

## ABSTRACT

### Exploring Implicit Anti-Arab Prejudice and Sociopolitical Correlates Through a Minimal Paradigm

Stephen R. Martin, Ph.D.

Chairperson: JoAnn C. Tsang, Ph.D.

The present paper contains three studies exploring whether political icons with known ideologies (i.e., Bernie Sanders, Donald Trump) alter implicit prejudicial attitudes toward Arab and Chinese cultural groups. Using a modified affect misattribution procedure, the studies suggest anti-Arab attitudes are detectable in speeded ratings of Arabic symbols, relative to Chinese symbols (Studies 1–2) and control symbols (Study 3). Sociopolitically relevant variables (e.g., social dominance orientation, right wing authoritarianism, explicit anti-Arab prejudice, generalized prejudice) predict negative ratings of Arabic symbols. Moreover, these variables predict positive and negative responses following a Trump and Sanders image, respectively. Brief exposure to political icons did not alter the negative ratings of Arabic symbols.

Exploring Implicit Anti-Arab Prejudice and Sociopolitical Correlates Through a  
Minimal Paradigm

by

Stephen R. Martin, B.S., M.A.

A Dissertation

Approved by the Department of Psychology and Neuroscience

---

Charles A. Weaver III, Ph.D., Chairperson

Submitted to the Graduate Faculty of  
Baylor University in Partial Fulfillment of the  
Requirements for the Degree  
of  
Doctor of Philosophy

Approved by the Dissertation Committee

---

JoAnn C. Tsang, Ph.D., Chairperson

---

Wade C. Rowatt, Ph.D.

---

Thomas A. Fergus, Ph.D.

---

Keith P. Sanford, Ph.D.

---

Grant B. Morgan, Ph.D.

Accepted by the Graduate School  
December 2018

---

J. Larry Lyon, Ph.D., Dean

Copyright © 2018 by Stephen R. Martin

All rights reserved

## TABLE OF CONTENTS

|                                                              |     |
|--------------------------------------------------------------|-----|
| LIST OF FIGURES .....                                        | vi  |
| LIST OF TABLES .....                                         | vii |
| ACKNOWLEDGMENTS .....                                        | x   |
| DEDICATION .....                                             | xi  |
| CHAPTER ONE                                                  |     |
| Introduction .....                                           | 1   |
| 1.1 <i>Anti-Arab Prejudice</i> .....                         | 2   |
| 1.2 <i>Political Symbols and Ideologies</i> .....            | 3   |
| 1.3 <i>Moral Foundations and Prejudicial Attitudes</i> ..... | 6   |
| 1.4 <i>Authoritarianism and Domination</i> .....             | 10  |
| 1.5 <i>Proposed Studies</i> .....                            | 12  |
| CHAPTER TWO                                                  |     |
| Stimulus Pilot .....                                         | 15  |
| 2.1 <i>Methods</i> .....                                     | 15  |
| 2.2 <i>Results</i> .....                                     | 18  |
| 2.3 <i>Discussion</i> .....                                  | 23  |
| CHAPTER THREE                                                |     |
| Study 1 .....                                                | 25  |
| 3.1 <i>Methods</i> .....                                     | 26  |
| 3.2 <i>Results</i> .....                                     | 29  |
| 3.3 <i>Discussion</i> .....                                  | 37  |
| CHAPTER FOUR                                                 |     |
| Study 2 .....                                                | 39  |
| 4.1 <i>Methods</i> .....                                     | 39  |
| 4.2 <i>Results</i> .....                                     | 39  |
| 4.3 <i>Discussion</i> .....                                  | 43  |
| CHAPTER FIVE                                                 |     |
| Study 3 .....                                                | 45  |
| 5.1 <i>Methods</i> .....                                     | 45  |
| 5.2 <i>Results</i> .....                                     | 46  |
| 5.3 <i>Discussion</i> .....                                  | 52  |
| CHAPTER SIX                                                  |     |
| Integrative Data Analysis .....                              | 55  |
| 6.1 <i>Reanalysis of RWA Models</i> .....                    | 58  |

|                                                          |    |
|----------------------------------------------------------|----|
| CHAPTER SEVEN                                            |    |
| General Discussion .....                                 | 61 |
| 7.1 <i>Limitations</i> .....                             | 65 |
| APPENDIX A                                               |    |
| Political Candidate Attitude Measure .....               | 69 |
| APPENDIX B                                               |    |
| Revised Anti-Arab Prejudice Scale .....                  | 71 |
| APPENDIX C                                               |    |
| Prime Images and Mask .....                              | 72 |
| APPENDIX D                                               |    |
| Chinese Symbols .....                                    | 73 |
| APPENDIX E                                               |    |
| Arabic Symbols .....                                     | 74 |
| APPENDIX F                                               |    |
| Random Glyph Symbols .....                               | 75 |
| APPENDIX G                                               |    |
| Thermometer Items (Reversed Generalized Prejudice) ..... | 76 |
| BIBLIOGRAPHY .....                                       | 77 |

## LIST OF FIGURES

|             |                                                                                                                                                                                                |    |
|-------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Figure 2.1. | Probability of responses for each item, from the multinomial model. Items 1–36 are Arabic images, 37–72 are Chinese, 73–108 are random glyphs. Rows correspond to the response categories..... | 21 |
| Figure 3.1. | Experimental flow of studies 1–3. The diagram represents the order of stimuli seen per trial. The dashed line represents the stimulus only observed during study 3. ....                       | 27 |
| Figure 3.2. | Model template graph. Interaction coefficients are not plotted for brevity.                                                                                                                    | 32 |
| Figure 5.1. | Fixed effect predictions for each condition, across each latent moderator. The y-axis is the probability, ranging from 0 to 1, of responding “pleasant”. ....                                  | 49 |

