areas were active when people contemplated making a charitable donation. Ventral striatum activation in charitable giving was replicated by Harbaugh, Mayr, and Burghart (2007). Social reward shares same substrate as monetary rewards in ventral striatum (Izuma, Saito, & Sadato, 2009). However, this research focuses on the reward participants feel when they act morally, rather than when judging the moral behavior of others.

Although past research in moral cognitive neuroscience has advanced understanding of the neural processes underlying moral judgment, one limitation of the research is the broad “moral versus morally irrelevant” approach taken by most studies. Understanding that morality is not homogenous (Haidt & Joseph, 2008; Shweder, et al., 1997) some moral neuroscience researchers have turned to understanding the domain specificity of moral judgment.

Domain specificity of moral judgment

The domain-specificity of moral judgment, that is, the extent to which there are distinct networks associated with specific moral violations or virtues, has been a topic of theoretical debate among moral psychologists (Haidt & Joseph, 2008). At the neural level, Moll et al. (2005b) suggest that the moral judgment network is bound by: “...norms and contextual elements of social situations, [and are] elicited in response to violations or enforcement of social preferences and expectations. Although the contextual cues that link moral emotions to social norms are variable and shaped by culture, these emotions evolved from prototypes found in other primates and can be characterized across cultures” (Moll et al., 2005b).

If morality can evoke intuitive responses based on the domain of content information presented, then there should be discrete neural patterns of activation for moral information from one content domain compared to moral information from other content domains, or compared to information from neutral (non-moral) domains. There is some preliminary evidence to support this domain-specific notion (c.f. Schaich Borg et al., 2006; Schiach Borg, Lieberman, & Kiehl, 2008; Wright, Shapira, Goodman, & Liu, 2004). Most recently, Parkinson et al. (2011) directly compared harmful, dishonest (fair), and disgusting (purity) violations. They found differences in neural processing between the types of transgressions, and have called for a systematic investigation into the domain-specificity of moral judgment, suggesting the foundations suggested by Haidt (harm, fairness, authority, loyalty, purity) could help organize this research. However, conceptual issues with distinguishing the five domains Haidt proposes may limit the ability to find different neural networks associated with each domain (for example, disgust responses associated with purity may also underlay responses in other moral domains). Therefore, it may be relevant to begin with Haidt's moral foundations, but also examine broader groupings of the foundations within broader domains.

Examining first the four domains of harm, fairness, authority, and loyalty, it is immediately apparent that these domains are not considered separate in all research. For example, Shweder et al. (1997), whose work provided background for Haidt's moral foundations, proposed two broad domains instead of four more narrow ones. Shweder et al. (1997) proposed that morality can be broken into the two broad domains of autonomy, associated with the rights of the individual; and community, associated with ethics that protect the group traditions and loyalty. Conceptually, the Haidt and Joseph (2008) domains of harm

and fairness fall into the “autonomy” domain, and authority and fairness into the “community” domain. The autonomy domain includes violations which take away others’ personal freedoms (such as volition or health; i.e. Haidt’s harm domain) and virtues which emphasize giving or expanding others’ personal freedoms (i.e. Haidt’s fairness domain; Haidt & Kesebir, 2010), whereas the community domain involves violations which go against society, specifically duty, hierarchy, interdependency, and group values including normative behaviors that are not explicitly involved in denying or upholding the personal rights of others.

This broader conceptualization of domains has been supported in non-neuroscientific work by Graham, Haidt, and Nosek (2009), suggesting that liberals place more emphasis on autonomy violations, and conservatives on community violations, and by Lewis et al. (2011) demonstrating the liberal and conservative students react differently to authority versus harm violations in media. It has also been conceptually echoed by Haidt, Graham, and Hersh (2006) who proposed that the difference between liberals and conservatives may be found by subtracting scores on authority, group loyalty, and purity from scores on harm and fairness. The resulting Trait progressivism scale thus indicates the extent to which individuals are different on the two broad domains of autonomy and community.

The evidence for neurological correlates of morality in the autonomy domain is extensive, as moral violations in the harm and fairness domains are perhaps the most studied areas of moral cognition. In the harm domain, several studies have shown that the anterior cingulate cortex (ACC) and anterior insula (AI) are activated both when people are in pain themselves and when they see others in pain (see Singer & Steinbeis, 2009, for overview). Wright et al. (2003) found that viewing pictures of mutilation (harm domain) caused greater

activation of the orbital temporal cortex and right superior parietal cortex than disgust-based pictures. Although Wright et al. (2003) did not frame their study as an investigation into morality; their findings closely follow the conceptualization of two domain-distinct areas that Haidt proposed. Heekeren et al., (2005) also found distinct patterns of activation for scenarios featuring physical harm when compared to neutral scenarios. In line with the fairness domain, the following results demonstrate neural activation during cooperation, trust, and fair play experiments: vmPFC, medial prefrontal cortex (mPFC: Decety & Jackson, 2004, McCabe, Houser, Ryan, Smith, & Trouard, 2001, Rilling et al. 2004; Tabibnia, Satpute, & Lieberman, 2009), whereas unfair and untrustworthy responses activate the insula (Sanfey, Loewenstein, McClure, & Cohen, 2003), the ventral striatum in the basal ganglia (de Quervain et al., 2004), the DMPFC (Decety et al., 2004). Parkinson et al. (2011) found that harmful moral transgressions preferentially activated the DMPFC, the supplementary motor area (SMA), the DLPFC, STS, intraparietal lobule (IPL), and cerebellum versus neutral transgressions, and that dishonest transgressions preferentially activated the DMPFC, TPJ, and precuneus versus neutral transgressions.

Although no work to date has examined specific authority violations (to this author's knowledge), viewing a superior individual versus an inferior individual differentially engaged bilateral occipital/ parietal cortex, ventral striatum, parahippocampal cortex, and DLPFC (Zink et al., 2008). Marsh, Blair, Jones, Soliman, and Blair (2009) demonstrated that VLPFC and STC was activated for high status versus low status cues. Similarly, although no work has specifically examined group adherence or betrayal violations, viewing ingroup (but not kin-related) racial faces versus outgroup racial faces elicited greater activity in the amygdala, fusiform gyri,

orbitofrontal cortex, and dorsal striatum (Bavel, Packer, & Cunningham, 2008). Platek and Kemp (2007) found that family faces, versus unknown faces, activated right supramarginal gyrus, right inferior parietal lobe, right precuneus, left middle and inferior frontal gyri, supporting the theory of a midline cortical network for discriminating family from non-family. Importantly, these regions are not generally implicated in either harm or fairness-type violations, indicating that these responses may be dissociated at the neural level.

This conceptualization of Haidt's domains as incorporated into Shweder's broader autonomy and community areas leaves only Haidt's domain of purity ungrouped with other domains. Although MFT identifies purity as one of five separate moral domains, there is reason to believe that neural patterns of activation may not distinguish purity violations from the other domains. Some research has shown that feelings of disgust might underlay all moral judgments (Haidt, Rozin, McCauley, & Imada, 1997; Schnall, Haidt, Clore, & Jordan, 2008; Moll et al., 2005b; Chapman, Kim, Susskind, & Anderson, 2009) thus making purity violations, which are theoretically related to the feelings of disgust, hard to distinguish from other domains. This evidence suggests that purity may underlie all other moral judgments. On the other hand, Horberg, Oveis, Keltner, and Cohen (2009) present self-report data contradicting this notion. In Horberg et al. (2009), purity violations were separate from harm or fairness violations. On the other hand, Shaich Borg et al. (2008) found that neural reactions to incest were separate from sociomoral or contamination reactions, thus suggesting that the domain-specificity of purity may be reliant on the operational definitions used in each study. As such, for the current study, purity will not be included as a separate domain. Instead, the study begins with the assumption

that purity-like responses (that is, insula and amygdala activation) should be equally present in all moral violation conditions.

Although it may seem that Haidt's domains may be distinguished from Shweder's based on the research presented here, it is important to note that the results supporting neural disassociation of Haidt's domains are disparate, have not been intentionally tested as separate domains, and are only suggestive of different networks involved in each domain. Behaviorally, the autonomy and community domains seem to function similarly in terms of determining participant responses to stimuli. Although it may be possible to distinguish neural responses to Haidt's domains, and that is indeed one of the goals of this line of research, at this point grouping the domains into broader domains is justified. It may well be easier at first to identify the broader (but still separate) networks involved in Shweder's domains than to focus specifically on the domains as described by Haidt. However, the failure to identify domain-relevant contextual effects (derived from theory) in prior studies has been a weakness that this study will attempt to address.

The current study

This study is designed to fill important gaps in the existing literature of media and moral psychology. It will examine the role of positive, negative, and domain-specific messages on moral judgments and person perception, as well as the role of moral salience in shaping these judgments. Primarily exploratory in nature, this study has the following goals:

First, the study examines the role of moral relevance of messages on moral judgments and associated neural activation. It is expected that morally relevant messages (messages with

moral content) will preferentially activate the moral processing network versus messages with no moral content.

Second, positively and negatively valenced moral messages will be contrasted with each other. This will allow for examination of moral versus morally irrelevant neural networks, as well as the unique contribution of moral valence to moral judgments and neural activation. It is expected that there will be unique patterns of activation for messages with positive versus negative moral valence.

Third, to examine the role of broad moral domains (autonomy, community) in moral processing, this study will contrast the effect of messages within each broad domain on moral judgments and neural activation. It is expected that there will be distinct patterns of activation involved in the processing of moral messages within each broad domain.

Fourth, to examine the effect of trait differences in moral foundation salience, trait progressivism will be used to form groups of participants who vary on the importance they place on moral domains in making moral decisions. Using trait progressivism (high, low) as between-subjects effects for the prior analyses will allow for an examination of the effect of trait individual differences on moral judgments and on the neural networks involved in moral processing.

Broadly, these findings will help researchers understand the role of moral intuitions in moral judgment from a neural level, as well as help determine the role of domain-specificity in moral judgments. These findings help us explain viewer responses to characters specifically and media narrative in general, and connect media psychology to a broad field of research on moral cognitive neuroscience and person perception.

Method

Procedure overview

First, an online survey was administered to students in Communication courses at Michigan State University. Participants for this online survey completed personality questionnaires measuring individual differences in moral foundations and empathy (see Appendix A for consent form and all measures). A tertiary split of scores on the moral foundations questionnaire was used to identify participants high and low in trait progressivism. These participants were contacted to take part in the fMRI experiment. These participants were additionally screened for handedness, pregnancy, and weight in accordance with the safety procedures recommended by the Cognitive Research Facility at Michigan State (see Appendix B for consent form and measures). If accepted for the study, participants were scheduled for the fMRI task between one and three months after taking the online survey. Upon arrival at the fMRI facility, participants were greeted, administered a new consent form for the fMRI experiment, and trained on a trial run of the experimental procedure prior to scanning (see Appendix C for protocol). After training, participants completed the fMRI task. In this task, participants were presented with 144 moral judgment trials, broken into 6 functional runs of 24 trials each. Trials consisted of the presentation of a person's face, then a short statement describing a moral or non-moral behavior, and finally a ratings screen in which participants rated the perceived morality of the behavior or the person (see Appendix D for faces). Post-scanning, participants rated the moral statements for the extent that the statements evoked

specific moral emotions while in the scanner (see Appendix E). After completing these ratings, participants were debriefed, thanked, and compensated for their time.

Selection Survey

Participants ($n = 590$, $F = 386$) completed an online survey including questions about moral domain salience, empathy, and demographics for course credit (Appendix A). Based on scores on the moral domain salience questionnaire, 188 women were contacted to participate in the fMRI survey. From those contacted, 20 right handed female undergraduate students (10 liberal, 10 conservative, *Mean age* = 20.2 years, *SD* = .12) eventually participated in the fMRI experiment. All participants completed informed consent forms approved by the Michigan State University Institutional Review Board's Human Research Protection Program prior to both the survey and the fMRI experiment (Appendix A).

Measures

Trait Progressivism. The Moral Foundations Questionnaire (MFQ; Haidt, Graham, & Hersh, 2006) was used to determine trait progressivism. The MFQ is a 32-item measure designed to measure the importance to individuals of the five moral domains identified by Haidt and Joseph (2004, 2007). Each domain is measured by three "relevance" items with the stem "When you judge an action as right or wrong, how important are the following considerations in your decision" anchored at 1 (*not at all important*) to 6 (*very important*), as well as three "statement" items which ask participants to rate the extent to which they agree with statements regarding the domain on a 7-point Likert-type scale. Scores on all items were averaged to form a composite salience score for each domain ($M_{harm} = 4.39$, $SD = .74$; $M_{fair} = 4.12$, $SD = .70$; $M_{loyalty} = 3.90$, $SD = .67$; $M_{authority} = 3.92$, $SD = .70$; $M_{purity} = 3.53$, $SD = .87$).

Following the procedure described in Haidt, Graham, and Nosek, (2010), trait progressivism scores for all participants were calculated by subtracting the sum of the community (authority, loyalty, purity) domain scores from the sum of the autonomy (harm and fairness) domain scores ($M_{\text{Trait Progressivism}} = -2.85$, $SD = 1.76$). A tertiary split on trait progressivism scores was used to identify female subjects at either end of the scale for the fMRI experiment. Participants were considered liberal if they scored above -2.66 on the trait progressivism measure. Participants were considered conservative if they scored below -3.69 on the trait progressivism measure. For the 20 participants who participated in the fMRI experiment, the overall mean was -2.67 ($M_{\text{liberal}} = -.83$, $SD = 1.18$, $M_{\text{conservative}} = -4.97$, $SD = .86$).

Stimuli

Faces. Faces were selected from a dataset of 300 computer generated faces available online (<http://webscript.princeton.edu/~tlab/databases/>). Face stimuli consisted of computer generated portrait pictures previously rated along dimensions important to social judgment (i.e. attractiveness, likeability, trustworthiness, competence, extroversion, dominance, meanness, fright, and threat). Ratings for all faces used can be reviewed in Todorov, Baron, and Oosterhof (2008) and are thus not replicated here. Todorov also made available masculinity ratings for all faces in the database (personal communication, September 9, 2010). These ratings included a “probability of being female” score for all faces, on a scale of 0 (0% probability of being female) to 1 (100% probability of being female). The 144 faces selected were those that scored 0 on this measure, that is, faces that were considered unambiguously male. All faces used in the study are included in Appendix D.

Statements. Twelve statements were created for this study. Each statement describes the action of one male person who “went out of his way” to perform an act. Statements were created to vary along moral dimensions with four positive statements (care, fair, loyal, respect), four negative statements (harm, unfair, disloyal, disrespect), and four morally irrelevant statements (collect, get rid of, succeed, fail). Two of these statements (collect, get rid of) were created based on examination of past morally irrelevant stimuli (Heekeren et al., 2005). The other two morally irrelevant messages (succeed, fail) were included to provide emotionally-valenced but morally irrelevant contrasts to the “true neutral” messages provided by ‘collect’ and ‘get rid of’¹. The moral (positive and negative) statements were based on items from the moral foundations questionnaire (Haidt et al., 2006). The moral statements varied along four of the moral domains (harm, fairness, authority, group loyalty) developed by Haidt and Joseph (2008). The positive moral statements reflected the most salient moral virtue in each domain (care, fair, respect, loyal). The negative moral statements reflected the most salient moral vice in each domain (harm, unfair, disrespect, disloyal; Haidt & Joseph, 2008). Based on the moral domains, the moral statements were further grouped into the broader moral domains (autonomy and community) defined by Shweder et al. (1997).

Some of the twelve statements were collapsed in different combinations to form two variables for use in different analyses. The first variable was labeled moral relevance, which had two conditions, morally relevant and morally irrelevant. For this variable, the four positive

¹ For the current analyses, all morally irrelevant statements were collapsed into the morally irrelevant category. A study separate from this dissertation will separate positive (succeed) and negative (fail) morally irrelevant statements from the true neutral morally irrelevant (collect, get rid of) to examine hypotheses related to emotional valence.

moral valenced statements were combined with the four negative valenced statements to create the morally relevant condition, and the four neutral moral statements were combined to create the morally irrelevant condition. The second variable was labeled moral valence, which had two conditions, positive, and negative. For this variable, the four positive moral statements were combined and the four negative moral statements combined to create the separate conditions of positive and negative moral valence. See Table 1 for all statements.

Table 1

Moral Statements

Statement	Moral Relevance	Moral Valence	Broad Domain	Domain
He went out of his way to care	Moral	Positive	Autonomy	Harm/care
He went out of his way to harm	Moral	Negative	Autonomy	Harm/care
He went out of his way to be fair	Moral	Positive	Autonomy	Fairness
He went out of his way to be unfair	Moral	Negative	Autonomy	Fairness
He went out of his way to respect	Moral	Positive	Community	Authority
He went out of his way to disrespect	Moral	Negative	Community	Authority
He went out of his way to be loyal	Moral	Positive	Community	Group
He went out of his way to be disloyal	Moral	Negative	Community	Group
He went out of his way to collect items	Morally Irrelevant	Morally Irrelevant	N/A	N/A
He went out of his way to get rid of items	Morally Irrelevant	Morally Irrelevant	N/A	N/A
He went out of his way to succeed	Morally Irrelevant	Morally Irrelevant	N/A	N/A
He went out of his way to fail	Morally Irrelevant	Morally Irrelevant	N/A	N/A

Statements were pilot tested using a separate sample of 121 undergraduate students to measure arousal level, intentionality and moral emotions (see Appendix F for pilot study, and Table 2 for all descriptive data). The pilot testing was intended to ensure that statements varied in predictable manner between moral valence and moral domain. As expected, arousal levels were higher for the moral statements (both positive and negative; $M = 4.55$, $SD = .32$) than the morally irrelevant statements ($M = 4.00$, $SD = .67$), although the mean for morally irrelevant statements was inflated due to the “He went out of his way to succeed/fail” statements, which were judged to be as arousing as other moral statements. All statements were judged within a 95% confidence interval for intentionality ($M = 5.25$, $SD = .40$) except for “He went out of his way to fail” which was outside the CI ($M = 4.21$, $SD = 1.82$). Due to the difficulty of constructing a true morally irrelevant positive statement as a foil for succeed, this message was retained for analysis purposes.

Moral emotions were derived from Haidt (2003) and included positive moral emotions (admiration, elevation, gratitude) as well as negative moral emotions (disgust, anger, indignation, contempt). Examining the scores for the moral emotions generated by each statement, there is a clear pattern of moral statements (care, fair, respect, loyal) being rated higher on the positive moral emotions than negative moral emotions (see Figure 1). This pattern is reversed for the immoral statements (harm, unfair, disrespect, disloyal). Succeed and Fail statements were judged similarly to moral and immoral statements, respectively.

After examining the figure, moral emotion ratings were combined to form average positive ($M = 2.44$, $SD = 1.12$) and negative ($M = 2.31$, $SD = .93$) emotion ratings. A 2 Moral Valence (positive, negative) x 2 Emotion Valence (positive, negative) repeated measures

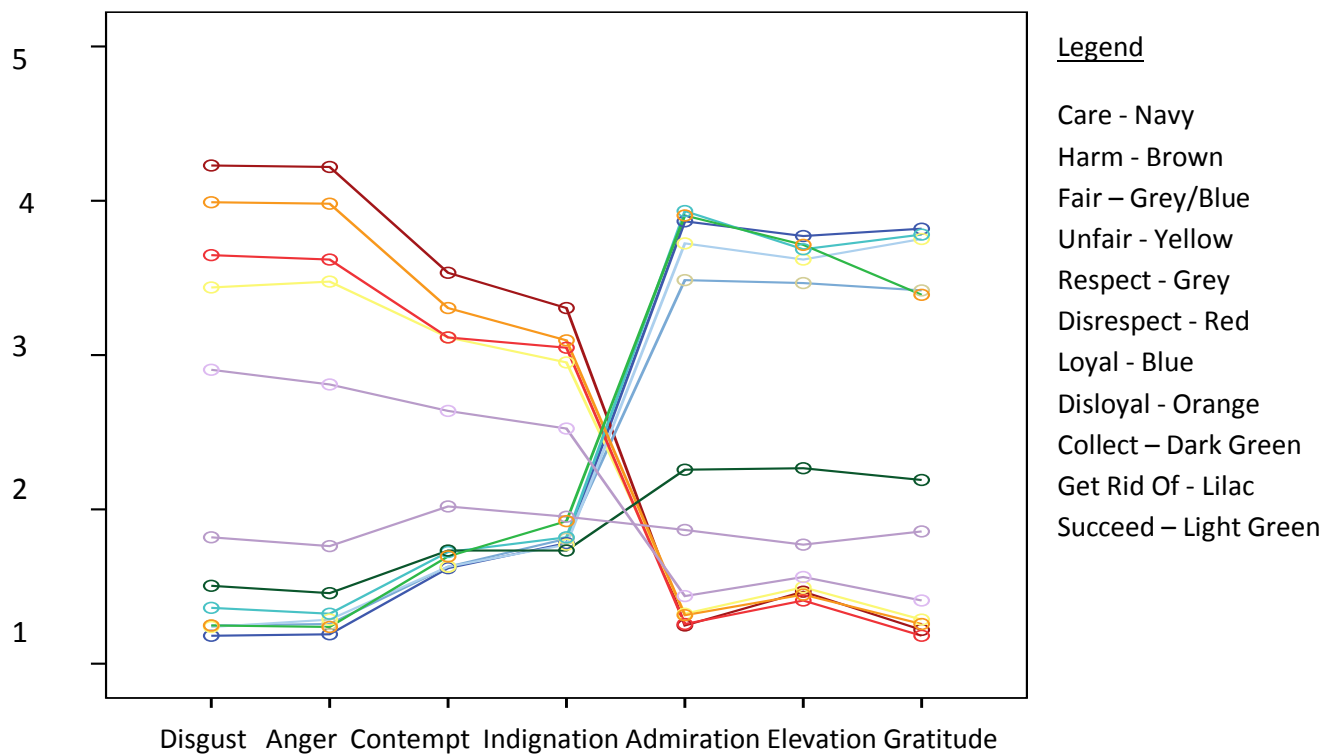
ANOVA was conducted on these scores. There was an interaction of statement valence and moral emotion, $F(2, 9) = 27.76$, $p < .01$, $\eta^2 = .86$, such that positive moral statements (care, fair, loyal, respect) were rated higher on positive moral emotions ($M = 3.68$, $SD = .28$) than negative moral emotions ($M = 1.05$, $SD = .16$). Negative moral statements (harm, unfair, disloyal, disrespect) were rated higher on negative moral emotions ($M = 3.48$, $SD = .16$) than positive emotions ($M = 1.34$, $SD = .28$). Morally irrelevant statements were rated lower on positive emotions than positive statements ($M = 2.30$, $SD = .28$) and higher on negative emotions than negative statements ($M = 1.95$, $SD = .16$). See Figure 2 for graph of results.