## LIST OF TABLES

|            |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |    |
|------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Table 2.1. | Multinomial logistic mixed model expected values (and posterior standard deviations). Each row corresponds to a response category. Columns correspond to parameter estimates. The modeled value is the log-odds of choosing the response category, relative to choosing Arabic. “Not” is short for “Not a language”. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    | 20 |
| Table 2.2. | Estimated political leanings of Sanders and Trump on a scale ranging from 1 (Very liberal) to 7 (Very conservative) in four domains. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    | 23 |
| Table 3.1. | Parameter estimates for the six models predicting pleasant and unpleasant responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. For example, the fifth column (RWA), fourth row (Arabic) indicates that for each one unit increase in RWA, the effect of Arabic symbol decreases by -.22. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is <i>positive</i> (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal). . . . . | 34 |
| Table 3.2. | Parameter estimates for the six models predicting peaceful and aggressive responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is <i>positive</i> (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal). . . . .                                                                                                                                                           | 35 |

|            |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |    |
|------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Table 4.1. | Parameter estimates for the six models predicting pleasant and unpleasant responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is <i>positive</i> (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal)..... | 40 |
| Table 4.2. | Parameter estimates for the six models predicting peaceful and aggressive responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is <i>positive</i> (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal)..... | 41 |
| Table 5.1. | Parameter estimates for all models. Expected a-posteriori values provided, with [95% credible intervals] and (ppp). AA = anti-Arab prejudice, SDO = social dominance orientation, RWA-R = right-wing authoritarianism, TW = Trump warmth, SW = Sanders Warmth, Therm = Warmth toward other groups. First column represents the base model, with no moderators. Subsequent columns represent the moderating effect of the column construct on the corresponding row's effect. ....                                                                                                                                                                                                 | 48 |
| Table 5.2. | Correlation matrix of random effects (lower triangle), standard deviation of random effects (diagonal), and ppp (above diagonal) for study 3. ....                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                | 49 |
| Table 5.3. | Parameter estimates for the baseline model with trial order interacting with each predictor. Because the coefficients for Trial Order are < .002, these coefficients are multiplied by 108. Therefore, the Trial Order column indicates how much the respective row's effect would change by the final trial. ....                                                                                                                                                                                                                                                                                                                                                                | 51 |
| Table 5.4. | Correlations (lower diagonal) of latent factors with all other latent factors (with posterior probability that the correlation is positive) on the upper diagonal. Several 1s and 0s are present; this is because to .0001 precision, the value would be rounded to 0 or 1. ....                                                                                                                                                                                                                                                                                                                                                                                                  | 52 |

|            |                                                                                                                                                                                                                             |    |
|------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Table 5.5. | Coefficients for model including all predictors and moderators simultaneously.....                                                                                                                                          | 53 |
| Table 6.1. | Fixed effects meta-analysis on shared linear components across studies 1–3 [with 95% credible intervals] and (posterior probability of positivity, where notable) for the pleasant vs unpleasant response option data. .... | 57 |
| Table 6.2. | Fixed effects meta-analysis on shared linear components across studies 1–2 [with 95% credible intervals] and (ppp), for peaceful vs aggressive response option data. ....                                                   | 57 |
| Table 6.3. | Re-analysis of the moderating effect of RWA split by facets Authoritarian aggression/submission (AAS) and Conservatism (C). ....                                                                                            | 59 |
| Table 6.4. | Fixed effects meta-analytic estimates of the two RWA-R facets' effects. ...                                                                                                                                                 | 60 |
| Table 7.1. | Simple summary of hypothesis support across studies. ✓:Full support, /: Partial support (substantively, or probability > .9), ?: Uncertain due to noise, X: Unsupported. ....                                               | 61 |

## ACKNOWLEDGMENTS

I thank my wife Michelle Martin, and the rest of my wonderful family for their encouragement and support throughout my education. I am grateful to Dr. Tsang for her mentorship throughout the past several years, and to Dr. Rowatt for his wisdom and support. Thank you, my graduate school family and labmates, for your friendship and guidance.

## DEDICATION

To my beautiful wife, Michelle

## CHAPTER ONE

### Introduction

The 2016 election cycle was seemingly among the most divisive, controversial American elections in recent history. The political division across the nation, the polarizing candidates and positions, and the social issues that were so prominently displayed throughout the election cycle created a rich, albeit ephemeral, environment for observing social psychological phenomena. Among these phenomena are the social judgments and prejudices toward Arab individuals. Furthermore, whether current political icons affect individuals' judgments of Arab individuals is important not only to assess the impact of possible candidates on social attitudes, but more importantly how corresponding ideologies implicitly affect social attitudes. This proposal primarily aims to examine two phenomena. The first phenomenon is whether implicit anti-Arab prejudice is detectable in a minimal paradigm. The second phenomenon is whether exposure to political icons corresponding to popular political ideologies can influence judgments within the minimal paradigm. A secondary interest is in how sociopolitically relevant individual difference variables can influence the first and second phenomena.

First, previous literature on explicit and implicit anti-Arab prejudice is summarized and discussed. Constructs relevant to such prejudices are discussed (e.g., right-wing authoritarianism, social dominance orientation), along with broader theories of prejudicial attitudes. Second, the role of social and political symbols in social group dynamics and signaling is briefly summarized. In particular, the importance of symbols in intragroup and intergroup relations and ideological signaling will be examined. Third, the relevancy of the current political zeitgeist to anti-Arab prejudice and prejudice theory is summarized in order to justify the stimuli used in the three studies. Finally, three studies (and one pilot study) are presented with specific a priori predictions made for each study.

### *1.1 Anti-Arab Prejudice*

Prejudice toward Arab individuals is well-established. Western populations consistently perceive Arabs as threatening and aversive. Explicit prejudice toward Arabs was more widespread in the West than other prejudices even prior to high-profile terrorist attacks (e.g., 9/11; Strabac & Listhaug, 2008). Following the terrorist attacks on 9/11 and during the Iraq war, US citizens were explicitly less concerned about fair treatment of Arabs and perceived them as threatening (Coryn, Beale, & Myers, 2004; Hall, 2003; Oswald, 2005). After playing video games that contained non-Arab terrorists, anti-Arab attitudes increase, suggesting individuals associate Arab individuals with terrorism (Saleem & Anderson, 2013). Experimental evidence and field studies suggest that Arabs are much less likely to be considered for employment relative to other ethnicities across several Western cultures (Agerström & Rooth, 2009; Derous, Nguyen, & Ryan, 2009; Derous, Ryan, & Nguyen, 2012; Widner & Chicoine, 2011).

In addition to anti-Arab prejudice, a strong anti-Muslim prejudice is pervasive across Western nations. Arabs and Muslims are distinct groups; nevertheless they are often conflated targets of prejudicial attitudes. Non-Muslim Americans implicitly prefer Christians over Muslims (Rowatt, Franklin, & Cotton, 2005). Moreover, non-Arab individuals implicitly prefer all other people compared to Arabs and Muslims (Nosek et al., 2007). By merely priming Arab or Muslim categories, participants demonstrate a heightened threat response on a shooter bias task (Mange, Chun, Sharvit, & Belanger, 2012). These implicit and explicit prejudices are heightened if one is high in social dominance orientation (Cohrs & Asbrock, 2009; Henry, Sidanius, Levin, & Pratto, 2005; Johnson, Labouff, Rowatt, Patock-Peckham, & Carlisle, 2012; McFarland, 2010; Pratto et al., 1994; Rowatt et al., 2005; Uenal, 2016), right-wing authoritarianism (Altemeyer & Hunsberger, 1992; Echebarria-Echabe & Guede, 2007; Henry et al., 2005; Johnson et al., 2012; McFarland, 2010; Pedersen & Hartley, 2012; Rowatt et al., 2005; Uenal, 2016), and sensitivity to threat (Uenal, 2016). The media only

worsen this prejudice. Across political parties, the extent to which one consumes news media predicts greater anti-Muslim prejudice (Shaver, Sibley, Osborne, & Bulbulia, 2017).

Both implicit and explicit prejudice toward Arabs and Muslims are well established in the literature. The question remains of whether contemporary political icons and the normalization of their ideologies can affect such prejudices. In the context of the 2016 US election, a new political icon emerged and was elected into office in part on the premise of securing Americans from the perceived threat of the Arab world and the Islamic faith. To the degree that society believes a particular demographic is threatening or associate the demographic with negative emotions, individuals will believe that discrimination is justified (Pereira, Vala, & Costa-Lopes, 2010; Talaska, Fiske, & Chaiken, 2008). Moreover, given that Trump won the election, potentially normalizing expressions of anti-Arab prejudice, exposure to Trump may evoke such attitudes even moreso than prior to his election to office (Crandall, Miller, & White, 2018; Crandall & White, 2016). It is therefore plausible that the political icons and their respective ideologies may alter prejudicial attitudes.

## *1.2 Political Symbols and Ideologies*

Symbols are important, pervasive communicative tools in human social living. From the perspective of cultural evolution, symbols serve several intragroup and intergroup purposes. Symbols broadcast one's ideologies and group allegiances (Alcorta & Sosis, 2005). In doing so, symbols permit value sets to be communicated summarily across time and space. Given that a symbol represents a shared ideology or group, individuals may find solidarity and desired cooperative relationships by seeking individuals who broadcast their identity through a shared, valued symbol (Bulbulia, 2012). Symbols are therefore one of several communicative tools that promote cooperative actions (See Ahn, Janssen, & Ostrom, 2004, for a review).

Symbols are not merely communicative tools for groups. Symbols are associated with ideologies, which inform attitudes and behavior. Exposure to known cultural symbols

influence cooperativeness with strangers and ingroup members in a manner congruent to the ideology represented by the symbol (Wong & Hong, 2005). Exposure to ingroup symbols (e.g., national flags, confederate flags) enhance group cohesion and promote attitudes concordant with the represented ideology (Becker, Enders-Comberg, Wagner, Christ, & Butz, 2012; Butz, 2009; Butz, Plant, & Doerr, 2007; Callahan & Ledgerwood, 2016; Sibley, Hoverd, & Duckitt, 2011). Symbols of authority predict compliant behaviors (Bushman, 1984). Exposure to symbols representing prejudicial ideologies (e.g., the confederate flag) can evoke prejudicial evaluations of minorities (Ehrlinger et al., 2011). Symbols are not only descriptive communications, but also prescriptive given that the mere exposure to them can elicit attitudes and behaviors that are concordant with the ideology represented by the symbol.