Table 2

Descriptive statistics for all statements

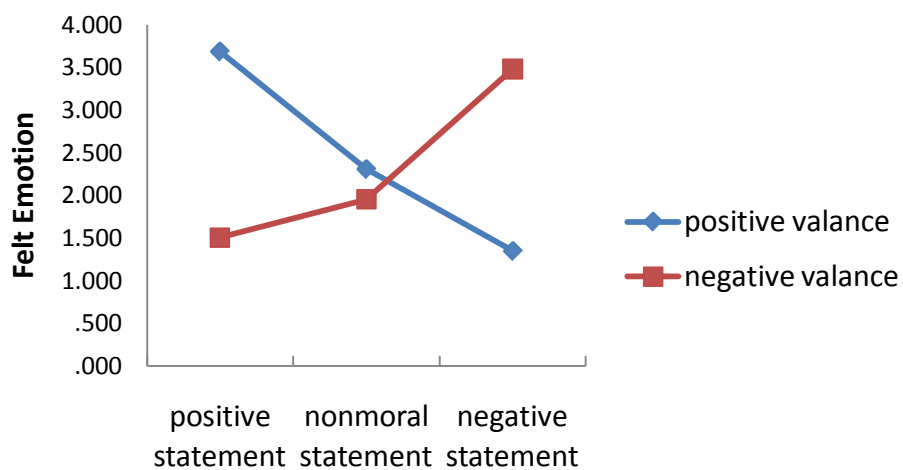
	Arousal		Intention		Anger		Contempt		Indignation		Admiration		Elevation		Gratitude		Disgust	
	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>	<u>M</u>	<u>SD</u>
Care	4.65	1.30	5.31	1.41	1.20	0.60	1.61	1.02	1.78	1.13	3.86	0.98	3.75	1.07	3.84	1.11	1.21	0.62
Harm	4.94	1.34	5.76	1.39	4.19	1.03	3.49	1.37	3.23	1.41	1.27	0.70	1.49	1.02	1.25	0.69	4.17	1.01
Fair	4.13	1.30	5.10	1.34	1.29	0.73	1.64	1.03	1.81	1.16	3.47	1.11	3.46	1.15	3.46	1.08	1.27	0.73
Unfair	4.32	1.43	5.01	1.51	3.50	1.05	3.09	1.08	2.97	1.17	1.37	0.87	1.54	1.01	1.30	0.73	3.43	1.07
Respect	4.28	1.51	5.25	1.47	1.31	0.77	1.67	1.06	1.77	1.11	3.69	1.10	3.59	1.11	3.73	1.12	1.30	0.82
Disrespectful	4.43	1.31	4.95	1.69	3.61	0.97	3.09	1.11	3.02	1.09	1.27	0.64	1.42	0.92	1.20	0.55	3.64	1.01
Loyal	4.63	1.31	5.51	1.35	1.33	0.82	1.74	1.13	1.83	1.22	3.91	1.00	3.66	1.11	3.80	1.10	1.35	0.89
Disloyal	5.02	1.43	5.44	1.65	3.99	1.05	3.25	1.26	3.07	1.35	1.32	0.79	1.47	1.00	1.26	0.72	3.95	1.07
Collect	3.38	1.48	5.47	1.44	1.47	0.88	1.74	0.99	1.72	1.01	2.25	1.09	2.26	1.18	2.19	1.21	1.53	0.94
Get Rid Of	3.58	1.68	5.36	1.52	1.84	1.15	2.03	1.21	1.98	1.18	1.86	1.10	1.81	1.01	1.91	1.15	1.88	1.19
Succeed	4.88	1.37	5.63	1.45	1.29	0.74	1.74	1.14	1.94	1.22	3.88	0.96	3.71	1.10	3.41	1.26	1.31	0.81
Fail	4.16	1.54	4.21	1.82	2.80	1.27	2.58	1.18	2.49	1.14	1.43	0.81	1.57	0.99	1.40	0.76	2.92	1.30

Figure 1
Pilot Ratings of Moral Emotions, Visual Representation



Note. For interpretation of the references to color in this and all other figures, the reader is referred to the electronic version of this dissertation.

Figure 2
Pilot Ratings of Moral Emotions x Moral Valence



Experimental Procedure

Upon arriving at the scanning facility, participants were greeted by an experimenter, administered a consent form and short questionnaire, and then trained on the experimental task prior to entering the scanner (see Appendix C for fMRI experiment protocol). During training, participants were told the following:

“In the past year, over 100,000 news media stories were analyzed based on their most common themes. You will be presented with representative short statements which have been isolated from these common themes. You will also be presented with representative faces which have been graphically manipulated to form composites of the faces of the people associated with the stories. This is so you can put a face with the behavior, to help it seem more real to you. We will be asking you to rate both the people and their behaviors in the scanner based on your own personal sense of right and wrong.”

News media stories were chosen for three reasons. First, many of the participants knew the experimenter had been analyzing a large body of news stories, so it was a plausible scenario. Second, as Zillmann and Knobloch (2001) found, news stories are able to evoke the types of affective responses central to disposition formation processes found in other narrative media. Third, the experimenters felt that insisting that these were “real people” performing these acts helped participants stay engaged with the task during the experiment.

After training, participants entered the scanner and were trained on the use of the response glove to enter the ratings. The response glove is an input device within the scanner that allows participants to interact with and respond to prompts on screen without moving the

body or the head. After demonstrating proficiency with the response glove, determined by successful input following an on-screen prompt, the scanning procedure began. During the scan, participants performed the moral judgment task for six functional runs. After each functional run, participants relaxed for approximately two minutes before the beginning of the next functional run. After all six functional runs were completed, participants relaxed and closed their eyes for eight minutes while anatomical images of their brain were collected. Finally, participants were removed from the scanner and escorted to a laboratory room outside the scanning area. In this room participants completed a post-scan survey, in which they rated the emotions they felt while in the scanner. In addition, participants responded to an open-ended prompt asking them to relate what they thought about when rating each statement (see Appendix E for post-scan measure). After completing the survey, participants were debriefed, thanked for their time, and paid 40\$.

Moral judgment task

Experimental trials were presented using E-Prime (Psychology Software Tools). An event-related design was employed, with trials presented in a fixed randomized order. Trials consisted of the presentation of a face centered on a black screen for 2 seconds, along with a letter P (indicating a person judgment) or A (indicating an action judgment)². Next, a statement appeared in white text below the face for 2 seconds. The moral or morally irrelevant word in each statement was underlined and bolded (i.e. **harm**, **care**, **succeed**). Finally, a response screen appeared, during which participants rated either the action or the person on a 9 point

² For the current analysis, person and action prompts were combined.

response scale from ‘very immoral’ to ‘very moral’ using the prompt “How do you judge this person?” or “How do you judge this action?” See Figure 3 for trial structure. Participants had 6 seconds to respond to this prompt using a response glove before the end of the trial, and were encouraged to take the full 6 seconds to make their rating. On average, participants took 3.5 seconds to make their ratings. All moral judgment ratings and reaction time data are available in Table 3. After the ‘relax’ screen a fixation cross was presented on a black screen for a jittered interval from 3-8 seconds long. Total presentation time for each run was 6 minutes, 42 seconds. Each session consisted of 144 trials spread over six functional runs, with 24 trials per scan. Presentation of faces was randomized using the built-in randomization function of the E-Prime software.

Figure 3

Trial Structure

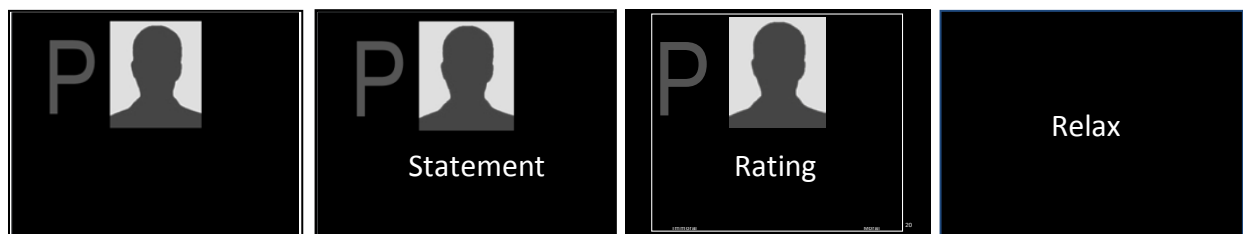


Table 3
Moral Judgment and Reaction Time Data

	Moral Judgment		Reaction Time	
	<i>Mean</i>	<i>SD</i>	<i>Mean</i>	<i>SD</i>
Care	7.42	1.227	3504.54	1024.74
Harm	2.38	1.095	3547.38	1059.08
Fair	7.72	1.125	3540.29	982.17
Unfair	1.81	1.186	3331.58	1019.19
Respect	7.47	1.111	3554.33	1101.10
Disrespect	2.63	1.234	3550.37	1117.10
Loyal	7.43	1.067	3462.08	1010.02
Disloyal	2.37	1.077	3438.87	1069.86
Collect	5.27	1.135	3056.62	1224.29
Get Rid Of	5.17	1.051	3142.10	1139.82
Succeed	6.89	1.085	3483.43	1020.25
Fail	3.18	1.068	3460.46	1091.08

Image acquisition

Data were collected on a 3-T GE Signa EXCITE scanner (GE Healthcare, Milwaukee, WI) with an eight-channel head coil. BOLD contrast was obtained with a gradient-echo echo-planar imaging (EPI) sequence (General Electric scanner; field strength of 3 Tesla; whole brain coverage with 30 interleaved slices; slice size 4mm with 0.4mm gap; TR = 2000ms; TE = 27.2 ms; flip angle = 77°, field of view 22 × 22 cm², matrix size 64 × 64).

fMRI data analysis

All analyses were carried out using AFNI software (<http://afni.nimh.nih.gov/afni>). Individual subject data for the functional scans were slice-time corrected for temporal offsets in the acquisition of slices, and manually motion corrected to the functional image closest in time to the acquisition of the high-resolution anatomical images. This optimized alignment. Blood-

oxygen-level-dependent (BOLD) signal time-series data were then converted to percentage signal change (PSC) for each subject. Percent signal change is an estimate of effect size using the mean of the hemodynamic response function over a specified time series for a voxel. PSC was computed by dividing each time series value, voxelwise, by the mean signal value for that run and multiplying by 100. Single-subject time series were then analyzed using a general linear model (GLM) with separate regressors for each condition, formed by convolving the first six seconds of each stimulus duration with an ideal hemodynamic response function. Resultant beta weights were used to calculate voxelwise contrasts between moral and morally irrelevant trials. Contrast maps were then registered to the anatomical images for each subject, transformed into a standardized template, and entered into voxelwise group GLM analyses to identify brain regions showing reliable differences between moral valence and broad domain conditions.

Post-scan survey

After the scanning procedure, participants viewed each statement again and were asked to rate the extent to which each statement evoked a particular moral emotion on a 10-point scale from 1 (not at all) to 10 (the most you could feel; See Appendix F). Emotions were derived from Haidt (2003) as well as pilot survey data and included positive moral emotions (admiration, gratitude, elevation) as well as negative moral emotions (disgust, anger, contempt). From the pilot survey, 'indignation' was eliminated as it was too close to anger. One emotion (boredom) was added to capture participant interest in the statements, and three more to capture basic emotional responses without a moral component (joy, sadness, fear). All descriptive data for these ratings is available in Table 4. As can be seen from the table, positive

moral statements elicited higher ratings on positive moral emotions than negative moral emotions. Negative moral statements elicited higher ratings on negative moral emotions than positive. The neutral statements elicited low scores on all moral emotions. As these findings are secondary to the main goals of the current study, these ratings are not analyzed further, although they are discussed briefly in the discussion section where relevant to the major findings.³

³ Post-scan moral ratings were analyzed using a 2 Trait Progressivism (liberal, conservative) x 2 Broad Moral Domain (autonomy, community) x 2 Moral Emotion (positive, negative) mixed ANOVA. There was no main effect for Domain, $F(1, 18) = .99, p = .34, \eta^2 = .05$, Moral Emotion $F(1, 18) = .08, p = .78, \eta^2 = .00$, Trait Progressivism $F(1, 18) = 3.79, p = .06, \eta^2 = .17$, although Trait Progressivism approached significance. The only interaction effect was between Domain x Moral Emotion such that for statements in the autonomy domain, positive emotions ($M = 3.31, SD = .25$) were rated less intense than negative emotions ($M = 3.44, SD = .26$) whereas in the community domain, positive emotions ($M = 3.38, SD = 3.44$) were rated as more intense than negative emotions ($M = 3.16, SD = .28$), $F(1, 18) = 6.73, p < .05, \eta^2 = .27$. No other interaction effects were significant.

Table 4

Means of Post-Scan Emotion Ratings

	Disgust	Anger	Contempt	Admiration	Elevation	Gratitude	Boredom	Sadness	Joy	Fear
Care	1.33	1.14	2.14	5.58	4.11	5.64	1.94	1.14	5.78	1.14
Harm	6.53	6.44	4.36	1.28	1.72	1.33	1.92	5.81	1.28	5.17
Fair	1.25	1.33	2.17	5.58	4.39	5.19	1.72	1.28	5.33	1.42
Unfair	5.42	5.31	4.00	1.22	1.25	1.28	1.72	3.92	1.31	2.33
Respect	1.17	1.14	1.94	5.83	4.47	5.28	2.03	1.17	5.58	1.17
Disrespect	5.06	5.69	3.58	1.31	1.61	1.22	1.69	3.97	1.19	2.39
Loyal	1.39	1.28	2.08	6.03	4.42	5.56	1.75	1.31	5.25	1.42
Disloyal	5.25	5.22	4.14	1.25	1.33	1.22	1.53	4.53	1.31	2.39
Collect	1.72	1.69	1.69	1.86	1.83	1.72	2.42	1.36	1.94	1.72
Get Rid Of	1.36	1.47	1.44	2.19	1.94	1.92	2.22	1.61	2.03	1.50
Succeed	1.14	1.19	2.03	5.92	4.44	3.44	1.69	1.19	5.17	1.36
Fail	3.86	3.94	3.44	1.39	1.53	1.47	2.14	4.75	1.33	1.75

Results

Overview

To ensure that the moral judgment ratings were consistent with expectations, moral judgment ratings were subjected to a 2 Trait Progressivism (liberal, conservative) x 3 Moral Valence (positive valence, negative valence, morally irrelevant) mixed ANOVA. To examine the effect of moral domain on moral judgments, moral judgment ratings were analyzed with a 2 Trait Progressivism (liberal, conservative) x 2 Moral Domain (autonomy, community) mixed ANOVA. Finally, a 2 Trait Progressivism (liberal, conservative) x 2 Moral Domain (autonomy, community) x 2 Moral Valence (positive, negative) ANOVA was conducted to examine these results in greater detail.

Imaging data were analyzed using four separate designs. First, a 2 Trait Progressivism (liberal, conservative) x 2 Moral Relevance (morally relevant, morally irrelevant) ANOVA was run to examine the main effect of moral relevance, trait progressivism, and the interaction between the two. Next, a 2 Trait Progressivism (liberal, conservative) x 2 Moral Valence (positive, negative) ANOVA was run to examine the role of positive and negative moral valence on moral judgment. After this, to examine the role of moral domains in moral judgment, a 2 Trait Progressivism (liberal, conservative) x 2 Moral Domain (autonomy, community) ANOVA was conducted. Finally, a contrast-based ANOVA combining group, domain, and valence was conducted. This was run by subtracting positive from negative moral statements within autonomy and community domains to construct contrast composites (negative-positive within autonomy, negative-positive within community). These composites were then used in a 2

Contrast Composite (negative-positive within autonomy, negative-positive within community) x 2 Trait Progressivism (liberal, conservative) ANOVA.

Moral judgment ratings

First, a 2 Trait Progressivism (liberal, conservative) x 3 Moral Valence (positive, negative, morally irrelevant) mixed ANOVA was conducted on all moral judgment ratings to ascertain the perceived morality of positive moral, negative moral, and morally irrelevant statements. There was a main effect for morality such that positive moral statements ($M = 7.51, SE = .12$) were rated as more moral than morally irrelevant statements ($M = 5.13, SE = .05$) which were in turn rated as more moral than negative moral statements ($M = 2.30, SE = .12$), $F(1.13, 20.45) = 465.99, p < .001, \eta^2 = .96$ (Greenhouse-Geisser correction employed). This was interpreted to mean that the manipulation of moral valence was successful.

There was also a main effect for trait progressivism, $F(1, 18) = 7.27, p < .05, \eta^2 = .288$ such that liberals ($M = 5.05, SE = .04$) judged all statements as more moral than conservatives ($M = 4.90, SE = .04$). There was no interaction effect, $F(1, 18) = 1.17, p = .32, \eta^2 = .06$. This indicated that the trait progressivism measure was successful in identifying differences between groups that were relevant for moral judgment processes.

To examine the impact of moral domains on moral judgment a 2 Trait Progressivism (liberal, conservative) x 2 Broad Moral Domain (autonomy, community) ANOVA was conducted. There was a main effect for broad moral domain such that autonomy statements ($M = 4.82, SD = .17$) were rated as less moral than community statements ($M = 4.97, SD = .14$), $F(1, 18) =$

14.16, $p < .001$, $\eta^2 = .44$. This was interpreted to mean that differences between the broad moral domains appear to affect moral judgment processes.

There was a main effect for trait progressivism, $F(1, 18) = 4.32$, $p = .05$, $\eta^2 = .19$, such that liberals ($M = 4.87$, $SE = .04$) judged all statements as more moral than conservatives ($M = 4.95$, $SE = .04$). There was no interaction effect, although it approached significance, $F(1, 18) = 3.28$, $p = .08$, $\eta^2 = .15$. Again, this highlights the differences between individuals high and low on trait progressivism.

In order to investigate the relationship between moral domain and valence more closely, a 2 Trait Progressivism (liberal, conservative) x 2 Broad Moral Domain (autonomy, community) x 2 Moral Valence (positive moral, negative moral) ANOVA was conducted. Besides the main effects for broad moral domain $F(1, 18) = 10.51$, $p < .01$, $\eta^2 = .36$, and moral valence $F(1, 18) = 478.11$, $p < .01$, $\eta^2 = .96$, already discussed, there was a two-way interaction effect. That is, positive autonomy statements ($M = 7.56$, $SD = .56$) were judged as more moral than positive community statements ($M = 7.43$, $SD = .65$), and negative autonomy statements ($M = 2.09$, $SD = .54$) were judged less moral than negative community statements ($M = 2.50$, $SD = .59$), $F(1, 18) = 10.53$, $p < .05$, $\eta^2 = .36$. This finding seems to indicate that messages in the autonomy domain are considered more extremely moral than messages in the conservative domain. That is, violations in the autonomy domain are judged more immoral than violations in the community domain, and moral behaviors in the autonomy domain are judged as more moral than behaviors in the community domain.

fMRI data

Morally relevant versus morally irrelevant results. A 2 Trait Progressivism (liberal, conservative) x 2 Moral Valence (moral, morally irrelevant) ANOVA was run to examine the main effect of moral relevance, trait progressivism, and the interaction between the two on brain activation. Results indicate that there is a distinct pattern of activation for morally relevant acts compared to morally irrelevant acts, involving areas previous implicated in moral judgment studies such as the left SMA, the left STS, both left and right DLPFC, left operculum, left amygdala, and right temporal pole (See Table 5). Additionally, the left thalamus, left precentral and poscentral gyri, right cerebellum/fusiform gyri, and left occipital gyri showed greater activation for moral versus morally irrelevant stimuli. These findings indicate that the moral judgment network was activated for messages containing moral information.

Morally irrelevant stimuli preferentially activated the left anterior STS and left paraoccipital zone.

Table 5

Brain Regions Showing a Significant Main Effect of Moral Relevance (Moral vs. Morally Irrelevant)

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	x	y	z		
Moral > Morally Irrelevant					
SMA	-4	-4	56	2288 ¹	L
THA	-17	-10	8	1672 ¹	L
PRCG	-37	-10	56	1264 ¹	L
POCG	-35	-32	52	1207 ¹	L
CERBL/FUSIG	18	-54	-18	415 ¹	R
CERBL	27	-47	-30	409 ¹	R
STS	-40	-58	6	298 ¹	L
PINS	-41	-32	21	780	L
MCC	-13	-24	43	579	L
PRCG	46	0	33	353	R
OPRC	-41	-1	25	338	L
AMYG	-14	-3	-13	294	L
TPJ	50	-38	33	281	R
OCG	-22	-84	-7	283	L
PDLPFC	-22	-4	47	272	L
PDLPFC	35	-8	59	237	R
THA	11	-18	12	198 ¹	R
ADLPFC	-27	33	37	126 ¹	L
Morally irrelevant > Moral					
TPJ	-42	-68	27	221	L
POTZ	-37	-74	43	278	L

Note. p<.05, corrected, ¹voxelwise p<.001.

Examining the effect of trait progressivism across all trial types, liberals showed more activation of the left operculum than conservatives overall, whereas conservatives

demonstrated more activation in the right SPL/angular gyrus (see Table 6). This provides evidence that trait progressivism can identify differences among participants.

Table 6
Brain Regions Showing a Significant Main Effect of Trait Progressivism

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	X	y	Z		
OPRC	Liberals > Conservatives				
	-47	11	0	278	L
SPL/ANGG	Conservatives > Liberals				
	26	-65	42	177 ¹	R

Note. $p < .05$, corrected, ¹small-volume correction applied.

There was also a significant interaction effect (see Table 7). Liberals, compared to conservatives, demonstrated increased activation in the following areas during moral trials, after subtracting the activation present in the morally irrelevant trials: the left lingual gyrus/fusiform area, the precuneus, the superior medial DLPFC, the left superior parietal DLPFC, the left SPL/TPJ, the left IPL, left operculum/insula, right insula, and left anterior and posterior STS. This supports the notion that, not only are liberals and conservatives processing messages differently, but that these differences are particularly salient for judgments of moral messages.

Table 7

Brain Regions Showing a Significant Morality x Trait Progressivism Interaction

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	x	y	z		
Liberals (moral-morally irrelevant) > Conservatives (moral-morally irrelevant)					
LNGG/FUSIG	-15	-69	0	1141	L
PRCN/MCC	2	-38	39	1071	R
SMDLPFC	6	35	37	751	R
SDLPFC	-24	16	52	707	L
PRCN	7	-62	31	373	R
SPL/TPJ	-53	-38	41	254	L
MCC	5	-19	37	231	R
PRCN	-6	-64	45	228	L
IPL	-47	-62	41	208 ¹	L
OPRC/INS	-27	30	5	199 ¹	L
SRPFC	-4	61	12	173 ¹	L
SMDLPFC	-2	16	58	167 ¹	L
ADLPFC	-36	52	17	139 ¹	L
RPDLPFC	32	15	48	125 ¹	R
PSTS	-41	-70	27	130 ¹	L
ASTS	-62	-9	-8	121 ¹	L
AINS	44	19	6	117 ¹	R

Note. $p < .05$, corrected, ¹ small-volume correction applied.

Moral Valence. A 2 Trait Progressivism (liberal, conservative) x 2 Moral Valence (positive moral, negative moral) ANOVA was run to examine the role of positive and negative moral valence on moral judgment. Positive moral statements preferentially activated areas of the right lingual gyrus, left operculum, and precuneus. Negative moral statements, on the other

hand, preferentially activated areas of the left lingual gyrus, and left medial occipital gyrus (see Table 8). These findings indicate that textual processing is occurring (as expected when using textual stimuli), as well as the novel finding that moral valence is involved in inducing a lateralization effect in the lingual gyrus.

Table 8

Brain Regions Showing a Significant Main Effect of Moral Valence (Negative vs. Positive)

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	x	y	z		
	Negative > Positive				
LNGG/FUSIG	-5	-68	0	5240	L
LNGG/FUSIG	-26	-52	-1	598	L
MOGG	-28	-78	2	209	L
	Positive > Negative				
LNGG/FUSIG	18	-52	-10	418	R
LNGG	15	-74	5	158 ¹	R
OPRC	-55	4	15	147 ¹	L
PRCN	-18	-39	36	118 ¹	L

Note. $p < .05$, corrected, ¹small-volume correction applied.

An interaction analysis and follow-up tests indicated the following regions as showing preferential involvement in negative moral trials, after subtracting the positive moral trials, for the conservative compared to the liberal group: the left superior medial DLPFC, the left superior medial PFC, and the right hypothalamus (see Table 9). This suggests that not only is the moral (versus morally irrelevant) content of the message important for identifying differences between liberals and conservatives, but also that positive and negative moral messages are processed differently for liberals versus conservatives.

Table 9

Brain Regions Showing a Significant Moral Valence x Trait Progressivism Interaction

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	X	y	Z		
Conservatives (negative-positive) > Liberals (negative-positive)					
SMDLPFC	-10	46	41	968	L
SMPFC	-16	10	61	437	L
HPThal	-3	-10	-3	359	L

Note. $p < .005$, corrected. ¹ small-volume correction applied.

Broad Moral Domain. Next, to examine the role of broad moral domains in moral judgment, a 2 Trait Progressivism (liberal, conservative) x 2 Broad Moral Domain (autonomy, community) ANOVA was conducted (see Table 10). The the right precuneus and the left calcarine cortex showed preferential involvement in autonomy versus community trials. For community versus autonomy trials, the superior medial DLPFC was preferentially activated. This indicates that broad moral domains activate different aspects of the moral processing network. It further suggests that these differences may lay in the degree of self versus other processing that occurs when processing messages in each domain.

Table 10

Brain Regions Showing a Significant Main Effect of Broad Moral Domain

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	X	y	Z		
Autonomy > Community					

Table 10

Brain Regions Showing a Significant Main Effect of Broad Moral Domain (con't)

PRCN	-4	-65	33	215	L
LCLC	-1	-58	12	148 ¹	L
Community > Autonomy					
SMDLPFC	-12	23	51	160 ¹	L

Note. $p < .005$, corrected. ¹ voxelwise $p < .001$.

The same contrasts were used to examine the role of trait progressivism. The right temporal pole, the left precuneus, and the left superior medial DLPFC showed preferential involvement in autonomy trials, after subtracting community moral trials, for the liberal, as compared to the conservative, participants (see Table 11). This offers more support for the notion that both broad moral domain and trait progressivism are important for processing moral messages.

Table 11

Brain Regions Showing a Significant Trait Progressivism x Broad Moral Domain Interaction

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	x	y	z		
Liberals (Autonomy-Community) > Conservatives (Autonomy-Community)					
TPL	47	17	-11	213	R
PRCN	-2	-55	22	215	L
SMDLPFC	-14	39	51	206	L

Note. $p < .005$, corrected.

Finally, examining the 2 Contrast (negative-positive within autonomy, negative-positive within community) x 2 Trait Progressivism (liberal, conservative) results, there was a main

effect for the contrast composite. The difference between negative and positive acts, within the autonomy domain, preferentially activated the left IPL when compared to the same contrast within the community domain (see Table 12). An interaction effect with trait progressivism was also significant. This same contrast activated the right IPL for liberals, when compared to conservatives (see Table 13). This again suggests that differences between the broad domains is involved in the self-other differences between the domains. It also suggests that the MFQ is accurately identifying differences in these domains along content lines relevant for distinguishing groups that process moral information differently.

Table 12

Brain Regions Showing a Significant Main Effect of Moral Contrasts (Negative-Positive Autonomy > Negative-Positive Community)

Brain Region	Talairach Coordinates			Size (mm) ³	Hemi
	<i>x</i>	<i>Y</i>	<i>Z</i>		
Negative-Positive Autonomy > Negative-Positive Community					
IPL	-40	-50	44	351 ¹	L

Note. (p<.05, corrected) ¹ voxelwise p<.001.

Table 13

Brain Regions Showing a Significant Contrast (Negative-Positive Autonomy > Negative-Positive Community) x Trait Progressivism Interaction

Brain Region	Talairach Coordinates			Size	Hemi
	x	y	Z	(mm) ³	
Progressives > Conservatives					
IPL	34	-55	39	110 ¹	R

Note. p < .005, corrected. ¹ small-volume correction applied.

Discussion

This study began with multiple goals regarding better understanding the role of moral valence, moral domain, and trait differences in moral salience in the processing of moral and messages.

First, findings suggest that moral content activates distinct neural areas based on moral relevance. In line with past moral neuroscientific findings, morally relevant content versus morally irrelevant content activated a “moral judgment” network in the brain. In addition to the immoral versus morally irrelevant activation found in previous research, this study further contributes the finding that activation of the moral judgment network depends on the valence of the moral stimuli. That is, positive and negatively valenced statements elicited different neural activation among participants. This suggests the importance of including both positive and negative moral content in future research, rather than confining inquiry to negative moral content.

The second major finding of this study is that moral judgments, as well as neural activation, vary based on the domain of the moral behavior. This corroborates mounting evidence against morality being a homogenous construct. Furthermore, results support theoretical distinctions between autonomy and community domains in that autonomy domains deal with predominantly self-centered versus other-centered moral judgments, whereas community domains deal with other-centered judgments versus self-centered judgments. Results also suggest that behavior in the autonomy domain is more salient as “moral” behavior than behavior in the community domain. These results are important for confirming the

salience of domains in moral processing, as well as highlighting conceptual differences between major moral domains.

The third major finding of this study is that both moral judgments and neural activation vary based on participants' self-reported salience of moral intuitions. That is, liberal and conservative participants, as distinguished by scores on the MFQ, showed distinct patterns of neural activation across all conditions. Importantly, this indicates that the MFQ can be used to detect important differences in moral judgments between participants. Additionally, it suggests that liberals and conservatives use different processes when judging moral and immoral behavior.

The following section describes these findings in detail, beginning with the moral valence results, moving on to the moral domain results, and finishing by describing the findings relevant to individual differences in moral processing. Areas for future research along these lines are also described, as well as caveats to the current study. Finally, the overall implications of this study are discussed.

Moral Relevance

The findings for the impact of moral relevance on neural activation are in line with previous work examining the effect of moral versus morally irrelevant stimuli on brain activity. There was a distinct pattern of activation for all morally relevant acts compared to all morally irrelevant acts involving the following areas: the left SMA, STS, left and right DLPFC, left operculum, left amygdala, right cerebellum, left thalamus, left insula, and right TPJ. All the areas above have been implicated in the “moral judgment network” identified by Moll et al.

(2005b) and have been identified as being strongly involved in moral and emotional processing in past research.

For example, the left SMA has been implicated in decision making under pressure, as well as judging harmful transgressions (Parkinson et al., 2011). The STS has shown to be involved in the processing of prosocial emotions, as well as judging moral versus non-moral behaviors, as has the left operculum and left amygdala (Moll et al 2005b; Moll et al 2002; Schiach Borg et al., 2006). The TPJ has been implicated in person-perception judgments relevant to moral decision making (Young & Saxe, 2009). Both the left and right DLPFC have been implicated in most moral neuroscience. Specifically, the left DLPFC has been shown to be involved in making decisions regarding the perceived responsibility of acts (Buckholz et al. (2008), moral versus neutral images (Schiach Borg et al., 2006), compliance to social norms (Spitzer, Fishbacher, Herrnberger, Gron, & Fehr, 2007), broad involvement in social judgment and social dilemmas (Heereken et al., 2003, 2005), judgment of moral harm (Greene et al., 2001), and harmful as well as dishonest moral transgressions (Parkinson et al., 2011). The right DLPFC has been implicated in the processing of emotional information during moral judgment (Glenn et al., 2007). Although not implicated in all moral studies, the right cerebellum has shown up in studies of disgusting stimuli (Parkinson et al., 2011) as well as indignation (Moll et al., 2002; Moll et al., 2005a). The left insula has been implicated in studies of disgusting stimuli (Moll et al., 2005a) as has the left thalamus (Moll et al., 2002).

Taken together, these results suggest that the manipulation of morally relevant versus morally irrelevant stimuli was successful, with morally relevant stimuli preferentially activating regions of interest thought to be associated with moral and emotional processing. Despite the

brevity of the stimulus materials, participants appeared to process the information using the moral networks identified in past research. Furthermore, although the moral judgment ratings demonstrated a clear and expected linear pattern for moral ratings, both types of moral valence (positive and negative) activated this moral judgment network when compared to morally irrelevant statements.

Positive and Negative Moral Valence

When negative and positive moral statements were compared directly, there was distinct neural activation for each type of moral statement based on moral valence. Negative statements preferentially activated areas of the left lingual gyrus/fusiform gyrus, involved in word processing and face recognition, and the left medial occipital gyrus, also involved in language recognition. Positive statements, on the other hand, preferentially activated areas of the right lingual gyrus, involved in visual processing. Both the right and the left lingual gyri have been implicated in emotional and face processing (Moll et al., 2002a) as well as moral processing (Schaich Borg et al., 2006) in addition to word processing.

Although these results may be a function of using text-based stimuli, the bilateral effect for positive versus negative statements has not been seen previously, and may prove interesting for further exploration. Positive statements also preferentially activated the left operculum, which is part of a larger “empathy network”, and the precuneus, which is involved in self reflections (Kjaer, Nowak, & Lou, 2002), visuospatial processing (Kawashima, Roland, & O’Sullivan, 1995), self consciousness (Vogeley et al., 2004) self-reflection (Johnson et al., 2006) and episodic memory and self agency (Cavanna & Trimble, 2006). This seems to indicate that

positive moral acts preferentially activate an emotional or self-reflective response compared to negative moral acts. This is a tentative finding, however, and requires future investigation.

Understanding the neural response to positive moral stimuli is especially important for media researchers, as disposition theory suggests that the positive outcomes experienced through viewing media (i.e., enjoyment) result from viewing characters receiving morally-appropriate rewards and punishments (Zillmann, 2000). The appropriateness of the outcome is based primarily on viewers' moral approbation of characters' behaviors, and the emotional involvement viewers feel with characters. Although past research on person perception has touched upon the importance of positive moral information in person perception, moral cognitive neuroscience has been based in understanding the moral judgments of negative behaviors (although see the discussion on moral beauty by Takahashi et al., 2008). If future research does support the notion that positive behaviors evoke greater emotional response than negative behaviors, it could pave the way for greater understanding about the mechanisms behind character involvement and empathic reactions. Past research has shown a definite "negativity effect," such that negative information about character morality is more salient than positive when making assessments about character liking and morality (Tamborini, Weber, Eden, Bowman, & Grizzard, 2010). However, results from the current study suggest there may be a corresponding "positivity effect" in which viewers may have greater emotional involvement for positive moral behavior. This would then increase their empathic response towards characters acting in a positive manner, and help explain the great enjoyment viewers feel when seeing moral characters rewarded in narrative.

It may be that the focus on immoral behavior in media research is due to the messages that have traditionally been the focus of this research. For example, disposition theories have focused on suspenseful narrative and violent drama, where antisocial behavior provides the main source of conflict for the characters. The focus on the immoral behaviors of villains (rather than the moral behavior of heroes) is more central to this type of narrative. However, in other types of narratives, such as tragedy, we may see different motivating forces. For example, recent research on the sharing of stories in the New York Times demonstrated that the stories shared the most are those featuring “elevating” content (Berger & Milkman, 2010). The findings from the current study regarding positively valenced moral messages suggest that there is a distinct neural response to these types of elevating messages that is distinct from other types of moral message.

Although these findings are intriguing, and certainly important for understanding responses to the valence of moral behaviors, they do not address more interesting questions about morality such as whether or not morality is homogeneous or based in separate domains. Therefore, further analyses examined the domain-specificity of moral judgment using the broad moral domains put forth by Shweder et al. (1997).

Broad Moral Domains: Autonomy and Community

Examining results grouped by domain, autonomy statements were judged overall as less moral than community statements, although this is most likely due to the extremely low moral ratings for the ‘unfair’ statement ($M = 1.81$) in the autonomy domain compared to the ratings for the other negative moral statements (harm = 2.38, disrespect = 2.63, disloyal = 2.37). In fact, overall autonomy statements were judged more extremely than community statements. That

is, positive autonomy statements were judged as more moral than positive community statements, and negative autonomy statements were judged less moral than negative community statements. This may indicate that for all participants, harm and fairness are more salient in terms of morality than group loyalty and authority.

Examining the domain-specific neural data, there was a main effect for domain such that autonomy statements preferentially activated the right precuneus whereas community statements preferentially activated the superior medial DLPFC. As previously discussed, the precuneus is implicated in self-reflection processes, whereas the superior medial DLPFC is more often implicated in mediating social judgment processes outside the self. These findings are in line with theoretical reasoning by Shweder et al. (1997) and Haidt and Kesebir (2010) suggesting that autonomy domains are related to the self, whereas community domains are related to social environment. Shweder et al., (1997) specifically describes this dichotomy as that between the autonomous self (separate from society) and the connected self (a part of society). This suggests that domain-specific inquiries into moral processing should explore the self-other dichotomy between these domains, in addition to pure content differences.

The indication that autonomy judgments involve self-referential memory also supports the notion that harm and fairness judgments are central to individual conceptualizations of morality. This emphasis on harm and fairness being the primary concerns of morality is reflected in the preoccupation of moral psychologists with researching moral violations in the autonomy domain (Haidt & Joseph, 2004; Haidt & Kesebir, 2010). Indeed, moral foundations theory (as well as the cultural morality modal proposed by Shweder et al., 1997) was proposed as a counter to overwhelming preoccupation of moral psychologists with harm and fairness

concerns.⁴ Haidt and Kesebir (2010) argued that the (predominantly liberal) academy was ignoring moral concerns that were less salient to the researchers investigating moral phenomena. However, it may simply be that the academy was reflecting a true underlying bias towards harm and fairness in moral cognition. This dominance of harm and fairness can also be seen in the distribution of moral salience across liberals and conservatives. Both liberals and conservatives value harm and fairness in making moral decisions; however conservatives *additionally* value the community domains of group loyalty and authority (Graham, Haidt, & Nosek, 2010). Recently, Wright and Baril (2011) supported the notion that harm and fairness are central to the concept of morality. In their experiment, under conditions of cognitive load and self-regulation depletion, conservatives de-emphasized considerations of group loyalty, authority, and purity, but maintained their considerations of harm and fairness. Wright and Baril (2011) suggested that conservatives and liberals both start from a moral system focused on harm and fairness issues, and that conservatives broaden from those foundations based on situational affordances.

The centrality of harm and fairness to morality is reflected in media psychology research as well, where most research examining morality does so focusing on violence (Raney & Bryant, 2002; Raney, 2005), justification (Zillmann & Bryant, 1974), or both (Tamborini et al., in press). This preoccupation with harm and fairness is also featured in news analyses (Ehrlich, Weller, & Eden, 2005; Kerbel, 2000). As the messages in the current study were purported to be common themes from news stories, the centrality of harm and fairness may simply reflect the

⁴ Harm/Care is also a relatively recent addition, added in the early 1980s by Gilligan (1982). Prior to this point, moral psychologists had been concerned primarily with moral behavior in the fairness domain.

dominance of harm and fairness concerns in news. It may well be that framing the study as an analysis of entertainment or comedic messages, rather than news stories, would have an effect on which statements were judged more strongly moral or immoral.

Indeed, recent research applying moral foundations to media, has broadened the focus away from harm and fairness concerns (c.f. Tamborini, 2011). For example, research on heroes and villains as moral exemplars suggests that authority and group loyalty are key components in understanding the attraction of anti-heroes (Eden, Oliver, Tamborini, Woolley, & Limperos, 2009; Tamborini, Grizzard, Eden, & Lewis, 2011). The study of genre-specific content such as soap opera (Tamborini, Enrique, Lewis, Grizzard, & Mastro, 2011) suggests that different genres are concerned with addressing issues in specific moral domains. It may be that the application of broader moral concerns is only relevant for specific viewers, though, or under specific conditions (such as defining protagonists as anti-heroes). In order to address issues such as these, the final analysis examined differences in moral processing between liberal and conservative participants.

Trait Progressivism

In order to examine the influence of moral intuitions on moral judgments, participants' scores on the MFQ were combined to form a trait progressivism score. Participants high on this scale were considered liberal, and participants low on this scale were considered conservatives. Across all moral judgments, liberals judged statements as more moral than conservatives. Looking at the neural data, across all trial types (both moral and morally irrelevant), liberals showed more activation of the left operculum than conservatives, whereas conservatives demonstrated more activation in the right SPL/angular gyrus. The operculum has been

implicated in the processing of basic emotions, such as joy or sadness, when compared to moral statements (Moll et al., 2005; Takahashi et al., 2009). It has also been associated with empathic reactions (Eslinger, Moll, & deSouza, 2002; Decety, 2004). The right SPL/angular gyrus, on the other hand, has been associated with negative outgroup judgments (Cikara & Fiske, 2011) as well as more general negative moral judgments (Greene et al., 2001; Greene et al., 2004; Raine & Yang, 2006).

This indicates that perhaps liberal participants were relying more on emotional concerns while performing the experimental task, whereas conservative participants may have been more concerned with making the judgments themselves. This effect in the self-report data appears to be mainly driven by the extremely low moral ratings conservatives gave to the harm statement, however this is consistent with the neural findings as well. Including the moral domains, however, liberals showed preferential activation for moral processing regions during autonomy versus community trials: The right temporal lobe and the left superior medial DLPFC were activated. This suggests that inclusion of valence and domain, as well as trait progressivism, is required to understand the interaction of moral intuitions and moral judgments.

Looking more closely at this three way interaction, the following regions showed preferential involvement in moral versus morally irrelevant trials for the liberal compared to the conservative group: the left nucleus accumbens and fusiform face area, the precuneus, the superior medial DLPFC, the left TPJ, left operculum, left IPL, middle cerebral cortex, the anterior and posterior DLPFC, the left posterior STS, and the right insula. Interestingly, these areas are the same ones implicated in the comparison of harmful and dishonest moral transgressions

with neutral behaviors found by Parkinson et al., (2011), as well as the areas implicated in the overall moral versus morally irrelevant analysis. This suggests that liberals, more so than conservatives, may be responding to all moral behaviors using the moral judgment network used to judge behaviors in the autonomy domains.