The primary election cycle saw two popular anti-establishment candidates who quickly gained in popularity, and whose messages resonated with many citizens. However, the activism and values promoted by these two candidates were remarkably opposite of one another. Bernie Sanders gained popularity based on his message of inclusivism, social equality, and economic justice, whereas Donald Trump gained popularity based on his message of exclusivism, returning to previous American values, and economic scrupulousity. Both Trump and Sanders positioned themselves as political symbols for new political and social movements across the nation.

A primary concern of each political symbol and a common talking point amongst the various candidates was immigration reform. On the leftmost side of the political spectrum was the notion that literal and figurative borders should be blurred, and an emphasis was placed on amnesty and a more efficient immigration process (“A fair and humane immigration policy”, 2017); this represents an inclusive ideology that blurs the lines between outgroups and ingroups. Conversely, on the rightmost side was the notion that borders should be strengthened and made more concrete, quite literally with the proposal to build a large concrete wall on the southernmost border (“Immigration reform that will make America great

again”, 2018); this represents an exclusive ideology that perpetuates ingroup and outgroup differences. This debate continued into 2018 with Democrats and Republicans battling over new legislation to replace the deferred action for childhood arrivals act (DACA) following its 2017 repeal (“Trump ends DACA but gives Congress window to save it”, 2017).

Likewise, the political extremes have recently and repeatedly battled over immigration policies and visitation permits from predominantly Muslim nations out of fear of “radical Islamic terrorists” (“The latest: Trump says ‘radical Islamic terrorism’ must end”, 2017) and immigrants usurping jobs and wages from citizens (“Trump unveils legislation limiting legal immigration”, 2017). With regard to the former, left-wing symbols have explicitly avoided the term, so as to not conflate Islam with terrorism, whereas right-wing symbols have explicitly employed the term, perhaps mimicking each ideology’s perspective on outgroup and symbolic threats (“Obama: Why I won’t say ‘Islamic terrorism’”, 2016; Perry, Sibley, & Duckitt, 2013; Uenal, 2016). The notion that immigrants take jobs from US citizens is comparable to the perspective of those high in SDO (Perry et al., 2013). While the predominantly Republican administration led by Trump attempted to implement “travel bans” of individuals from Muslim nations, federal courts, Democratic legislators, and left-wing activists protested and attempted to overturn such a ban (“Trump travel ban: Timeline of a legal journey”, 2018).

Both ends of the political spectrum have grown increasingly disparate from one another over decades (Baldassarri & Gelman, 2008; Layman, Carsey, & Horowitz, 2006; Poole & Rosenthal, 1984), in part due to changing media consumption habits (Iyengar & Hahn, 2009; Lawrence, Sides, & Farrell, 2010; Levedusky, 2013; Prior, 2013) and media selection biases (Bolin & Hamilton, 2018; Crawford, Jussim, Cain, & Cohen, 2013). Within US politico-ideological groups, individuals have become more homogenous over time (Baldassarri & Gelman, 2008). Politico-ideological leaders, such as Trump and Sanders, have been employed as symbols representing the whole of their respective ideologies. From a social psychological perspective, a key question is whether exposure to such polarized

political symbols could implicitly influence individuals' behaviors and attitudes to align with the espoused ideologies. Recent evidence suggests that the targeting of Muslims and immigrants by the Trump campaign indeed has normalized anti-Muslim and anti-immigrant attitudes (Crandall et al., 2018).

Sanders and Trump are both political icons representing differing ideologies. Though they are not linguistic or geometric symbols, they are nevertheless icons representing two drastically different social and political ideologies. The two ideologies espoused by Sanders and Trump differ markedly in attitudes toward immigrants, and in particular toward Arabs and Muslims. With the apparent influence of symbols on attitudes and behaviors, it is possible that exposure to Sanders or Trump may similarly affect attitudes according to the ideologies espoused by each. Consequently, exposure to these two political icons may implicitly influence evaluations of Arab individuals.

Different ideologies may evoke differing moral concerns, schemata, and social values that affect evaluations of Arabs. For example, the ideologies of Sanders and Trump may differ in which moral domains are valued, and may evoke different information used in moral judgments. Likewise, the ideologies of Sanders and Trump may differ in the value of maintenance and creation of social hierarchies. These differing topologies of moral and social concerns may cause differing attitudes toward Arabs and Muslims. These examples are discussed in greater detail below.

### *1.3 Moral Foundations and Prejudicial Attitudes*

Moral foundations theory (MFT; J. Graham et al., 2011; Haidt, 2007) argues that there exist five primary domains of moral concern across cultures. Just as a “big five” personality structure is presumed to exist, Haidt argues for a “big five” structure of moral concerns. Haidt argues that the five moral concerns are evolutionarily derived and innate to humans (J. Graham et al., 2011). Factor analytic methods suggest that persons vary in their concern for:

*Harm/Care* Concern for others' well-being.

*Fairness* Concern for acting justly based on shared rules.

*Loyalty/ingroup concern* Concern for privileging one's own group over others'.

*Respect for authority* Concern for hierarchical structures, tradition, and obeying authentic authorities.

*Sanctity/purity* Concern for violating rules of sanctity of self or others; abhorrence of "disgusting" actions.