In contrast, conservatives seem to use this network primarily when considering positive versus negative moral judgments. For example contrasting positive and negative moral judgment, conservatives, more so than liberals, showed preferential activation for the left superior medial DLPFC, the left superior medial PFC (generally considered part of moral judgment networks) and the right hypothalamus. Activity in the hypothalamus has been shown to correlate with exposure to violations of social values (Zahn et al., 2008) and to processing of moral issues related to care (Robertson, Snarey, Ousley, Harenski, Bowman, Gilkey, & Kilts, 2007). These results indicate that conservatives show more activation in moral judgment areas than liberals in response to moral violations. This might be a reflection of previous research suggesting that conservatives are more sensitive to threats to social order and moral infractions in general than liberals (Jost, Federico, & Napier, 2009).

Combining all analyses, an interaction effect between trait progressivism, moral valence and broad moral domain was found such that liberal participants had less variance between their judgments of both positive and negative autonomy and community statements, whereas conservative participants had greater variance in autonomy than community judgments. This effect was mirrored in the neural data. When looking at trait progressivism, broad moral domain, and moral valence (negative, positive) in the same analysis, for the negative – positive contrasts, the left inferior parietal lobule was activated for autonomy versus community trials.

The left inferior parietal lobule has been implicated in studies of agency (Chaminade & Decety, 2002), when the subject is *not* the agent of the action. When the subject is the agent of the action, however, the right inferior parietal lobule is activated. Activation in the right IPL can also show up when subjects are making a self-other distinction (Uddin, Molnar-Szakacs, Zaidel, & Iacoboni, 2006). In this study, the right IPL was activated for liberals during negative autonomy trials, despite general activation of the left IPL for all negative autonomy trials, suggesting that liberals were making a self-other distinction during autonomy moral violations that conservatives did not make.

These results reveal fascinating differences in how liberals and conservatives respond to moral information. Liberals appear especially sensitive to domain context, whereas conservatives are more sensitive to the valence of the behavior. When considering media effects, this finding has important implications for the types of behaviors viewers may be attending to and learning from media. For example, social cognitive theory proposes that individuals may learn social norms vicariously through media exposure. Results from the current study might suggest that liberals may be more prone to attend to (and potentially learn from) behaviors in the harm and fairness domains, and be relatively unaffected by behaviors in the group loyalty and authority domains. Conversely, conservatives may be more prone to attend to and learn from moral violations across all domains, and less affected by viewing moral behaviors than liberals. For example, health campaign or political messages targeting liberals may well emphasize the harmful or unfair issues with adopting particular legislation, but may be equally well served in addressing the caring or fairness related issues as well. Conversely, conservatives may respond well to negative information across all domains.

Although speculative, understanding these types of effects stemming from differences in moral processing would be incredibly helpful for designers of campaign and edutainment materials, and could also provide greater clarification into understanding pro- and anti-social effects from media exposure.

Limitations

This study was designed to be a broad exploration of moral valence, moral domain, and moral intuitions as they relate to moral processing in the media arena. Although successful in identifying areas of the brain involved in moral processing, there are some limitations that may hinder the broad applicability of this study to work on media and morality.

First, due to the focus on media theory, the reliance on text stimuli is a limitation that should be addressed in future research. Although results indicated that the textual stimuli were processed using the moral judgment network found in past research, in order to be more relevant to media work it would be best to examine these phenomena further using audio-visual stimuli. Additionally, the effects found for moral valence involved word recognition areas, therefore it would be important to replicate these effects using non-textual stimuli. A second limitation was the reliance on a tertiary split of a student sample for the selection process. Previous research indicates that students may be limited in the amount of variance displayed in moral judgments, and thus even a tertiary split would not provide the needed variance between groups. However, given the robust findings in this area, the selection process appears to have been justified. The use of an all-female sample may also be a limitation, as women tend to be more sensitive to harm violations than men (Schiach Borg et al., 2006). Future research should examine these effects in male as well as female subjects in order to

broaden the findings. Third, the focus on news stories as opposed to other types of genre may limit the applicability of these effects. Further investigations in this line of research should address moral messages from a broad range of media genres. Fourth, two of the messages used as morally irrelevant (success and failure) were rated similarly to moral messages along dimensions of intentionality and arousal. As stated, these messages were included to provide emotionally-valenced but morally irrelevant contrasts to the “true neutral” messages (get rid of, collect) in future analyses. However, collapsing them here with the other morally irrelevant messages may have led to less distinction of moral emotional and morally irrelevant networks than could be achieved.

The majority of these limitations were understood prior to this investigation. They were accepted as trade-offs for the advantages offered by the design selected for this investigation. In this regard, they can be thought of less as problems with the current study, but rather areas of need for future investigations in this line of research. To that end, we now turn to discussing future directions for research in the three areas implicated by this study: Moral cognitive neuroscience, moral psychology, and media psychology.

Future Directions

This study supports past findings in moral cognitive neuroscience that suggest moral processing is associated with a specialized network of neural activation. It also addresses the relative lack of findings dealing with positive moral messages in past research. Furthermore, it suggests that the neural substrates for judging behaviors in different content domains are distinguishable using imaging techniques. Future analyses using the data collected in this study will focus on isolating these networks at a more specific domain level (such as the specific

moral domains proposed by Haidt & Joseph, 2008) and contrasting specific activations for messages within each moral domain with morally irrelevant messages. This will allow for a fuller understanding of the mechanisms behind specific moral domains, rather than the broad domains discussed in the current study. The person versus action contrast will be included to examine differences between judging people and their behavior. Extensive research suggests that there are distinct processes for judging people and action behaviors. Isolating the contributions of each type of judgment may help parse out which areas of the brain are activated due to moral judgment versus other types of judgment. Finally, the moral emotions will be regressed onto the results from these analyses in order to determine to what extent moral emotions shape our moral judgments at a neural level. These future studies will help explicate the different neural substrates involved with each determinant of moral judgment, in order to attain a much more complete picture of how the brain processes moral information.

In the area of moral psychology, the main contributions of this study are in supporting the domain-specific model of morality proposed by Shweder et al., (1997) and Haidt and Joseph (2008). The suggestion that there is a self-other dichotomy behind these different broad domains is one that should be explored more fully in the future by manipulating self and other judgments within these domains. Also, the finding that autonomy domains may be considered more strongly “moral” is one that begs for future investigation. Is this effect dependent on context? Is it dependent on the individual makeup of the person making the judgments? Is this a function of a lay understanding of morality or does it echo deeper truths about the importance of the respective domains? Future research should attempt to distinguish the relative importance of the domains in different contexts and for different

participants. Finally, this study offers support for the validity of the MFQ as an instrument with which to distinguish different groups of people who vary on moral judgment processes. Future research should examine the extent to which the MFQ is correlated with other measures, in order to build a better picture of the traits and qualities that impact moral judgment. For example, Jost et al., (2009) found that self-reported liberalism and conservatism varied with tolerance for ambiguity, openness to experience, threat perception, and uncertainty avoidance. It would be interesting to measure these scales along with the MFQ in order to attempt to determine the extent to which moral salience drives these individual differences, versus the extent to which individual differences affect the salience of moral intuitions.

Finally, in terms of media psychology, this study demonstrated that moral messages can impact moral judgments of others and their behaviors. Specific implications in the need to examine both positive and negative moral judgments in media texts, as well as focusing on the role of all moral domains in media research, have been discussed previously in the discussion section. However there is room for much exploration in this area. For example, future studies could vary the source of the moral message in order to examine potential genre effects. It may be that news messages are particularly good at evoking moral responses because they deal with real people; whether fictional characters would evoke the same response would be worth investigating. Abraham, von Cramon, and Schubotz, (2008) examined the role of fictionality in neural responses to narrative, finding that fictionality elicits a “distancing” effect when compared to real people. By replicating this fictionality effect in a moral domain, media psychology would be able to specifically investigate the type of neural response elicited from viewing fictional narratives. In addition, future work should use media-specific stimuli such as

persuasive messages to examine the effect of individual differences in attending to such messages. Although we might speculate how moral salience affects responses to these messages, it would be worth investigating the specific effect of morality on the processing of these types of texts. Although there has been recent work in entertainment theory focusing on the role of moral salience in responses to media, branching into persuasion would allow for an examination of cognitive and emotional influences on media message processing that are affected by moral judgment processes. This could gain us fuller understanding of the role of morality and moral framing in persuasive texts.

Conclusions

In conclusion, this study was a broad, exploratory examination of several concepts important to moral neuroscience, moral psychology, and media researchers. As such, it generated almost as many questions as it answered. However, reexamining the initial goals of the study, it is clear that moral information has an effect on moral judgments. This process is affected by moral valence, broad moral domains, and the salience of moral intuitions relevant to the behavior. Taken together, these results are promising for several reasons. First, this study supports and extends existing research in moral cognitive neuroscience by offering more evidence for a distinct moral processing network. It extends moral psychological theory by offering support for a domain-specific model of morality. Also, it allows media theorists to expand our understanding of the media experience beyond what is available through existing self-report or behavioral measures. Thus we are able to gain a fuller picture of how viewers respond to narrative content.

Results from this study also suggest the MFQ is accurately capturing a robust difference in moral processing between subgroups using self-report data focused on moral intuitions. Despite the limitations of self-report data, especially for measuring something as elusive as moral intuitions, it appears that the MFQ is at least reflecting processing differences salient for moral judgments at the neurological level. These neural findings are hugely important to help validate the measure, especially given the sometimes weak or inconsistent findings provided by self-report and reaction-time data. Additionally, these findings lend support to the media work suggesting that differences among the moral domains (as measured by the MFQ) can be used to examine “morality subcultures” among viewers (Eden and Tamborini, 2010; Tamborini, 2011; Lewis, Grizzard, Eden, & Tamborini, 2011).

Finally, these results suggest that greater understanding of media theory can be attained through incorporating neuroscientific findings into research on media and morality. Although this science is still in its initial stages, there is great room for expanding and testing moral media theory through neuroscience. These results have particular import for established theories of media and morality, such as disposition theory and social cognitive theory. In addition, this study supports newer models, such as the model of media exemplars and moral intuitions (Tamborini, 2011). This model suggests that moral intuitions can be primed by exposure to media content, resulting in short- and long-term effects on morality and behavior. This study lays the groundwork to neuroscientific studies testing these theories and model, which will in turn greatly expand our understanding of media processes.

APPENDICES

Appendix A

Person Perception Processes

MICHIGAN STATE UNIVERSITY, DEPARTMENT OF PSYCHOLOGY CONSENT TO ACT AS A RESEARCH SUBJECT

Investigators

Dr. Issidoros Sarinopoulos, sarinopo@msu.edu, 517-290-4327, 171B Radiology Bldg., MSU, East Lansing, MI; Allison Eden, edenalli@msu.edu, 517-614-5314, 552 CAS, East Lansing, MI.

Purpose of Research

The focus of the research study is to gain a better understanding of the neural underpinnings of person perception processes.

Procedures to be followed during this study

There are two phases to this study. There are two phases to this study. The first phase of this study is an online survey of opinions and attitudes regarding various social issues. If you agree to participate in this phase of the study, you will be asked to complete several sections of different questionnaires. This survey will take about 30 minutes to complete. Following the completion the online survey, your voluntary agreement to participate in the second phase of the study, which is a functional magnetic resonance imaging (fMRI) study, may be solicited. If you are solicited, you will fill out a secondary consent form for the fmri portion of the study.

Risks of the study

The first phase of this study consists of an online survey and involves no foreseeable serious risks. However, you will be asked a series of survey questions about your opinions and attitudes regarding various social and political issues. These questions can sometimes make people uncomfortable. You do not have to answer any questions that you do not wish to answer. Just skip to the next question.

Benefits of the study

Participation in the first phase of the study will not provide any direct benefit to you. However, you will receive class credit as a result of your participation. Please note that research participation is not a course requirement as you can earn class credit in alternative ways that do not require research participation, such as submitting a written paper. Please see your instructor for alternate assignments. Your participation in this phase of the study may help contribute to scientific knowledge and theory.

Based on the results of the survey questions in the first phase of this study, you may be solicited to participate in the second phase of this study. Participation in this phase of the study does not guarantee any beneficial results to you. However, you will be paid up to a maximum of \$40 for your participation if you complete the entire second phase of this study. If you decide to

quit early you will receive \$20/hour for your time. Your participation in this phase of the study may help us learn more about how the brain functions thereby gaining scientific knowledge that may help other people in the future.

Your rights as a participant in this study

As a participant in this experiment, you have certain rights of which you should be aware.

1. First, even if you agree to participate, you have the right to change your mind and to not participate before the experiment begins, and to do so without penalty.

2. Second, even if you agree to participate and have begun the experiment, you have the right to discontinue your participation at any time and for any reason, and to do so without explanation or penalty. This will be the case whether you decide to discontinue participation during the survey or the fMRI part. If you decide to discontinue participation, all questionnaires and other information about you will be destroyed, and your name will be deleted from our records. However, your consent form will be saved for Institutional Review Board records. Your participation may be ended without your consent if: a. the investigator believes that it is in your best interest. b. the project is terminated. c. you no longer meet study criteria.

3. Third, you have the right to confidentiality. To ensure this, you will be assigned a subject number, which will be the only way to identify you in any reports about the study. The investigators listed on the project will have access to files and lists that can be used to link a name with a subject number. Additionally, any member of the Institutional Review Board may have access to files and lists. Any records identifying study participants will be destroyed three months following data collection. De-identified data will be stored in a secure computer file for a period of 5 years after the last publication from this project. The Michigan State University's Human Research Protection Program could have access to all files associated with this study. Your privacy will be protected to the maximum extent allowable by law.

If you are injured as a result of your participation in this research project, Michigan State University will assist you in obtaining emergency care, if necessary, for your research related injuries. If you have insurance for medical care, your insurance carrier will be billed in the ordinary manner. As with any medical insurance, any costs that are not covered or are in excess of what are paid by your insurance, including deductibles, will be your responsibility. The University's policy is not to provide financial compensation for lost wages, disability, pain or discomfort, unless required by law to do so. This does not mean that you are giving up any legal rights you may have. You may contact Dr. Issidoros Sarinopoulos (517-290-4327, sarinopo@msu.edu) with any questions or to report an injury.

If you have questions or concerns about your role and rights as a research participant, would like to obtain information or offer input, or would like to register a complaint about this study, you may contact, anonymously if you wish, the Michigan State University's Human Research Protection Program at 517-355-2180, Fax 517-432-4503, or e-mail irb@msu.edu or regular mail at 207 Olds Hall, MSU, East Lansing, MI 48824.

Questions about the study

After the study is completed if you have any questions or if you would like a written summary of the general results, you are invited to contact the investigators at their University offices (Dr. Sarinopoulos: 517-884-3283, sarinopo@msu.edu; Allison Eden: 517-614-5314, edenalli@msu.edu).

Acknowledgement of consent

If you agree to participate in this study, please click "Next Page". On the following page you will be asked to enter your Research ID. This will be kept separate from your data in a locked file, and will only be used to a) contact you for participation in the second part of the study and b) to provide credit to your course instructors.

If you agree to participate, please click the "Next Page" button below to begin.

Moral Foundations Questionnaire

Part 1. When you decide whether something is right or wrong, to what extent are the following considerations relevant to your thinking? Please rate each statement using this scale:

[0] = not at all relevant (This consideration has nothing to do with my judgments of right and wrong)

[1] = not very relevant

[2] = slightly relevant

[3] = somewhat relevant

[4] = very relevant

[5] = extremely relevant (This is one of the most important factors when I judge right and wrong)

1. _____ Whether or not someone suffered emotionally
2. _____ Whether or not some people were treated differently than others
3. _____ Whether or not someone's action showed love for his or her country
4. _____ Whether or not someone showed a lack of respect for authority
5. _____ Whether or not someone violated standards of purity and decency
6. _____ Whether or not someone was good at math
7. _____ Whether or not someone cared for someone weak or vulnerable
8. _____ Whether or not someone acted unfairly
9. _____ Whether or not someone did something to betray his or her group
10. _____ Whether or not someone conformed to the traditions of society
11. _____ Whether or not someone did something disgusting
12. _____ Whether or not someone was cruel
13. _____ Whether or not someone was denied his or her rights
14. _____ Whether or not someone showed a lack of loyalty
15. _____ Whether or not an action caused chaos or disorder
16. _____ Whether or not someone acted in a way that God would approve of

Part 2. Please read the following sentences and indicate your agreement or disagreement:

[0]	[1]	[2]	[3]	[4]	[5]
Strongly disagree	Moderately disagree	Slightly disagree	Slightly agree	Moderately agree	Strongly agree

_____ Compassion for those who are suffering is the most crucial virtue.

_____ When the government makes laws, the number one principle should be ensuring that everyone is treated fairly.

_____ I am proud of my country's history.

_____ Respect for authority is something all children need to learn.

_____ People should not do things that are disgusting, even if no one is harmed.

_____ It is better to do good than to do bad.

_____ One of the worst things a person could do is hurt a defenseless animal.

_____ Justice is the most important requirement for a society.

_____ People should be loyal to their family members, even when they have done something wrong.

_____ Men and women each have different roles to play in society.

_____ I would call some acts wrong on the grounds that they are unnatural.

_____ It can never be right to kill a human being.

_____ I think it's morally wrong that rich children inherit a lot of money while poor children inherit nothing.

_____ It is more important to be a team player than to express oneself.

_____ If I were a soldier and disagreed with my commanding officer's orders, I would obey anyway because that is my duty.

_____ Chastity is an important and valuable virtue.

Empathy Scale

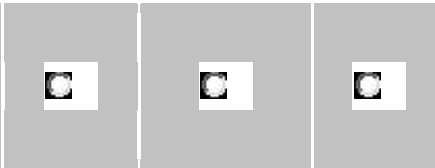
Please read the following sentences and indicate your agreement or disagreement.

	Strongly disagree	Moderately disagree	Slightly disagree	Neither agree nor disagree	Slightly agree	Moderately agree	Strongly agree
When I see someone being taken advantage of, I feel kind of protective towards them.	☐	☐	☐	☐	☐	☐	☐
Before criticizing someone, I try to imagine how I would feel if I were in their place.	☐	☐	☐	☐	☐	☐	☐
I sometimes try to understand my friends better by imagining how things look from their perspective.	☐	☐	☐	☐	☐	☐	☐
When I am reading an interesting story or novel, I imagine how I would feel if the story were to happen to me.	☐	☐	☐	☐	☐	☐	☐
I really get involved with the feelings and characters in a novel.	☐	☐	☐	☐	☐	☐	☐
I try to look at everyone's side of a disagreement before I make a decision.	☐	☐	☐	☐	☐	☐	☐

When I watch a good movie, I can very easily put myself in the place of the lead character.	☐	☐	☐	☐	☐	☐	☐
When I am upset at someone I usually try to put myself in his/her shoes for a while.	☐	☐	☐	☐	☐	☐	☐
I am often touched by things that I see happen.	☐	☐	☐	☐	☐	☐	☐
I become very involved when I watch a movie.	☐	☐	☐	☐	☐	☐	☐
I often have tender, concerned feelings for people less fortunate than myself.	☐	☐	☐	☐	☐	☐	☐
I cannot continue to feel OK if people around me are depressed.	☐	☐	☐	☐	☐	☐	☐
I would describe myself as a soft-hearted person.	☐	☐	☐	☐	☐	☐	☐
After acting in a play, or seeing a play or a movie, I have felt partly as though I were one of the characters.	☐	☐	☐	☐	☐	☐	☐
I become nervous if others around me seem nervous.	☐	☐	☐	☐	☐	☐	☐
I am the type of person who is	☐	☐	☐	☐	☐	☐	☐

concerned when
other people are
unhappy.

The people
around me have a
great influence on
my moods.



Please enter your age:

Please select your gender:



Male



Female

Please select the ethnic group that best represents your own:

Please select your country of origin:

Thank you so much for your help with our research. We just want to reiterate one thing. Based on the results of this study, you may be contacted by one of the primary researchers to participate in a functional magnetic resonance imaging (fMRI) study. Participation in the fMRI portion of the study is completely voluntary and does not affect your credit for taking part in this survey.

If you are interested in participating in the second part of the study, there will be an opportunity to sign up during your class period. If you are selected you will complete an additional signed consent form that details the risks and benefits to you as a participant in an fMRI study. fMRI studies do contain some risk to certain participants, therefore you will be asked to complete a secondary screening prior to participation to ensure your safety in the scanner.

Participation in the second phase of the study does not guarantee any beneficial results to you. However, you will be paid up to a maximum of \$40 for your participation if you complete the entire second phase of this study. If you decide to quit early you will receive \$20/hour for your time. Your participation in this phase of the study may help us learn more about how the brain functions thereby gaining scientific knowledge that may help other people in the future.

At this point you are finished with the survey. If you have any questions or concerns about the survey, please do not hesitate to contact the primary researchers; Allison Eden (edenalli@msu.edu) or Issidoros Sarinopolos (sarinopo@msu.edu).

Please click the "Submit Survey" button below to ensure we receive your responses. Thank you so much and have a great day!

Appendix B
fMRI Study of Person Perception Processes

**MICHIGAN STATE UNIVERSITY, DEPARTMENT OF PSYCHOLOGY
CONSENT TO ACT AS A RESEARCH SUBJECT**

Investigators

Dr. Issidoros Sarinopoulos, sarinopo@msu.edu, 517-290-4327, 171B Radiology Bldg., MSU, East Lansing, MI; Allison Eden, edenalli@msu.edu, 517-614-5314, 552 CAS, East Lansing, MI.