The MFT is thought to underpin many of the differences observed across the political spectrum. In particular, liberals appear to specifically value concerns about care and fairness, whereas conservatives are concerned about all five domains (J. Graham, Haidt, & Nosek, 2009; Haidt & Graham, 2007). Care and fairness are referred to as the "individualizing" foundations, whereas ingroup concern, respect for authority, and sanctity concerns are referred to as the "binding" foundations. The binding foundations are named as such due to their theoretical role in maintaining ever-growing social structures, and the individualizing foundations due to the focus on individual agents as the source or target of moral behaviors.

The MFT has gained support from both observational and experimental research. For example, in the debate over stem-cell research, those opposing stem cell research evoked arguments based on sanctity for human life, whereas those for stem cell research evoked arguments based on care (Clifford & Jerit, 2013). When pro-conservative (or pro-liberal) stances are couched within binding (or individualizing) concerns, conservatives (and liberals) are further entrenched into their own sociopolitical beliefs (Day, Fiske, Downing, & Trail, 2014). Conservatives and liberals both give more to charities when the charities frame the mission in terms of binding and individualizing moral foundations, respectively (Winterich, Zhang, & Mittal, 2012).

Moral foundations theory is not without dissidents, however (Suhler & Churchland, 2011). Most notably, Kurt Gray and colleagues argue that the innateness and modularity of

MFT should be reevaluated. The theory of dyadic morality (TDM; Schein & Gray, 2018) suggests that moral reasoning involves two perceived agents, and ultimately moral judgment is contingent on the harm experienced by either party. Instead of conceptualizing moral reasoning as occurring across five distinct domains, the TDM assumes there is ultimately one domain (harm to an agent), and individuals merely vary in what is considered perceivably harmful. For example, instead of suggesting that conservatives view unpatriotic behavior as morally repugnant because it violates the loyalty moral domain per se, the TDM suggests that conservatives perceive unpatriotic behavior as harmful, and it is the perceived harm that makes unpatriotic actions morally repugnant. Indeed, conservatives tend to rate purity violations among other acts as highly harmful (Gray & Keeney, 2015; Schein & Gray, 2015), with perceived harm mediating the relationship between an action and the rating of the action as immoral.

Whether MFT accurately describes our moral structure, or whether TDM accurately argues that immorality is equivalent to perceived harm, contemporary thinking on morality informs how individuals vary in their judgments of other groups' and their actions. Both political extremes utilize different moral foundations, different information, or different mechanisms of harm, which ultimately may lead to polarized opinions on the morality of immigration, and in particular Arab immigration. Conservatives may focus on the fairness to extant citizens' immediate economic well-being, whereas liberals may focus on the immigrants' well-being. Conservatives may focus on whether Arabs are directly harmful to citizens, whereas liberals may focus on the harm to Arabs. Conservatives may focus on the importance of maintaining a homogenous society, whereas liberals may emphasize a pluralistic society. Conservatives may focus on the values of loyalty, whereas liberals may focus on fluidity. Conservatives may focus on the sanctity and purity of some domains for the sake of disease avoidance (e.g., body, religious behaviors, as a behavioral immune system; Terrizzi, Shook, & McDaniel, 2013), whereas liberals focus on other domains (e.g., environmentalism, food practices; Feinberg & Willer, 2013).

From the MFT perspective, those who strongly value ingroup concerns, respect for authority, and purity may be threatened by the possibility of immigrants, *especially* when the immigrants hold perceivably different social and religious values from the concerning ingroup. For instance, if a right-wing person perceives Arabs to be predominantly Muslim (a “competing” religion with perceivably different moral values), who are loyal to a separate outgroup and set of ideologies (e.g., “sharia law”), and who wish to enact such ideologies across their nation, then this hypothetical person would perceive Arabs and Muslims to be morally threatening within all the binding foundations. Moreover, to the extent that Arab individuals are believed to be violent and reactionary, then Arabs would be perceived to violate both harm and fairness domains. Without invoking MFT, the principle holds from the perspective of TDM that Arabs would be perceived as moral threats to the extent that a person believes that Arab social ideologies and behaviors are harmful to others and to society. Regardless of the factual accuracies and inaccuracies of this information, Western people subscribing to such information will view the Arab people as a threat to their moral values.

Conversely, a left-wing person may operate on different moral domains (the individualizing foundations), or merely believe that a different set of ideas and behaviors are harmful. In the case of a left-wing person, discriminating against Arabs and Islamic belief may itself be a fairness violation, and disallowing Arab refugees from immigrating would be a harm violation. From the TDM perspective, a left-wing person may not believe that pluralistic societies are inherently threatening, that authority structures are necessary for social living, nor that people must be loyal to a particular group; in which case, Arab immigrants are not viewed as harmful nor a moral threat — Quite the opposite, disallowing immigration would be perceived as harmful.

Neither the MFT nor the TDM aim to explain how morally relevant information about Arabs originates or perpetuates. Nevertheless, both theories provide a formal understanding of how political ideologies strongly covary with moral concerns, which then interact with

information about Arab people to produce strong moral judgments and political values about Arab people. The interactive nature of moral concerns with information is further compounded by the tendency for those on the political extremes to selectively consume media that correspond to their political ideology and share their moral concerns (Bolin & Hamilton, 2018; Crawford et al., 2013), further entrenching individuals into their respective ideologies. Consequently, the ideologies espoused by Trump and Sanders may prompt differing moral concerns and morally relevant information to produce differing attitudes toward Arabs and Muslims.

#### *1.4 Authoritarianism and Domination*

Two additional constructs are relevant to understanding prejudice, and in particular anti-Arab prejudice: Right-wing authoritarianism (RWA) and social dominance orientation (SDO). RWA refers to the principle that social order and security should be preserved (Altemeyer & Hunsberger, 1992), with those high in RWA tending to endorse stereotypes and report heightened threat from outgroup members (Cohrs & Asbrock, 2009). Similar to RWA, SDO refers to the principle that groups are inherently competitive with one another, and one group is necessarily preferred over other groups (Pratto et al., 1994). In a sense, those high in RWA fight to preserve the extant social order, norms, and morals for the sake of cohesion and security, whereas those high in SDO wish to exert control over order, others' norms, and morals in order to secure their own group. For example, RWA predicts greater prejudice against groups manipulated to appear more competitive and threatening, whereas SDO did not predict *greater* prejudice against such groups (Cohrs & Asbrock, 2009), implying that those high in SDO already perceive other groups to be competitive or threatening. RWA predicts support for the prohibition of hate speech, but SDO does not (Bilewicz, Soral, Marchlewska, & Winiewski, 2017). RWA predicts negative ratings for media mentioning same-sex marriage (a symbolic threat to cultural norms), but SDO predicts negative ratings for media mentioning affirmative action (Crawford et al.,

2013). Consistently, RWA is associated with high threat reactions and world-as-dangerous worldviews, whereas SDO is associated with perceiving the social world as a zero-sum competition (Perry et al., 2013).

Both SDO and RWA are consistently predictive of both implicit and explicit prejudicial attitudes toward Arab persons (Rowatt et al., 2005). RWA predicts explicit negative attitudes toward Islamic immigrants and Arabs (Altemeyer & Hunsberger, 1992; Echebarria-Echabe & Guede, 2007; Manganelli Rattazzi, Bobbio, & Canova, 2007). The consistent relationship between religious fundamentalism and anti-Arab prejudice is mediated through RWA (Johnson et al., 2012). SDO predicts explicit generalized prejudice (Dunwoody & McFarland, 2018; McFarland, 2010) and support for anti-Muslim policies (Dunwoody & McFarland, 2018). SDO predicts greater justification for violence in the middle east (Henry et al., 2005). SDO predicts greater Islamophobia, mediated through perceived realistic and symbolic threats (Uenal, 2016). Indeed, a meta-analysis of 71 studies suggests that SDO and RWA correlate  $r = .55, .49$  respectively with prejudice measures (Sibley & Duckitt, 2008).

SDO and RWA are also both strongly associated with political leanings and moral foundations endorsements. In particular, those who identify as conservative tend to be higher in both SDO and RWA (Pratto et al., 1994). The difference between conservatives and liberals in emphasis on the moral binding foundations can be adequately explained by differences in RWA and SDO (Kugler, Jost, & Noorbaloochi, 2014). Furthermore, latent profiles of moral foundation responses suggest that those high in SDO tend to endorse all foundations equally, albeit moderately, whereas those high in RWA tend to strongly endorse all — Suggesting those high in RWA are concerned about moralizing, whereas those high in SDO are not (Milojev et al., 2014). In an effort to reframe the moral foundational differences between conservatives and liberals, an evolutionary-coalitional theory was examined by Sinn and Hayes (2017). The data suggested that the binding foundations include the theory-predicted outgroup antagonism and high threat sensitivity, whereas the individualizing foundations

included a sense of universalism. Importantly, universalism correlated negatively with SDO, and the dominating motive correlated strongly with authoritarianism. Conservatism was more sufficiently explained by SDO and authoritarianism than by the moral foundations alone (Sinn & Hayes, 2017).

To the degree that Trump (Sanders) represents an ideology laden with high (low) RWA and SDO, exposures to Trump and Sanders may evoke information, social goals, and subsequent judgments consistent with high and low RWA and SDO. Consequently, exposure to Trump may evoke expressions of anti-Arab and outgroup prejudices, whereas exposure to Sanders may not.

### *1.5 Proposed Studies*

The present paper explores whether these political icons affect automatic social judgments merely due to one's exposure to these icons and their respective stances. Capitalizing on the contemporary political and social zeitgeist, we focus on social judgments of Arab (the target group) and Chinese groups (control group) through a minimally representative paradigm. We explore the effects of these political icons and sociopolitically important personality variables (i.e., SDO and RWA) on the intuitive and implicit judgments of Arabic and Chinese symbols using the novel, minimal paradigm. In essence, we explore whether anti-Arab prejudice can be detectable in the evaluations of benign linguistic symbols, and how these evaluative prejudices are influenced by exposure to politico-ideological icons and sociopolitically relevant individual difference variables.

Three studies were conducted to explore the effect of these particular political icons and constructs on anti-Arab prejudice. Each study employed a modified affective misattribution procedure (AMP), a well studied and popular method for assessing implicit attitudes toward given stimuli (Gawronski & Ye, 2013; Payne, Cheng, Govorun, & Stewart, 2005; Payne, Hall, Cameron, & Bishara, 2010; Payne & Lundberg, 2014). The modified method allows for the detection of prejudice in a minimal manner, such that prejudice can be ob-

served even in the ratings of Arabic or Chinese script. Using this procedure, we assessed implicit attitudes toward Trump and Sanders, and toward Arabic and Chinese symbols. Importantly, we examined the interactions therein, such as the implicit prejudicial attitudes toward an Arabic symbol after being exposed to either Trump or Sanders. To better understand why individuals implicitly evaluate the candidates and the symbols as they do, we measure social dominance orientation (Pratto et al., 1994), right wing authoritarianism (Manganelli Rattazzi et al., 2007), explicit anti-Arab attitudes (Echebarria-Echabe & Guede, 2007), and various political attitudes. We assessed how these traits and stimuli influence not only judgments of Arabic symbols, but of both Arabic and Chinese symbols relative to random meaningless glyphs.