Purpose of Research

The focus of the research study is to gain a better understanding of the neural underpinnings of person perception processes.

Procedures to be followed during this study

There are two phases to this study. The first phase of this study was an online survey of opinions and attitudes regarding various social issues. If you have not completed these surveys, please notify the researcher at this time.

The second phase of the study is a functional magnetic resonance imaging (fMRI) experiment. The total amount of time required for this phase of the study is two hours, approximately one hour and 30 minutes of which will be spent in an MRI scanner. Once in the scanner you will be presented with trials that consist of people's faces as well as brief statements regarding these people. You will also be asked to indicate your feelings and perceptions of the people and their behaviors. The MRI phase will take place at the Radiology Department at MSU and will involve lying quietly inside the center of a large doughnut-shaped magnet (the MRI scanner).

If you agree to participate in the second phase of this study the following will happen:

MRI scanning

First, we will make sure that you do not have any risk factors for participating in the fMRI experiment. For example, we will ask you if you have a pacemaker, any metal pins, or other metal in your body such as metal pins in artificial joints or body piercing that you cannot remove. We cannot scan anyone who has a pacemaker or metal in their bodies. Similarly, we cannot scan a woman who is pregnant. If you are a woman of childbearing potential, you will need to take a pregnancy test. Any samples obtained for the pregnancy test will be discarded after results have been recorded. You may wear whatever clothes you like, and to be more comfortable you may want to take your shoes off. Because the electromagnet creates a strong

magnetic field, you may not take any metallic or magnetized objects into the magnet room. That includes keys, metal jewelry, wristwatch, or a belt with a metal buckle, and also credit cards. We will provide a safe place for you to store those possessions during the experiment.

Prior to entering the MRI scanner, you will be familiarized with the task by being shown examples of several trials a computer. Next, you will be asked to lie on your back on a firm but comfortable “bed” which is actually a movable dolly on tracks. Then, after we help you adjust pillows under your head and legs so that you’re completely comfortable, the “bed” is rolled into the magnet. Only the upper part of your body will be inside the scanner. The space around your head will be quite restricted, but nothing will touch your face. The scanner makes a loud noise while taking pictures of your brain, so you will be given earplugs to wear throughout the scan. This session will consist of several scans, each lasting approximately 8-9 minutes. During the MRI session we will also take pictures of your brain structure or anatomy. The MRI portion of the study will take approximately one hour and a half.

After Scan Ratings

After scanning is completed, you will be asked to rate the faces and behaviors presented during the fMRI scans. These ratings will take about 30 minutes to complete. However, the total duration of the second phase of the study, including the fMRI scans, should be two hours.

Risks of this Study

The second phase of this study involves magnetic resonance imaging (MRI). This form of imaging differs significantly from other scanning techniques used to measure structure and function of the brain. Other techniques, such as conventional X-ray, computed tomography (CT), or positron emission tomography (PET) use radiation generated by an X-ray machine or by chemical tracers. MRI does not use penetrating radiation. Instead, it uses a combination of radio frequency waves and a strong magnetic field generated by a large electromagnet to detect the distribution of hydrogen atoms in living tissues. Computers are then used to reconstruct the weak signals given off by the hydrogen atoms into high-quality anatomical images.

In summary, MRI presents no risks from ionizing radiation because no ionizing radiation is used. The magnetic coil that we will be using is approved by the Food and Drug Administration (FDA) for routine clinical uses, as well as for research purposes. There are, however, several other risks associated with the use of MRI. They are:

- 1) If the researchers or the participant are carrying any loose metallic objects, such as keys, these can be released in the vicinity of the magnet, and cause impact injuries;
- 2) For individuals wearing pacemakers or metallic prostheses, the magnetic field can induce electric currents in implanted wires in those devices;

- 3) For individuals with artificial joints, such as hip joints, the magnetic field can displace the metallic components (e.g. metal pins in the case of artificial joints) in those devices;
- 4) For individuals with electromagnetically programmed pacemakers, the magnetic field can erase the program code in these devices.
- 5) In addition, due to lack of knowledge of the effects of scans on a developing fetus, pregnant women will not be able to participate in this study.

We ensure that the 1st hazard is eliminated by removing all metallic objects (keys, watches, and the like) from the researchers and from you before the MRI experiment begins.

To eliminate the 2nd, 3rd, and 4th hazards, we do not allow anyone to participate who is wearing a pacemaker, neurostimulator, or any other implanted device.

Regarding the 5th issue, since the effects of the MRI on a developing fetus are not known, pregnant women will not be allowed to participate in this study. For that reason, if you are a woman of childbearing potential (women who are capable of becoming pregnant and are sexually active with men who are not vasectomized) you will need to take a pregnancy test (i.e., urine pregnancy test) prior to your entry into the study.

To ensure your safety, we ask all fMRI research participants to complete a screening form before entering the magnet. You will complete this form once you have signed the consent form, if you agree to participate in the study.

We want to mention two other possible risks. Certain individuals may feel “claustrophobic” once they are in the magnet. But, if after going into the magnet, you feel uncomfortable and want to stop the experiment, you can tell us, and we will immediately stop the experiment and let you decide whether or not you want to continue (see the section below on Your Rights as a Participant in an Experiment). Also, a small number of subjects perceive some dizziness as they are being moved into the scanner because they are passing through a magnetic field. This is completely harmless, but we will ask you to please tell us if you feel dizzy, and if you don’t want to continue you may of course stop at any point.

What If the MRI Reveals a Structural Abnormality in Your Brain?

There is one other consideration that we want to bring to your attention, and that is, what would we do if the scan of your brain reveals a structural abnormality of some kind? Looking for abnormalities is not the purpose of this research. Furthermore, many structural abnormalities are not clinically significant. However, any medical procedure carried out in the course of a medical check-up may turn up something that may warrant further examination. If something abnormal and potentially clinically significant shows up in your case, we are ethically bound to inform you, or if you prefer, your primary care physician. If you would want us to inform your

physician, please indicate his or her name and address (if you know it) in the space below. If you would rather we inform you, you do not have to write down anything.

(physician's name / address)

Benefits of the Study

Participation in the fMRI phase of the study does not guarantee any beneficial results to you. However, you will be paid up to a maximum of \$40 for your participation if you complete the entire second phase of this study. If you decide to quit early you will receive \$20/hour for your time. Your participation in this phase of the study may help us learn more about how the brain functions thereby gaining scientific knowledge that may help other people in the future.

Your Rights As a Participant in this Study

As a participant in an experiment, you have certain rights of which you should be aware:

1. First, even if you agree to participate, you have the right to change your mind and to *not* participate before the experiment begins, and to do so without penalty. Participation in this study is completely voluntary and your participation in this research project will not involve any additional costs to you or your health care provider.
2. Second, even if you agree to participate and have begun the experiment, you have the right to discontinue your participation at any time and for any reason, and to do so without explanation or penalty. This will be the case whether you decide to discontinue participation during the first part or the part performed at the MRI Center. If you decide to discontinue participation, all questionnaires and other information about you will be destroyed, and your name will be deleted from our records. However, your consent form will be saved for Institutional Review Board records.

Your participation may be ended without your consent if:

- a. the investigator believes that it is in your best interest.
 - b. the project is terminated.
 - c. you no longer meet study criteria.
3. Third, you have the right to confidentiality. To ensure this, from this point forward you will be assigned a subject number, which will be the only way to identify you in any reports about the study. The investigators listed on the project will have access to files and lists that can be used to link a name with a subject number. Additionally, any member of the Institutional Review Board may have access to files and lists. Any records identifying study participants will be destroyed three months following data collection. De-identified data will be stored in a secure computer file for a period of 5 years after the last publication from this project. The Michigan State University's Human Research

Protection Program could have access to all files associated with this study. Your privacy will be protected to the maximum extent allowable by law.

If you are injured as a result of your participation in this research project, Michigan State University will assist you in obtaining emergency care, if necessary, for your research related injuries. If you have insurance for medical care, your insurance carrier will be billed in the ordinary manner. As with any medical insurance, any costs that are not covered or are in excess of what are paid by your insurance, including deductibles, will be your responsibility. The University's policy is not to provide financial compensation for lost wages, disability, pain or discomfort, unless required by law to do so. This does not mean that you are giving up any legal rights you may have. You may contact Dr. Issidoros Sarinopoulos (517-290-4327, sarinopo@msu.edu) with any questions or to report an injury.

If you have questions or concerns about your role and rights as a research participant, would like to obtain information or offer input, or would like to register a complaint about this study, you may contact, anonymously if you wish, the Michigan State University's Human Research Protection Program at 517-355-2180, Fax 517-432-4503, or e-mail irb@msu.edu or regular mail at 207 Olds Hall, MSU, East Lansing, MI 48824.

Questions About the Study

If you have any concerns or questions about this research study, such as scientific issues, how to do any part of it, or to report an injury, please contact the researcher Dr. Issidoros Sarinopoulos, at the Lab of Social and Affective Neuroscience, Department of Radiology, Michigan State University, 517-884-3283, sarinopo@msu.edu).

Signature and Acknowledgment

I have received a copy of this form.

I have read this consent form and I voluntarily agree to participate.

Subject Name (Please Print)

Signature of Subject

Date of Birth

Today's Date

HAND USAGE QUESTIONNAIRE

INSTRUCTIONS: Please indicate below which hand you ordinarily use for each activity.

With which hand do you:

- | | | | |
|---|----------|-----------|------------|
| 1. draw? | (1) Left | (2) Right | (3) Either |
| 2. write? | (1) Left | (2) Right | (3) Either |
| 3. use a bottle opener? | (1) Left | (2) Right | (3) Either |
| 4. throw a snowball to hit a tree? | (1) Left | (2) Right | (3) Either |
| 5. use a hammer? | (1) Left | (2) Right | (3) Either |
| 6. use a toothbrush? | (1) Left | (2) Right | (3) Either |
| 7. use a screwdriver? | (1) Left | (2) Right | (3) Either |
| 8. use an eraser on paper? | (1) Left | (2) Right | (3) Either |
| 9. use a tennis racket? | (1) Left | (2) Right | (3) Either |
| 10. use a pair of scissors? | (1) Left | (2) Right | (3) Either |
| 11. hold a match when striking it? | (1) Left | (2) Right | (3) Either |
| 12. stir a can of paint? | (1) Left | (2) Right | (3) Either |
| 13. On which shoulder do you rest a bat
before swinging? | (1) Left | (2) Right | (3) Either |

End of Questionnaire

Appendix C

<u>Participant, Experimenters</u>		
Subject #:	Date: ____/____/____	Exam #: ____ Day of Week: ____ Time started: ____
First Name	Gender: <i>M F</i> Age: ____ Phone numbers: _____	
Main Experimenter (ME):	Other Experimenters _____	
Scarlett Doyle phone number: (517)285-3313 Allison Eden phone number: (517)614-5314		

<u>Procedure</u>	<u>Scheduled Time</u>	<u>Procedure Time</u>
E105		
Consent form/questionnaires	9:00-9:15	:15
Instruction/training	9:15-9:35	:20
MRI		
Getting set up in scanner	9:35-9:45	:10
Initial placement scans	9:45-9:55	:10
1 st Functional scan	9:55-10:05	:10
2 nd Functional scan	10:05-10:15	:10
3 rd functional scan	10:15-10:25	:10
4 th functional scan	10:25-10:35	:10
5 th functional scan	10:35-10:45	:10
6 th functional scan	10:45-10:55	:10
Anatomical scan	10:55-11:05	:10
Getting out of scanner	11:05-11:10	:5
Post-scan responses and ratings	11:10-11:30	:20
Debriefing	11:30-11:35	:5

Ahead of time

- ☐ fMRI time confirmed, participant is phone screened, scheduled, knows where to go (see Phone screen).
- ☐ Make sure scan time is confirmed with Scarlett so she knows when you will bring the participant to the control room; remind her that you'll need their set of images to give to the participant

Back Room

- ☐ Grab some pens from Sid's office for the participants to use
- ☐ 4 documents printed and prepared for subject pre-scan: ☐ Consent, ☐ MRI-Screen, ☐ Hand Use, , Receipt form
- ☐ Log into computer (username: **isarinop**, password: **nice\$work**)
- ☐ Have training PowerPoint, E-prime training files, and post-scan ratings survey ready
 - Training PowerPoint and E-prime files located in folder **MorMRI Training** on desktop
 - Training PowerPoint file: **Training PPT**

Meeting Subject

___ Go to the atrium area to greet participant.

Hi, [participant's name], thanks for coming in. We really appreciate your help...

___ Take them to back room.

Please come with me and we can get you started on your paperwork.

___ Administer **Consent Form** and answer any questions.

Before getting started, please take as much time as you like to review this form and sign below if you agree to participate.

As explained in this form, we will first train you on the ratings task that you'll be performing in the MRI. Then we will take you to the MRI area to perform six 5-minute scans during which you will make ratings about people and their behaviors.

Before we get you out of the MRI scanner, you'll have one last anatomical scan. During this scan you can simply sit quietly and relax for about 10 minutes. Then, when your MRI scan is over, we'll have you complete a few post-scan ratings as well as complete some short answer responses. Do you have any questions that I can answer at this time?

After participant signs consent form

I know that you have already completed most of these items during the phone screen but we need a hard copy for our records. Can you please complete the following quick physical screens for the fMRI procedure?

___ Administer **MRI screen** [not sure about some item(s)? Note and ask Scarlett later]

___ Administer **Hand Usage Questionnaire** [if left hand use for more than 3 times, note and ask Sid later]

Basic Training (PowerPoint) and Rating Practice (E-Prime)

SCREEN 1

Now we would like to begin the training for the fMRI procedure. In the past year, over 100,000 news media stories were analyzed based on their most common themes. You will be presented with representative short statements which have been isolated from these common themes.

You will also be presented with representative faces which have been graphically manipulated to form composites of the faces of the people associated with the stories. This is so you can put a face with the behavior, to help it seem more real to you. We will be asking you to rate both the people and their behaviors in the fMRI. Do you have any questions?

SCREEN 2 and 3

Please take a minute to read through all the themes you will be rating. If you have any questions about the themes or what they mean, don't hesitate to ask me. Just click forward through the themes until you see a blank screen.

ALLOW PARTICIPANT TO READ THROUGH STORIES AT THEIR OWN PACE

If they have questions, refer to the following:

- Stories were collected from 40 demographic markets across the US over a 1 year period
- They were thematically analyzed based on common denominators such as actions and behaviors of individuals
- If you notice some are missing, again these are the most common themes, not every news story will be represented.
- Faces were made composite so they looked generic rather than like a specific person.
- If you are interested in any of these procedures, we will take your email and Allison will be happy to answer any further questions you might have.

SCREEN 4 - BLANK

As we've said before, you will be rating a person or their actions based on the information you are given regarding individuals in the news. First, we will show you the images you will see during the MRI scans. Then, we will also explain the ratings you need to complete during the scans. Finally, we will run you through a trial session in the MRI scanner.

SCREEN 5

First you will see a screen like this. There will be a face, an action, and the letter A or P next to the face.

Let's start with the face. This silhouette is a placeholder. You will actually see a composite photo of all the faces associated with this kind of story in the news, with hair and identifying facial features removed. But for now, let's just practice with the silhouettes.

Next you will see an action described below the face. These are the same actions we just showed you a minute ago. Make sure to pay attention to the key words which are underlined and in bold.

Finally, to the left side of the face, you will see the letters "A" or "P". This letter indicates which type of judgment you will make.

"A" stands for ACTION and "P" stands for PERSON

On this screen, you are being asked to make a judgment about this action.

SCREEN 6

Whereas on this screen, you are being asked to make a judgment about the person. You will see either an A or a P, not both. I just wanted to show you what each looked like, so you had an idea what to look for. Make sense?

SCREEN 7

And then finally, after seeing the face and reading the action, you will make a judgment based on how moral or immoral you believe this PERSON or ACTION is. You will have 6 seconds to rate each statement. You can't go back and change your answer once a selection has been made.

For example, here you would rate how moral or immoral the person is.

SCREEN 8

And on this one, you would rate how moral or immoral the action is.

In the scanner, you will have a glove on your right hand which will allow you to make the ratings. You will use your index finger to move the rating to the left, and you will use your middle finger to move your rating to the right.

*To lock in an answer, you will press your thumb.
It is important to know that when you lock in a response, you will automatically advance to the next screen.
Let's practice now. If you wanted to rate this guy a +2, what finger would you move?
What about a -2?
And how would you lock in your response?
Great!
Don't worry about getting it completely right now, we will have you practice this again in the magnet.*

SCREEN 9

*After each rating, you will see the word RELAX. This means the trial is over, and a new one is about to start. You will see a new face for the next trial, and you will again be asked to make a rating based on a description of them or an action they have performed.
After each Relax screen, you will see a little crossmark. This is just to keep you focused on the screen between trials.*

SCREEN 10

These next few slides show a summary of the order for each of the trials.

SCREEN 12

Again, this is the order of the trial that you will see while you are in the scanner. Do you have any questions?

SCREEN 13

Okay, let's go through a few quiz questions.

—[Have them verbally answer you. If they get any answers wrong, finish the PowerPoint and then go back and ask again the question(s) that they got wrong until they get them right.]

SCREEN 13

Q: What does the letter to the left of each face mean?

A: It indicates the type of judgment to be made. If you see a "P" then you must rate the person; if you see an "A" you must rate the action.

SCREEN 14

Q: When will you see the letters during the trial?

A: When I am seeing the person's face, when I am reading the statement, and when I am rating the statement.

SCREEN 15

Q: What does the cross-hair indicate?

A: That we are between trials

Q: When will you see the cross-hair?

A: At the end of a trial, after the 'Relax' slide

SCREEN 16

Q: What does the RELAX symbol mean?

A: That a trial has just ended and a new trial will begin. You will make a new judgment with a different face, and the rating will be based on their action or their description.

SCREEN 17

Q: *How do you make the ratings? For example, how would you make this guy a +3?*

A: You press your index finger to move to the left 3 times.

What about a -3?

A: You press your middle finger to move to the right 3 times.

Q: *How do you lock in a rating?*

A: You press your thumb

SCREEN 18

Q: *How long do you have to make each rating?*

A: 6 seconds, the rating screen will disappear after 6 seconds and you will not be able to return.

Q: *What will happen after you lock on a rating?*

The screen will change to the Relax screen.

SCREEN 19

Once they've successfully completed the questions:

"Okay. Thank you! Now we are going to move to the magnet and prepare you for scanning. We will practice the ratings again in the scanner. Please come with me now."

MRI Prep

__ Introduce Scarlett to participant and give her the Consent and MRI screening forms.

This is Scarlett. She will be the MRI technician in charge of running the scans. She will get you all set for the scanning session in a few minutes.

__ Secure subjects' belongings in one of the yellow lockers (get a key from Scarlett)

At this point we need to ask you to take off any metal jewelry, glasses, or other things and I will store them securely for you. Please empty your pockets and check everything. Bringing any metal into the scanner can be very dangerous, and we want to make sure you are safe. As a reminder, we need to ask you to take a pregnancy test. Anyone of child-bearing age must take one before entering an MRI machine to ensure safety.

__ Scarlett administers pregnancy test.

MRI Procedure

__ Log into computer in control room (username: sarinopoulos, password: sidis) and open E-prime

__ [Scarlett] performs a final MRI screening

__ [Scarlett] has subject lay on plank (bed) outside the bore.

__ [Scarlett] fits the response glove on right hand.

__ Ask how to make ratings (which fingers, time allotted for ratings etc...)

__ [Scarlett] slides subject in, inserts the LCD mirror and makes sure participant can see the LCD screen

Scan - Control Room

__ Ask Scarlett to call the participant and practice using the glove and make sure the box atop the computer is lighting up.

__ Run e-prime script with trial ratings

__ [Scarlett] after preliminary scans: *"the MRI procedure is about to begin"*

- ___ Log in to computer if not already done so (username: sarinopoulos, password: sidis) and open E-prime
- ___[TE] Runs scripts on the computer in coordination with Scarlett. Make sure she is ready for you before starting the run.
- ___To start, click on Running Man symbol in E-prime. Make sure to enter the correct Session Number and Subject Number.
 - MorMRI_2011-01-25_Run1.es2
 - o When the run is done, press SPACE BAR to save and transfer the output files
 - MorMRI_2011-01-25_Run2.es2
 - o Press SPACE BAR when done
 - MorMRI_2011-01-25_Run3.es2
 - o Press SPACE BAR when done
 - MorMRI_2011-01-25_Run4.es2
 - o Press SPACE BAR when done
 - MorMRI_2011-01-25_Run5.es2
 - o Press SPACE BAR when done
 - MorMRI_2011-01-25_Run6.es2
 - o Press SPACE BAR when done

End of Scan/Post Scan

- ___ Half-way through the last (anatomical) scan, remind Scarlett we need the brain images CD for the participant
- ___ [Scarlett] removes glove, belts etc and helps participant off the plank).
- ___Take participant to the yellow locker to retrieve belongings. (get a key from Scarlett)
- ___Take participant to E-105 (or back room depending on scheduling) for post-scan ratings and short responses

Post-Scan Ratings/Short Responses

- ___ Open Post-scan Ratings File: <http://research.adv.msu.edu/ss/wsb.dll/eden/Postscan.htm>
 - Fill in information on first screen, continue to second screen
 - If websurvey will not load, use hard copy of task and enter their answers online after they leave.
 - [Have participants perform ratings]
 - While participant is doing ratings, get money from cabinet and CD from Scarlett.
- ___Pay participant 40 dollars for their time, ask them to sign **Payment Receipt**.
- ___ Ask if they have any questions (there is a debriefing on Websurvey but if they have any further questions or comments).
- ___ Thank them for their time.