The method and data analytic tools provide a robust assessment of anti-Arab prejudicial responses. If the predicted effects can be consistently detected and replicated using the subtle stimuli and indirect assessment of prejudicial attitudes, both which attenuate effect sizes, then presumably anti-Arab prejudice is a large and robust phenomenon. Moreover, the fitted models are maximal models that permit uncertainty where possible (i.e., accounting for subject variance, measurement error). The combination of the minimal method with a model incorporating fully propagated uncertainty permits risky predictions, or severe tests, of the hypotheses, and helps establish boundaries for the anti-Arab prejudice effects.

A stimulus pilot and three studies were conducted. All studies employed the same methodology with similar stimuli. Participants were exposed to several modified AMP trials. On each trial, participants saw a picture of Sanders, Trump, or a control. Following a delay, participants saw either random Arabic script or a random Chinese character. Participants were instructed to rate the linguistic symbol as either unpleasant/pleasant, or peaceful/aggressive, the response options which varies between subjects. Study 3 participants additionally rated control linguistic characters consisting of randomly generated glyphs corresponding to no language, and thus no social groups. A pilot study was performed in order to verify that the population sampled interprets Arabic script as Arabic, Chinese

script as Chinese (or at least east-Asian), and the random glyph as unknown, to ensure that the ratings of the characters correspond to the perceived group.

We propose the following hypotheses:

- (1) Arabic symbols will receive less positive ratings than Chinese symbols. (A priori in study 1).
- (2) Exposure to Trump will moderate the effect, such that Arabic symbols will receive especially less positive ratings than Chinese symbols following a Trump prime. (A priori in study 1).
- (3) Exposure to Sanders will moderate the effect in the opposite direction, such that Arabic symbols will receive ratings equal to those of Chinese symbols following a Sanders prime. (A priori in study 1).
- (4) Arabic symbols will receive less positive ratings than Chinese symbols especially if one is high in SDO and RWA. (Confirmatory in study 2 and 3).
- (5) Judgments following a Trump prime will be moderated by SDO, RWA, and anti-Arab prejudice, such that individuals high in such traits will respond more positively following a Trump prime. (Confirmatory in study 2 and 3).
- (6) Judgments following a Sanders prime will be moderated by SDO, RWA, and anti-Arab prejudice, such that individuals low in such traits will respond more positively following a Sanders prime. (Confirmatory in study 2 and 3).
- (7) SDO and RWA will predict positive judgments of random symbols compared to both Arabic and Chinese symbols. (A priori in study 3).
- (8) Generalized prejudice will predict positive judgments of random symbols compared to both Arabic and Chinese symbols.

## CHAPTER TWO

### Stimulus Pilot

The purpose of the pilot study is not to test any substantive hypotheses or estimate effects, *per se*. Instead, the purpose of the pilot is to examine the experimental materials and methods in order to refine the inferences permissible by studies 1–3.

The importance of studies 1–3 is predicated on the assumption that any negative ratings of Arabic symbols relative to other symbols is due to anti-Arab prejudice, thus suggesting that anti-Arab prejudice is observable in a minimal manner in the mere ratings of effectively an Arabic font. In order for this predicate to be valid, individuals must recognize the Arabic script as Arabic, and Chinese script as Chinese. The pilot aims to assess the degree to which individuals categorize each symbol as the correct category.

A second aim to assess how well individuals explicitly recognize Arabic and Chinese symbols in the paradigm used for studies 1–3. Ideally, any observed effects of symbol type on the rating is due to an intuitive, associative link between the symbol type and prejudicial attitudes toward the group corresponding to the symbol. However, it is possible that individuals readily recognize the quickly presented symbols and have enough time to reflectively (instead of reflexively) indicate an explicit response. Therefore, participants will be tasked with simply sorting each symbol as Arabic or Chinese under the timings used throughout studies 1–3.

#### *2.1 Methods*

Two hundred participants were recruited from Amazon MTurk and paid \$.50 for their participation. Participants were randomly assigned to one of two tasks for the purposes of validating stimuli used in studies 1–3.

### 2.1.1 Stimulus Creation

All studies used a mask, control image, Trump images, Sanders images, Chinese symbols, and Arabic symbols; study 3 additionally required random glyphs. The mask was a black and white pattern with both local and global noise. The control image was a gray square. Trump and Sanders images were sourced from Google image searches. We attempted to match all images of the politicians in emotional expressions and angle of the picture. Twelve images of each person were selected, yielding a total of 25 prime images (one control, 12 Trump, 12 Sanders). All images were cropped to a square dimension, and resized to 300x300 pixels for consistency and to improve the speed of loading each image. See Appendix C for the prime images and mask to be used.

The symbols presented were either Chinese pictographs or Arabic “words.” To generate the symbols, a lorem ipsum generator (<http://generator.lorem-ipsum.info/>) generated random strings of Chinese and Arabic script. Each unit was then separated (an Arabic word defined as a space-separation, a Chinese symbol) and a custom script converted each unit to a 300x300 px image with a transparent background. This yielded over a hundred possible Arabic or Chinese images. From this pool of images, we selected those that are most recognizably Chinese glyphs (e.g., 秋, as opposed to ☒) and most recognizably Arabic script (e.g., متماقة, as opposed to و). See Appendices D and E for the Arabic and Chinese symbols.