Figure 4

Faces Used in fMRI Experiment

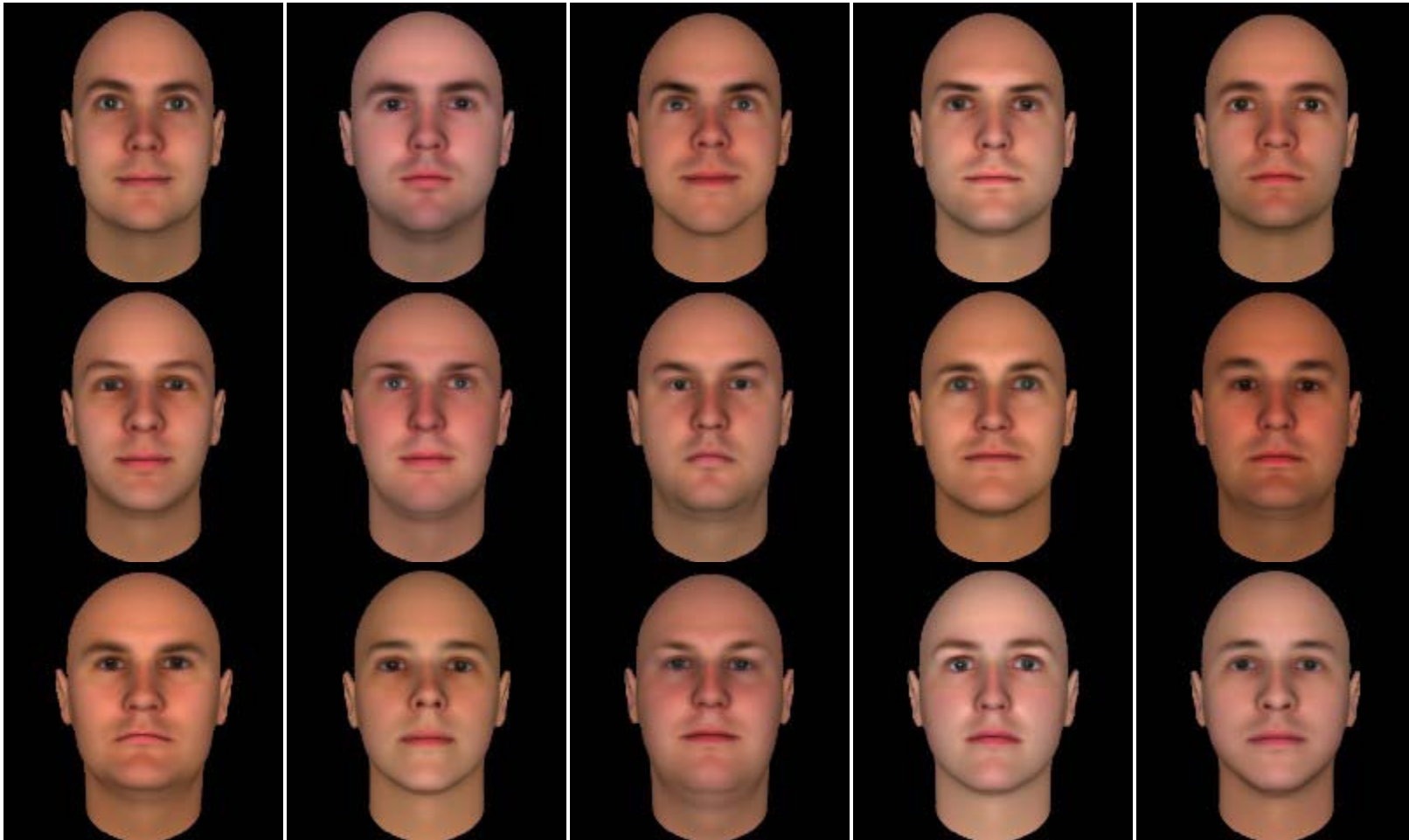


Figure 4 (con't)

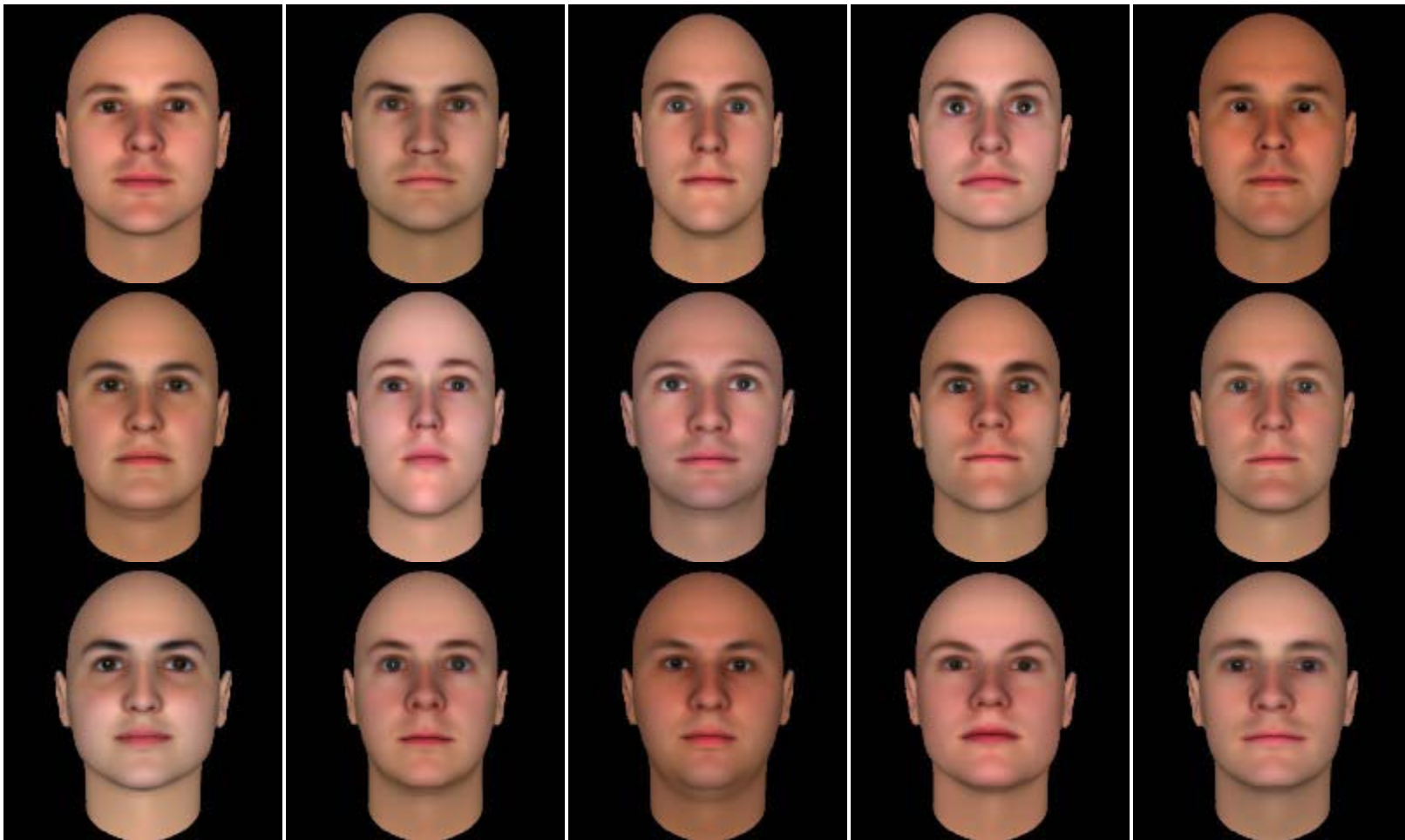


Figure 4 (con't)

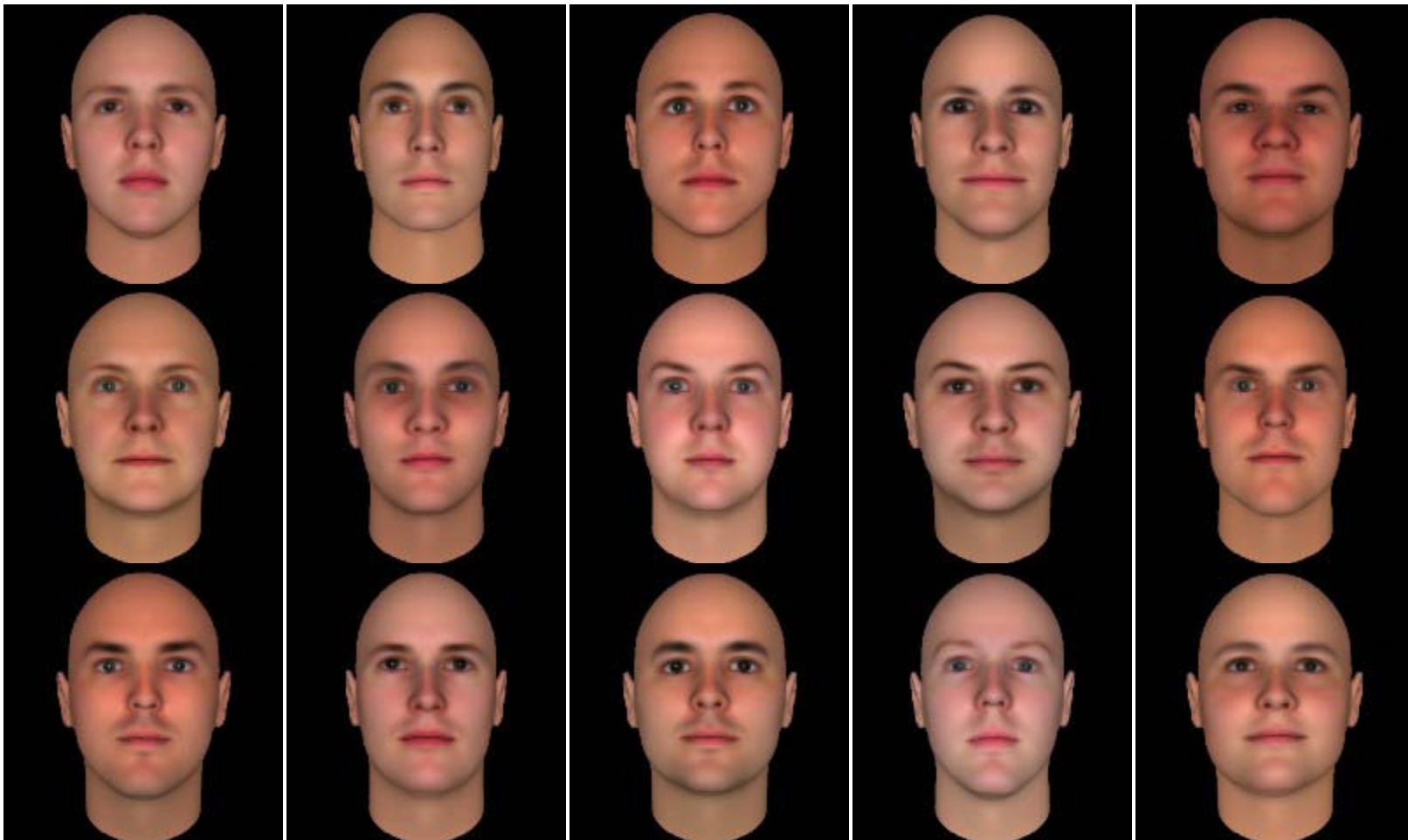


Figure 4 (con't)

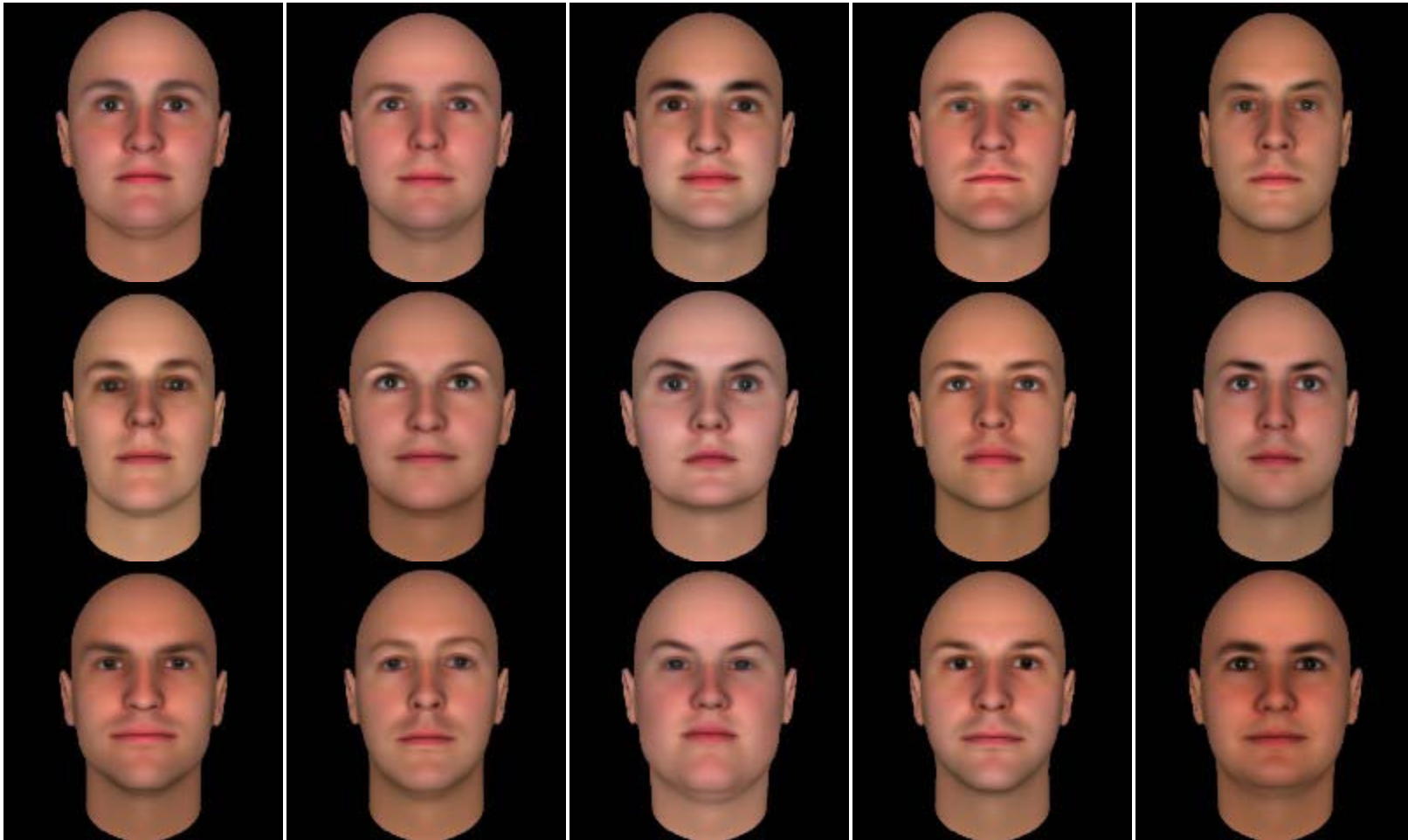


Figure 4 (con't)

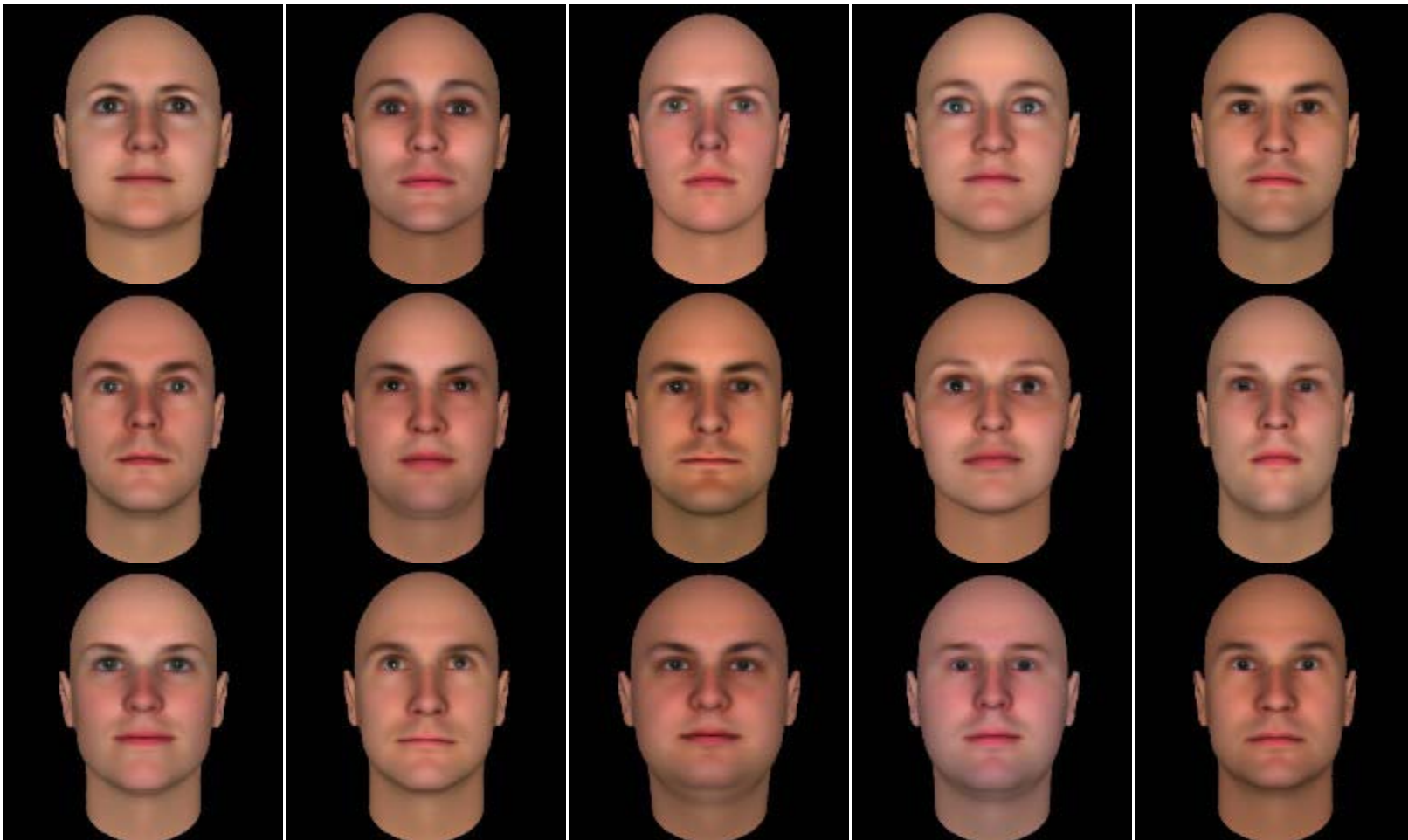


Figure 4 (con't)

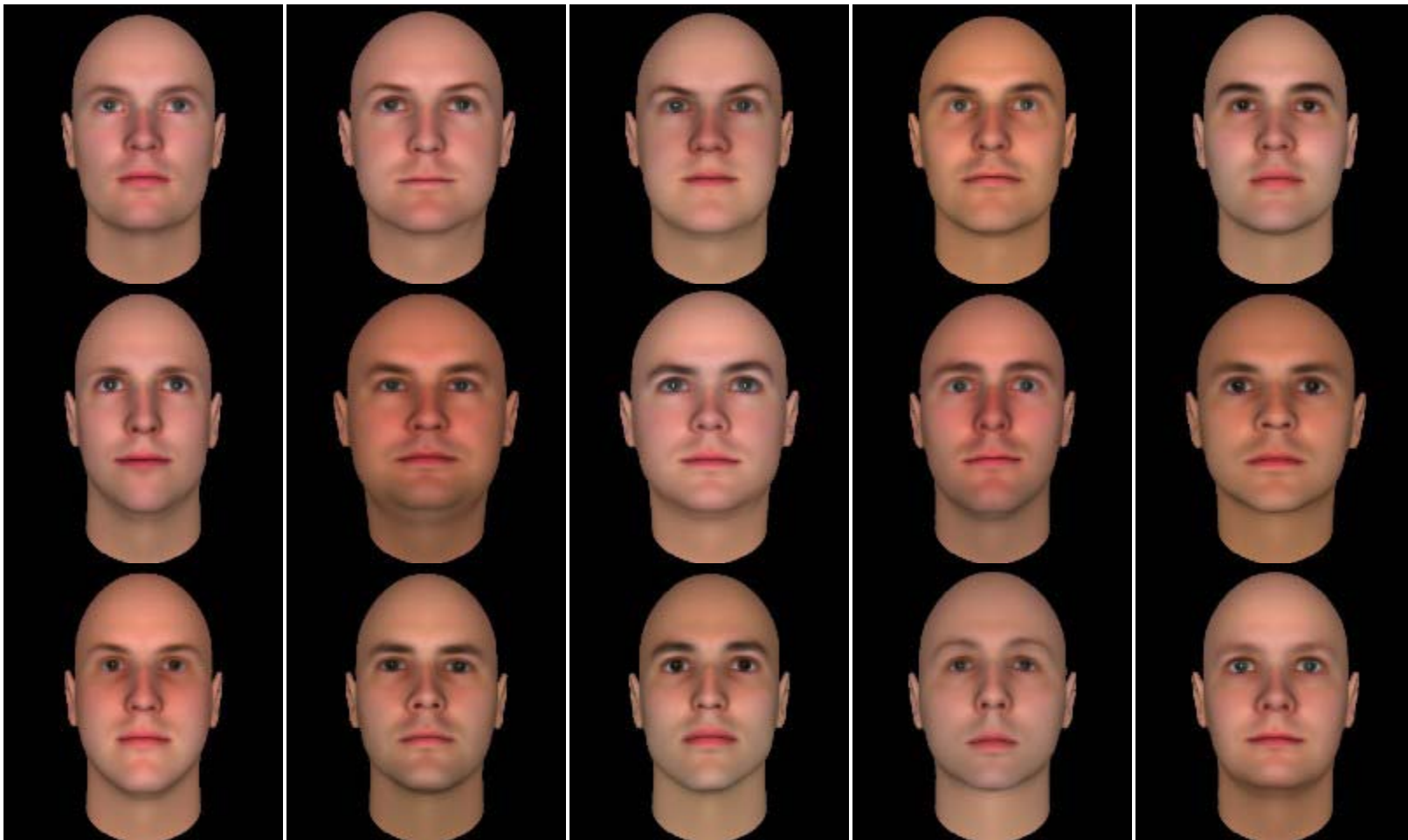


Figure 4 (con't)

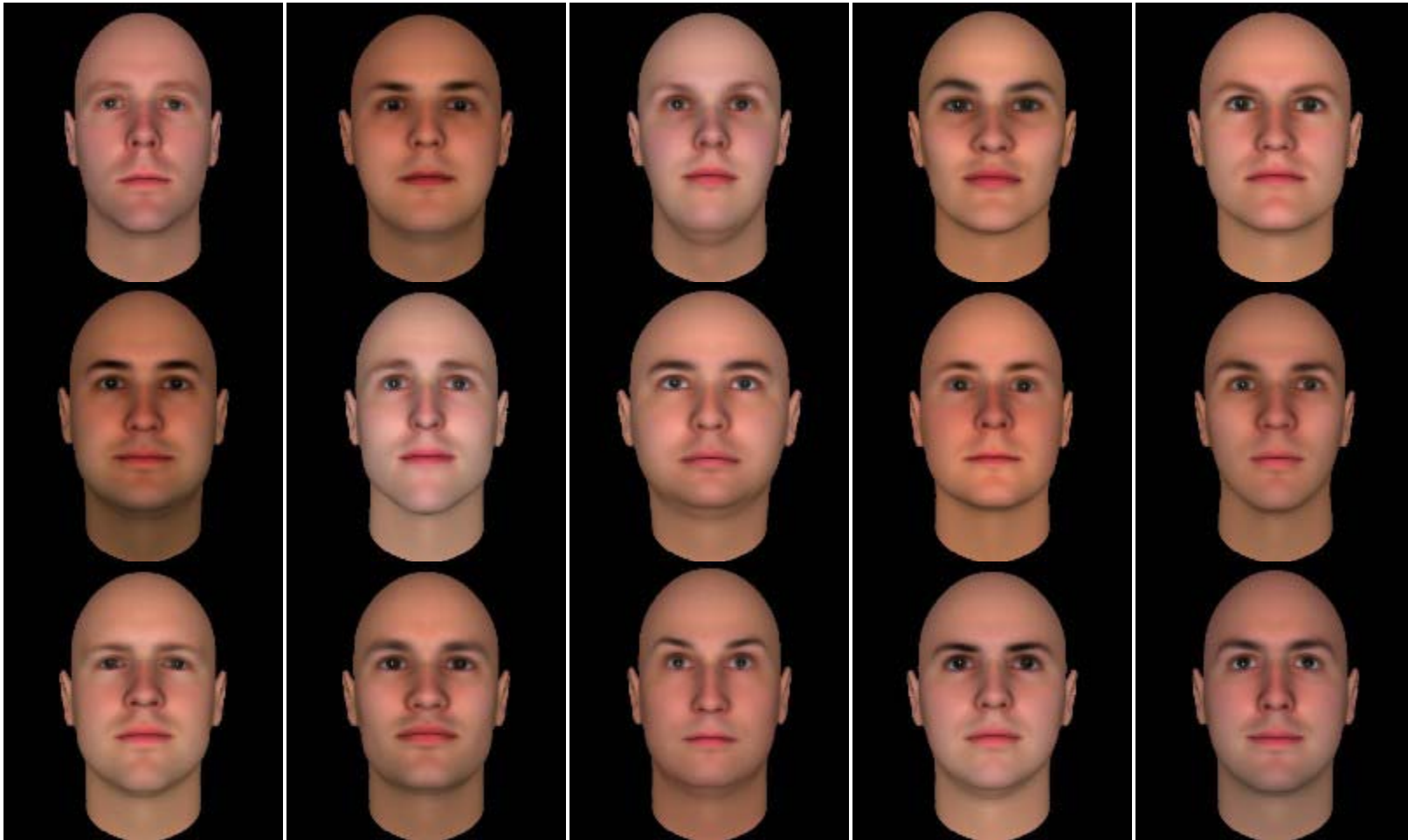
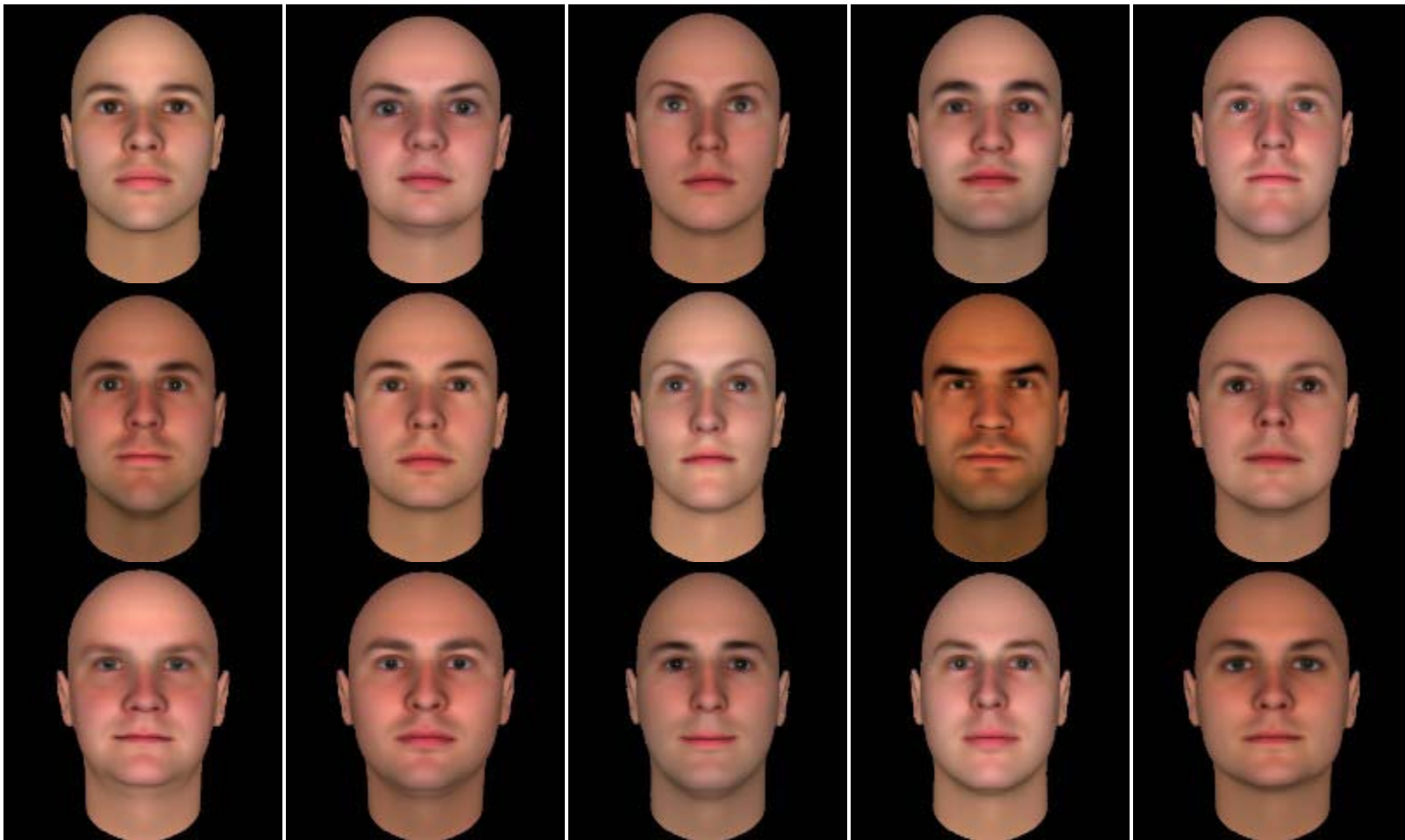


Figure 4 (con't)



Figure 4 (con't)



Appendix F

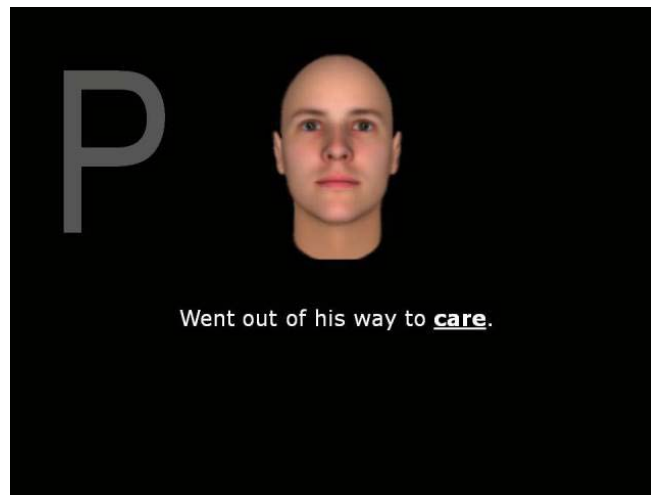
Post-scan Survey

As you know, the statements that you saw in the scan were based on the most common themes of 100,000 news stories. In this phase of the study, you will again be presented with statements that will reflect the ones you saw previously. This time, we would like you to think back and tell us how you felt while viewing those statements in the scanner. On this sheet, you will be asked to rank certain emotions on a scale of 1 to 10, and to also write down any thoughts that may have occurred with those emotions.

Please let the experimenter know if you have any questions at this time. When you are ready to begin, please turn the page and notify the experimenter.

Figure 5

Post-scan Survey



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)

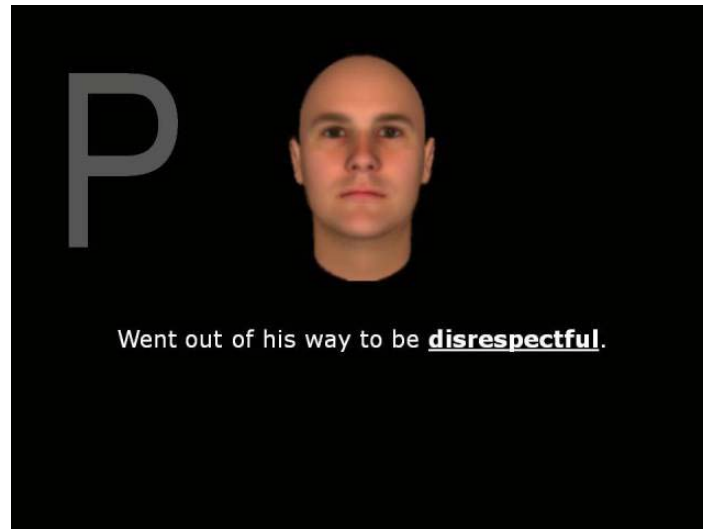


Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)

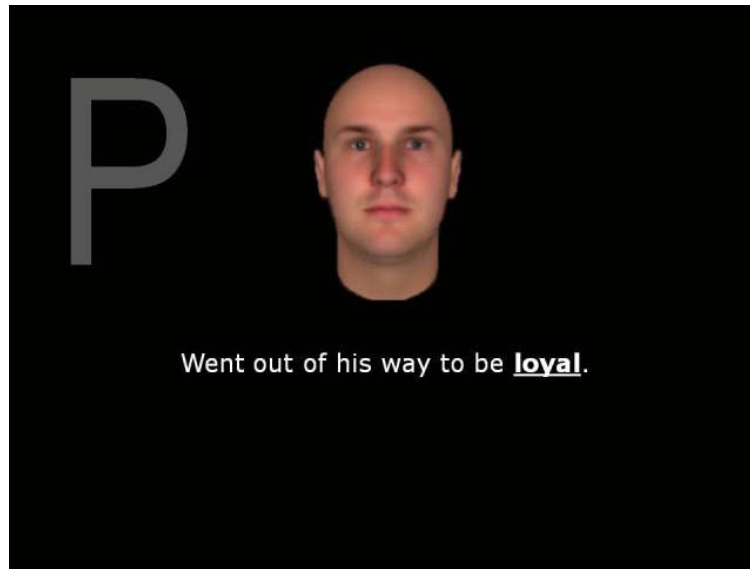


Please rank the following emotions that this statement may have evoked in you, with 1 indicating you Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)

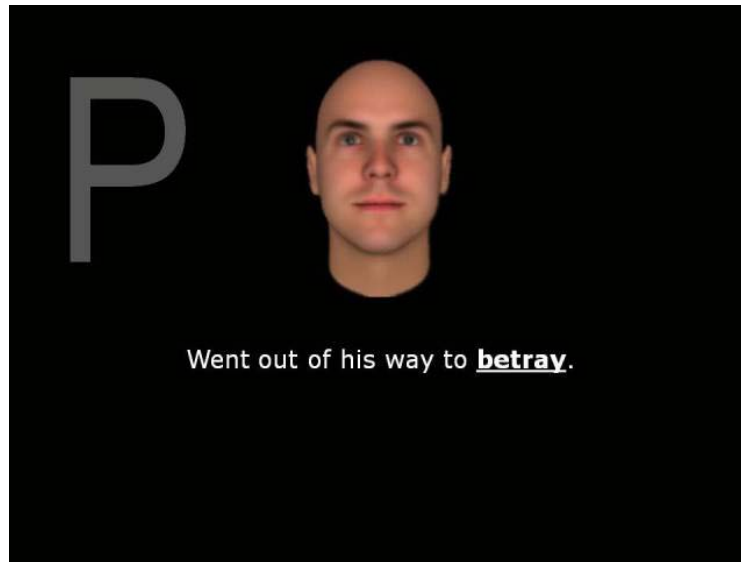


Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)

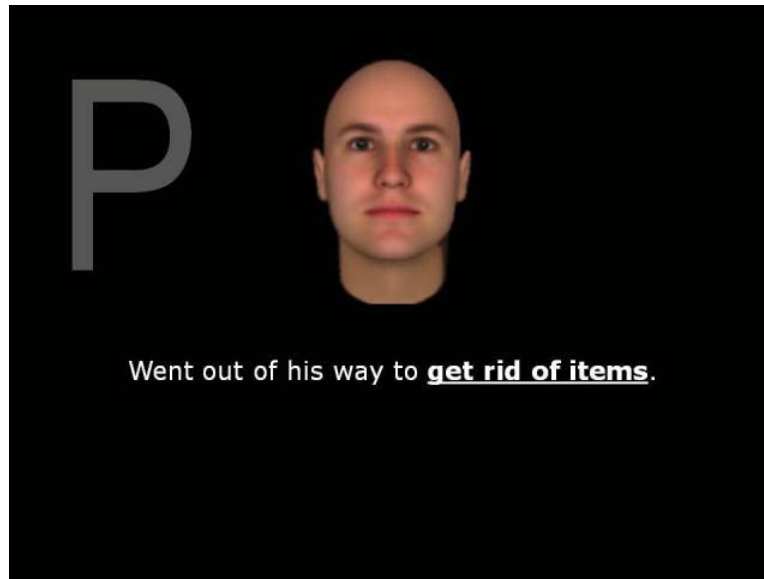


Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Figure 5 (con't)



Please rate the following emotions that you may have experienced when reading this statement **in the scanner**. 1 indicates you felt this emotion “the least possible” and 10 indicates you felt this emotion “the most possible.”

	Disgust
	Anger
	Contempt
	Indignation
	Admiration
	Elevation/Awe
	Gratitude
	Boredom
	Sadness
	Joy

When you read this statement in the scanner, please describe any **thoughts or feelings** that may have accompanied your emotions. Were there any news stories or specific examples that this statement evoked?

Appendix F

In the past five years, over 600 news media stories were analyzed based on their most common themes. Below you will be presented with representative short statements which have been isolated from these common themes, as well as a few extra statements we have added for comparison purposes. We would like you to rate the behavior described in the statement on a few different dimensions relevant to understanding media enjoyment.

First, you will be presented with all the statements. Please read each statement carefully. Imagine that each statement describes a fictional character performing these actions or a story based on the behavior described.

Next, you will be asked to rate the *behavior* described in the sentence along the following dimensions. We have provided definitions for the dimensions to help you rate the sentences.

When you are ready to begin, please turn the page.

Below are statements representing fourteen of the most common media narratives. Please read each statement carefully, thinking of narrative situations that might be similar to the behavior described in the statement.

The person went out of his way to care.

The person went out of his way to be fair.

The person went out of his way to be respectful.

The person went out of his way to fail.

The person went out of his way to be loyal.

The person went out of his way to harm.

The person went out of his way to be unfair.

The person went out of his way to succeed.

The person went out of his way to betray.

The person went out of his way to be disrespectful.

The person went out of his way to collect items.

The person went out of his way to get rid of items.

We would like you to rate how arousing each statement is. This indicates *how exciting or calm* you found the behavior.

Please use the following scale:

1	2	3	4	5	6	7	8	9
Extremely				Neither				Extremely
Boring				Exciting nor Boring				Exciting

The person went out of his way to care. _____

The person went out of his way to be fair. _____

The person went out of his way to be respectful. _____

The person went out of his way to fail. _____

The person went out of his way to be loyal. _____

The person went out of his way to harm. _____

The person went out of his way to be unfair. _____

The person went out of his way to succeed. _____

The person went out of his way to betray. _____

The person went out of his way to be disrespectful. _____

The person went out of his way to fail. _____

The person went out of his way to collect items. _____

The person went out of his way to get rid of items. _____

Next we would like to know how intentional you found the behavior to be. Intentionality means how much do you think the character *meant to do it* versus *how much do you think the behavior was an accident*?

Please use the following scale:

1	2	3	4	5	6	7	8	9	
Extremely			Neither						Extremely
Intentional			Intentional nor						Accidental

The person went out of his way to care. _____

The person went out of his way to be fair. _____

The person went out of his way to be respectful. _____

The person went out of his way to fail. _____

The person went out of his way to be loyal. _____

The person went out of his way to harm. _____

The person went out of his way to be unfair. _____

The person went out of his way to succeed. _____

The person went out of his way to betray. _____

The person went out of his way to be disrespectful. _____

The person went out of his way to fail. _____

The person went out of his way to collect items. _____

The person went out of his way to get rid of items. _____

Next, we would like you to rate each statement along several dimensions of emotion. You will be presented with each statement and then asked to indicate how much each statement evokes the particular emotion described. In order to help you, we have included definitions for each of the emotions below:

Disgust - a strong distaste; nausea; loathing

Anger - a strong feeling of displeasure aroused by a wrong

Contempt - the feeling with which a person regards anything considered mean, vile, or worthless

Indignation- strong displeasure at something considered unjust, offensive, insulting; righteous anger

Admiration - a feeling of wonder, pleasure, or approval

Elevation/Awe - a warm/open/pleasant feeling in the chest, even “choked up” with increased intensity, and a tendency to want to emulate and improve the self

Gratitude - the feeling of being grateful or thankful

Figure 6

Pilot survey emotion ratings

Please rate each statement to the extent it evokes the listed emotion. Use the following scale. Place an X or a check mark in the appropriate box.

1. The person went out of his way to care.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

2. The person went out of his way to be fair.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

Figure 6 (con't)

3. The person went out of his way to be respectful.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

4. The person went out of his way to fail.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

Figure 6 (con't)

5. The person went out of his way to be loyal.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

7. The person went out of his way to harm.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

Figure 6 (con't)

8. The person went out of his way to be unfair.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

9. The person went out of his way to succeed.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

Figure 6 (con't)

10. The person went out of his way to betray.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

11. The person went out of his way to be disrespectful.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

Figure 6 (con't)

12. The person went out of his way to fail.

	1 Least Amount	2	3	4	5	6	7	8	9 Most Amount
Disgust									
Anger									
Contempt									
Indignation									
Admiration									
Elevation/Awe									
Gratitude									

BIBLIOGRAPHY

BIBLIOGRAPHY

- Abraham, A., von Cramon, D.Y., & Schubotz, R. I. (2008). Meeting George Bush versus meeting Cinderella: The neural response when telling apart what is real from what is fictional in the context of our reality. *Journal of Cognitive Neuroscience*, 20, 965-976.
- Ames, D. L., Fiske, S. T., & Todorov, A. T. (2011). Impression formation: A focus on others' intents. In J. Decety & J. Cacioppo, (Eds.), *The Handbook of Social Neuroscience* (pp. 419-434). Oxford University Press.
- Amodio, D. M., Jost, J. T., Master, S. L., & Yee, C. M. (2007). Neurocognitive correlates of liberalism and conservatism. *Nature Neuroscience*, 10, 1246-1247.
- Bandura, A. (2001). Social cognitive theory of mass communication. *Media Psychology*, 3, 265-298
- Bavel, J. J., Packer, D. J., & Cunningham, W. A. (2008). The neural substrates of in-group bias: A functional magnetic resonance imaging investigation. *Psychological Science*, 19, 1131-1139.
- Berger, J., & Milkman, K. L. (2010). Social transmission and viral culture. Department of Marketing. University of Pennsylvania. Retrieved from http://opim.wharton.upenn.edu/~kmilkman/Virality_Feb_2010.pdf
- Berthoz, S., Armony, J. L., Blair, R. J. & Dolan, R. J. An fMRI study of intentional and unintentional (embarrassing) violations of social norms. *Brain*, 125, 1696–1708.
- Bowman, N. D., Dogruel, L, & Joeckel, S. (2011, May). Binding Americans and separating Germans: The influence of moral salience and nationality on media choices. Paper presented at the Annual Meeting of the International Communication Association, Boston, MA.
- Buckholz, C., Asplund, P., Dux, D., Zald, J., Gore, O., Jones, R., & Marios, R. (2008). The Neural correlates of third-party punishment. *Neuron*, 60, 930-940.
- Cavanna, A. E., & Trimble, M. R. (2006). The precuneus: a review of its functional anatomy and behavioral correlates. *Brain*, 129, 564–583.
- Chaminade, T., & Decety J. (2002). Leader or follower? Involvement of the inferior parietal lobule in agency. *NeuroReport*, 15, 1975-1978
- Chapman, H. A., Kim, D. A., Susskind, J.M., & Anderson, A.K. (2009). In bad taste: The evidence for the oral origins of moral disgust. *Science*, 5918,1222-1226.
- Cikara, M., & Fiske, S. T. (2011). Bounded empathy: Neural responses to outgroup targets' misfortunes. *Journal of Cognitive Neuroscience*, 1-13.

- Decety J., & Jackson P. L. (2004). The functional architecture of human empathy. *Behavioral Cognitive Neuroscience Review*, 3, 71–100.
- Delgado, M. R., Frank, R. H., & Phelps, E. A. (2005) Perceptions of moral character modulate the neural systems of reward during the trust game. *Nature Neuroscience*, 8, 1611-1618.
- de Oliveira-Souza, R., & Moll, J. (2000). The moral brain: functional MRI correlates of moral judgment in normal adults. *Neurology*, 54, 252.
- de Quervain, D. J-F., Fischbacher, U., Treyer, V., Schellhammer, M., Schnyder, U., Buck, A., & Fehr, E. (2004). The neural basis of altruistic punishment. *Science*, 305, 1254-1258.
- Eden, A. & Tamborini, R. (2010, June). *Moral modules as indicators of morality subcultures*. Paper presented at the Annual Meeting of the International Communication Association, Singapore.
- Eden, A., Oliver, M. B., Tamborini, R., Woolley, J., & Limperos, A. (2009, November). *Morality subcultures and perceptions of heroes and villains*. Paper presented at the Annual Meeting of the National Communication Association, Chicago, IL.
- Ehrlich, H., Weller, J., & Eden, A. (2009). The design of local TV news: If it's white, it's right. In H. Ehrlich (Ed.) *Hate Crimes and Ethnoviolence: The History, Current Affairs, and Future of Discrimination in America* (pp. 93-109). Westview Press: Boulder, CO.
- Eslinger, P. J., Moll, J., & de Oliveira-Souza, R. (2002). Emotional and cognitive processing in empathy and social behaviour. *Behavioral and Brain Science*, 25, 34-35.
- Frith, C. D., & Frith, U. (2006). How we predict what other people are going to do. *Brain Research*, 1079, 36-46.
- Gilbert D. T., & Malone P. S. (1995). The correspondence bias. *Psychological Bulletin*, 117, 21–38.
- Glenn, A. L., Raine, A., Schug, R. A., Young, L., & Hauser, M. (2007). Increased DLPFC activity during moral decision-making in psychopathy. *Molecular Psychiatry*, 14, 909-911.
- Greene, J. D., Sommerville, R. B., Nystrom, L. E., Darley, J.M., & Cohen, J. D. (2001). An fMRI investigation of emotional engagement in moral judgment. *Science*, 293, 2105–8.
- Greene, J. D., Nystrom, L. E., Engell, A. D., Darley, J. M., & Cohen, J. D. (2004). The neural bases of cognitive conflict and control in moral judgment. *Neuron*, 44, 389–400.
- Graham, J., Haidt, J., & Nosek, B. (2009). Liberals and conservatives rely on different sets of moral foundations. *Journal of Personality and Social Psychology*, 96, 1029-1046.

- Haidt, J. (2001). The emotional dog and its rational tail: A social intuitionist approach to moral judgment. *Psychological Review*, 108, 814 – 834.
- Haidt, J. (2003). The moral emotions. In R. J. Davidson, K. R. Scherer, & H. H. Goldsmith (Eds.), *Handbook of affective sciences* (pp. 852-870). Oxford: Oxford University Press.
- Haidt, J., & Graham, J. (2009). Planet of the Durkheimians, where community, authority, and sacredness are foundations of morality. In J. Jost, A. C. Kay & H. Thorisdottir (Eds.), *Social and psychological bases of ideology and system justification* (pp. 371 – 401). New York: Oxford.
- Haidt, J., Graham, J. & Hersh, M. A. (2006). *The Moral Foundations Questionnaire*. Retrieved January 31, 2008, from <http://faculty.virginia.edu/haidtlab/mf.html>
- Haidt, J., & Joseph, C. (2008). The moral mind: How 5 sets of innate moral intuitions guide the development of many culture-specific virtues, and perhaps even modules. In P. Carruthers, S. Laurence, and S. Stich (Eds.) *The Innate Mind* (pp. 367-393). Oxford: Oxford University Press.
- Haidt, J., & Kesebir, S. (2010) Morality. In S. Fiske, D., Gilbert, & G. Lindzey (Eds). *Handbook of Social Psychology*, 5th edition. (pp. 797-832.) Oxford: Oxford University Press.
- Haidt, J., Rozin, P., McCauley, C., & Imada, S. (1997). Body, psyche, and culture: The relationship of disgust to morality. *Psychology and Developing Societies*, 9, 107-131.
- Harenski, C. L., & Hamann, S. (2006). Neural correlates of regulating negative emotions related to moral violations. *Neuroimage*, 30, 313–324.
- Harbaugh, W. T., Mayr, U. & Burghart, D. R. (2007) Neural responses to taxation and voluntary giving reveal motives for charitable donations. *Science*, 5831, 1622-1625.
- Haxby, J. V., Hoffman, E. A., & Gobbini, M. I. (2000). The distributed human neural system for face perception. *Trends in Cognitive Sciences*, 4, 223-232.
- Haxby J. V., & Gobbini, M. I. (2007). The perception of emotion and social cues in faces. *Neuropsychologia*, 45, 1.
- Heider, F. (1958). *The psychology of interpersonal relations*. Hillsdale, NJ: Lawrence Erlbaum Associates.
- Heekeren, H. R., Wartenburger, I., Schmidt, H., Prehn, K. Schwintowski, H. P., & Villringer, A. (2005). Influence of bodily harm on neural correlates of semantic and moral decision-making. *Neuroimage*, 24, 887–897.
- Heekeren, H. R., Wartenburger, I, Schmidt, H., Schwintowski, H. P., & Villringer, A. (2003). An fMRI study of simple ethical decision-making. *Neuroreport*, 14, 1215–1219.

- Himmelweit, H. T., Oppenheim, A. N., & Vince, P. (1958). *Television and the child*. New York: Oxford University Press.
- Hoffner, C. & Cantor, J. (1991). Perceiving and responding to mass media characters. In J. Bryant & D. Zillmann (Eds.), *Responding to the screen: Reception and reaction processes* (pp. 63-91). New Jersey: Lawrence Erlbaum.
- Hofmann, W., & Baumert, A. (2010). Immediate affect as a basis for moral judgment: An adaptation of the affect misattribution procedure. *Cognition and Emotion*, 24, 522-535.
- Horberg, E. J., Oveis, C., Keltner, D., & Cohen, A. B. (2009). Disgust and the moralization of purity. *Journal of Personality and Social Psychology*, 97, 963-976.
- Izuma, K., Saito, D. N. & Sadato, N. (2009). The roles of the media prefrontal cortex and striatum in reputation processing. *Social Neuroscience*, 24, 1-15.
- Johnson, M. K., Raye, C. L., Mitchell, K. J., Touryan, S.R., Greene, E.J., & Nolen-Hoeksema, S. (2006). Dissociating medial frontal and posterior cingulate activity during self-reflection. *Social Cognitive Affective Neuroscience*, 1, 56-64.
- Jost, J. T., Federico, C. M. & Napier, J. L. (2009). Political ideology: Its structure, functions, and elective affinities. *Annual Review of Psychology*, 60, 307-333.
- Kanwisher, N., McDermott, J., & Chun, M. (1997). The fusiform face area: A module in human extrastriate cortex specialized for the perception of faces. *Journal of Neuroscience*, 17, 4302-4311.
- Kawashima, R., Roland, P. E., & O'Sullivan, B. T. (1995). Functional anatomy of reaching and visuomotor learning: a positron emission tomography study. *Cerebral Cortex*, 5, 111–122.
- Kerbel, M. R. (2000). *If it bleeds, it leads: An anatomy of television news*. Boulder, CO: Westview Press.
- King-Casas, B., Tomlin, D., Anen, C., Camerer, C. F., Quartz, S.R., & Montague, P. R. (2005) Getting to know you: reputation and trust in a two-person economic exchange. *Science*, 308, 78-83.
- Kjaer, T.W., Nowak, M., & Lou, H. C. (2002). Reflective self-awareness and conscious states: PET evidence for a common midline parietofrontal core. *Neuroimage* 17, 1080–1086.
- Konijn, E. A. & Hoorn, J. F. (2009). Some like it bad. Testing a model on perceiving and experiencing fictional characters. In: Gunter, B. & Machin, D. (Eds.), *Media Audiences* (pp. 107-143). SAGE Publications.

- Lewis, R., Grizzard, M., Bowman, N. D., Eden, A., & Tamborini, R. (2011, April). Intuitive Morality and Reactions to News Events: Responding to News of the Lockerbie Bomber's Release. Paper presented at the Media and Morality Symposium of Broadcast Education Association, Las Vegas.
- Liss, M. B., Reinhardt, L.C., & Fredrikson, S. (1983). TV heroes: The impact of rhetoric and deeds. *Journal of Applied Developmental Psychology*, 4, 175-187.
- Luo, Q., Nakic, M., Wheatley, T., Richell, R., Martin, A., Blair, R. J. R. (2006). The neural basis of implicit moral attitude - an IAT study using event related fMRI. *NeuroImage* 30, 1449-1457.
- Marsh, A. A., Blair, K. S., Jones, M. M., Soliman, N, & Blair, R. J. R. (2009). Dominance and submission: the ventrolateral prefrontal cortex and responses to status cues. *Journal of Cognitive Neuroscience*, 21, 713-724.
- McCabe K., Houser, D., Ryan, L., Smith, V., & Trouard, T. (2001). A functional imaging study of cooperation in two-person reciprocal exchange. *Proceedings of the National Academy of Science*, 98, 11832-11835.
- Mitchell J. P., Macrae C. N., & Banaji M. R. (2006). Dissociable medial prefrontal contributions to judgments of similar and dissimilar others. *Neuron* 50, 655–663.
- Mobbs, D., Yu, R., Meyer, M., Passamonti, L., Seymour, B., Calder, A.J., et al. (2009) A key role for similarity in vicarious reward. *Science*, 324, 900-904.
- Moll, J., de Oliveira-Souza, R., Moll, F. T., Ignacio, F. A., Bramati, I. E., Caparelli-Daguer, E. M., & Eslinger, P. J. (2005a). The moral affiliations of disgust – A functional MRI study. *Cognitive and Behavioral Neurology*, 18, 68–78.
- Moll, J., de Oliveira-Souza, R., Bramati, I. E., & Grafman, J. (2002a). Functional networks in emotional moral and nonmoral social judgments. *Neuroimage*, 16, 696–703.
- Moll, J., de Oliveira-Souza, R., Eslinger, P. J., Bramati, I. E., Mourao-Miranda, J., Andreiuolo, P. A., & Pessoa, L. (2002b). The neural correlates of moral sensitivity: A functional magnetic resonance imaging investigation of basic and moral emotions. *Journal of Neuroscience*, 22, 2730–2736.
- Moll, J., Krueger, F., Zahn, R., Pardini, M., de Oliveira-Souza, R., & Grafman, J. (2006). Human front-mesolimbic networks guide decisions about charitable donation. *Proceedings of the National Academy of Sciences*, 103, 15623-15628.
- Moll, J., Zahn, R., de Oliveira-Souza, R., Krueger, F., & Grafman, J. (2005b). The neural basis of human moral cognition. *Nature*, 6, 800-811.

- Parkinson, C., Sinnott-Armstrong, W., Koralus, P. E., Mendelovici, A., McGeer, V., & Wheatley, T. (2011). Is morality unified? Evidence that distinct neural systems underlie moral judgments of harm, dishonesty, and disgust. *Journal of Cognitive Neuroscience*, 31, 1530-8898.
- Prehn K., Wartenburger, I., Mériaux, K., Scheibe, C., Goodenough, O.R., Villringer, A., van der Meer, E., & Heekeren, H. R. (2008). Individual differences in moral judgment competence influence neural correlates of socio-normative judgments. *Social Cognitive Affective Neuroscience*, 3, 33-46.
- Platek, S. M., & Kemp, S. M. (2007) Is family special to the brain? An event-related fMRI study of familiar, familial, and self-face recognition. *Neuropsychologia*, 47, 849-58.
- Raine, A. & Yang, Y. (2006). Neural foundations to moral reasoning and antisocial behavior. *Social, Cognitive, and Affective Neuroscience*, 1, 203-213.
- Raney, A. A., & Bryant, J. B., (2002). Moral judgment as predictor of enjoyment . *Media Psychology*, 4(4), 305-322.
- Raney, A. A. (2004). Expanding disposition theory: Reconsidering character liking, moral evaluations, and enjoyment. *Communication Theory*, 14 (4), 348-369.
- Rilling, J. K., Sanfey, A. G., Aronson, J. A., Nystrom, L. E., & Cohen, J. D. (2004). The neural correlates of theory of mind within interpersonal interactions. *NeuroImage*, 22(4), 1694-1703.
- Robertson, D. C., Snarey, J., Ousley, O., Harenski, K., Bowman, F. D., Gilkey, R., & Kilts, C. (2007). The neural processing of moral sensitivity to issues of justice and care. *Neuropsychologia*, 45, 755-756.
- Sanfey, A. G., Loewenstein, G., McClure, S. M., & Cohen, J. D. (2006) Neuroeconomics: cross-currents in research on decision-making. *Trends Cognitive Science*. 10, 108-116
- Schiach Borg, J. S., Hynes, C., Van Horn, J., Grafton, S., & Sinnott-Armstrong, W. (2006). Consequences, action, and intention as factors in moral judgments: an fMRI investigation. *Journal of Cognitive Neuroscience*, 18, 803–17.
- Schaich Borg, J., Lieberman, D., Kiehl, K. A. (2008). Infection incest iniquity Investigating the neural correlates of disgust and morality. *Journal of Cognitive Neuroscience*, 209, 1529-1546.
- Schnall, S., Haidt, J., Clore, G. L., & Jordan, A. H. (2008). Disgust as embodied moral judgment. *Personality and Social Psychology Bulletin*, 34, 1096 – 1109.

- Shweder, R. A., Much, N. C., Mahapatra, M., & Park, L. (1997). The “big three ” of morality (autonomy, community, and divinity), and the “ big three ” explanations of suffering. In A. Brandt & P. Rozin (Eds.), *Morality and health* (pp. 119 – 169). New York: Routledge.
- Singer, T., Seymour, B., O'Doherty, J. P. Kaube, H., Dolan, R. J., & Frith, C. D. (2004). Empathy for pain involves the affective but not sensory components of pain. *Science*, 303, 1157–1162.
- Singer, T., & Steinbeis, N. (2009). Differential roles of fairness- and compassion-based motivations for cooperation, defection, and punishment. *Annals of the New York Academy of Sciences*, 1167, 41-50.
- Spitzer, M., Fishbacher, U., Herrnberger, B., Gron, G., & Fehr, E. (2007). The neural signature of social norm compliance. *Neuron*, 56, 185-196.
- Tabibnia, G., Satpute, A. B., & Lieberman, M. D. (2008). The sunny side of fairness: Preference for fairness activates reward circuitry (and disregarding unfairness activates self-control circuitry). *Psychological Science*, 19, 339-347.
- Takahashi, H., Kato, M., Matsuura, M., Koeda, M., Yahata, N., Suhara, T., & Okubo, Y. (2008). Neural correlates of human virtue judgment. *Cerebral Cortex*, 18, 1886-1891.
- Tamborini, R. (2011). Moral intuition and media entertainment. *Journal of Media Psychology*, 23, 39-45.
- Tamborini, R., Eden, A., Bowman, N., Grizzard, M., & Lachlan, K. (in press). *The influence of morality subcultures on the acceptance and appeal of violence*. *Journal of Communication*.
- Tamborini, R., Eden, A., Bowman, N. D., Grizzard, M., & Weber, R. (2009, November). *Predicting enjoyment from implicit morality*. Paper presented at the Annual meeting of the National Communication Association, Chicago, IL.
- Tamborini, R., Enriquez, M., Lewis, R. J., Grizzard, M. N., & Mastro, D. (2011). A content analysis of moral foundations presented in Spanish and English language soap operas. Paper presented at the Annual Meeting of the International Communication Association, Boston, MA.
- Tamborini, R., Grizzard, M., Eden, A., & Lewis, R. (November, 2011). Imperfect heroes and villains: Patterns of upholding and violating distinct moral domains and character appeal. Paper to be presented at the Annual Meeting of the National Communication Association, New Orleans, LA.
- Tamborini, R., Grizzard, M., Lewis, R., & Eden, A. (November, 2011). Priming intuitive morality: The effect of heroes and villains on immediate affective responses to moral

- stimuli. Paper to be presented at the Annual Meeting of the National Communication Association, New Orleans, LA.
- Tamborini, R., Weber, R., Eden, A. Bowman, N. D., & Grizzard, M. (2010). Repeated exposure to daytime soap opera and shifts in moral judgment toward social convention. *Journal of Broadcasting and Electronic Media*, 54, 621-640.
- Todorov, A., Baron, S. G., & Oosterhof, N. N. (2008). Evaluating face trustworthiness: A model based approach. *Social Cognitive and Affective Neuroscience*, 3(2), 119-127.
- Todorov, A., Gobbini, M. I., Evans, K. K., & Haxby, J. V. (2007). Spontaneous retrieval of affective person knowledge in face perception. *Neuropsychologia*, 45, 163-173.
- Todorov, A., Mandisodza, A. N., Goren, A., & Hall, C. C. (2005). Inferences of competence from faces predict election outcomes. *Science*, 308, 1623-1626.
- Uddin, L.Q., Molnar-Szakacs, I., Zaidel, E., & Iacoboni, M. (2006). rTMS to the right inferior parietal cortex disrupts self-other discrimination. *Social Cognitive and Affective Neuroscience*, 1, 65-71.
- van Leeuwen, F., & Park, J. H. (2009). Perception of social dangers, moral foundations, and political orientation. *Personality and Individual Differences*, 47, 169-173.
- Vogeley, K., May, M., Ritzl, A., Falkai, P., Zilles, K., & Fink, G. R. (2004). Neural correlates of first-person perspective as one constituent of human self-consciousness. *Journal of Cognitive Neuroscience*, 16, 817-827.
- Vuilleumier, P., & Pourtois, G. (2007). Distributed and interactive brain mechanisms during emotion face perception: Evidence from functional neuroimaging. *Neuropsychologia*, 45, 174-194.
- Winter, L., & Uleman, J. S. (1984). When are social judgments made? Evidence for the spontaneousness of trait inferences. *Journal of Personality and Social Psychology*, 47, 237 – 252.
- Wright, P., He, G., Shapira, N. A., Goodman, W. K., & Liu, Y. (2004). Disgust and the insula: fMRI responses to pictures of mutilation and contamination. *Neuroreport*, 15, 260-262.
- Wright, J. C., & Baril, G. (2011). The role of cognitive resources in determining our moral intuitions: Are we all liberals at heart? *Journal of Experimental Social Psychology*, 47, 1007-1012.
- Young, L., & Saxe, R. (2009). An fMRI investigation of spontaneous mental state inference for moral judgment. *Journal of Cognitive Neuroscience*, 21, 1396-1405

- Zahn, R., Moll, J., Paiva, M. M-F., Garrido, G. J., Krueger, F., Huey, E. D., & Grafman, J. (2009). The neural basis of human social values: Evidence from functional MRI. *Cerebral Cortex*, 19, 276-283.
- Zillmann, D. (2000). Basal morality in drama appreciation. In I. Bondebjerg (Ed.), *Moving images, culture and the mind* (pp. 53-63). Luton, England: University of Luton Press.
- Zillmann, D., & Bryant, J. (1975). Viewer's moral sanction of retribution in the appreciation of dramatic presentations. *Journal of Experimental Social Psychology*, 11, 572-582.
- Zillmann, D., & Cantor, J. R. (1976). A disposition theory of humor and mirth. In T. Chapman & H. Foot (Eds.), *Humor and laughter: Theory, research, and applications*. London: Wiley.
- Zillmann, D., & Knobloch, S. (2001). Emotional reactions to narratives about the fortunes of personae in the news theater. *Poetics*, 29, 189-206.
- Zink, C. F., Tong, Y., Chen, Q., Bassett, D.S., Stein, J. L., & Meyer-Lindenberg, A. (2008). Know your place: Neural processing of social hierarchy in humans. *Neuron*, 58, 164-165.