Hundreds of random glyphs were generated for use in study 3 using a modified free tool (Lourenço, 2015). The goal was to select and use symbols that are complex enough to be plausible linguistic symbols, but unique enough to not resemble any known language by accident. See appendix F for examples of random glyphs.

### 2.1.2 Procedure

The first task (explicit) was as follows. Participants saw each symbol (Arabic, Chinese, or random glyph) once. They were asked to categorize each symbol as Arabic, Hindi, Chinese, Korean, Vietnamese, English, German, or none. The reason these languages were

chosen specifically was to discern not only whether individuals can accurately categorize Arabic and Chinese, but specifically whether they can discern Arabic and Chinese from script that a naïve English speaker may not otherwise know. In particular, English speakers may not know the difference between Arabic and Hindi, nor discern Chinese from Korean or Vietnamese. Understanding how participants explicitly categorize these symbols provides insight into how to interpret the effects; for instance, if participants are poor at distinguishing Chinese from Korean or Vietnamese, then any effects with regard to the Chinese symbols may need to be interpreted as a comparison with East Asian cultures more broadly, rather than Chinese cultures specifically. Similarly, if participants categorize Arabic symbols as Arabic more frequently than Hindi or other languages, we can be reasonably certain that any effects with Arabic symbols are specific to participants' conception of Arabic cultures, as opposed to cultures in the Central and South East Asia more broadly. Participants then completed a similar task wherein they saw each image of Sanders and Trump, and were simply asked to indicate which images are of Trump and of Sanders.

The second task (implicit) was as follows. Participants engaged in a speeded modified AMP similar to studies 1–2. In brief, each trial presented a fixation point, a mask, a political icon (or a gray square), a symbol, and a mask (See chapter 3 methods for specific details). Participants indicated whether the presented symbol was Arabic or Chinese. The goal is to assess whether individuals recognize Arabic and Chinese symbols in the short 200ms exposure time to the symbol. Subsequently, participants repeated this task, but instead responded with whether they saw Trump or Sanders, instead of Arabic or Chinese symbols.

Prior to the explicit task and following the implicit task, participants were shown images of Trump and Sanders, and asked what each respective political icon's political leanings are for foreign policy, economic issues, social issues, and religious issues each on scales ranging from 1 (Very Liberal) to 7 (Very Conservative). The purpose of these questions is to examine whether participants are familiar enough with Trump and Sanders

to associate them with particular political leanings, based only on their image and not their name.

No predictions were made, but the ideal results for interpretation would be that participants in the *first* task accurately recognize Arabic symbols as Arabic, and Chinese symbols as *at least* East Asian; and participants in the *second* task do not easily recognize the symbols as Arabic or Chinese — I.e., an accuracy rate of approximately 50%. To clarify, if participants do not accurately categorize Arabic and Chinese symbols when provided ample time and options in the first task, then it would be unclear what differential responding toward Arabic and Chinese symbols would reflect. Any useful inferences about why Arabic symbols may receive worse intuitive ratings than Chinese symbols depends on what cultures the sampled population associates with such symbols. For example, if participants tend to indicate that the Arabic symbols are Hindi, then any differential responding may not be due to an anti-Arab attitudes, but rather anti-Hindi attitudes. Conversely, if participants are readily able to accurately categorize Arabic and Chinese symbols in a short timespan and with both forward and backwards masking (the second task), then the task may fail to be a solely *implicit* measure of attitudes toward the symbols. That is, if participants can explicitly recognize and reflect on the symbols, then participants' responses may be a function of both implicit and explicit attitudes, rather than implicit attitudes alone. Together, these tasks explore whether the stimuli used in studies 1–3 are associated with the target social groups (task 1) and whether the evaluation of the stimuli is implicit or explicit (task 2).

## 2.2 Results

After excluding those who reported that they have experience reading Arabic or Chinese, the final sample size across both tasks was 160,  $M_{\text{age}} = 38.30$ ,  $SD_{\text{age}} = 12.82$ , 43.13% male.

Data from each task were analyzed separately. The first task (explicit recognition,  $N = 90$ ) includes both categorization of symbols and subsequently the categorization

of the political icons. A response-based modelling approach is taken, in which the raw responses are modeled using a multinomial logistic mixed model with random effects of images and participants. This permits insight into fixed response patterns for each symbol type (Arabic, Chinese, and Random), and into which images produce notably different response patterns than intended.

Only two responses are possible when categorizing political icons. Therefore, an accuracy-based logistic mixed model was performed in order to ensure both Sanders and Trump were categorized equally accurately, and whether any particular image is notably categorized incorrectly.

The second task ( $N = 70$ ) has only two response options for both linguistic symbols (Arabic, Chinese) and political icons (Trump, Sanders). For both types of stimuli, an accuracy-based logistic mixed model is employed, with random effects of subjects and image.

### 2.2.1 Explicit Task - Symbols

A multinomial Bayesian logistic model was fit to the raw responses, with conditional random effects of subjects and images, using the brms package (Bürkner, 2017). The multinomial logistic model simultaneously fits  $K - 1$  models, where  $K$  is the number of categories, and  $K = 8$  for this task. The model sets one category as a reference category (e.g., 1 = “Arabic”), and predicts the log-odds of a response  $k \in \{2, \dots, 8\}$ , relative to 1, where  $k$  differs per model.

$$\log \left( \frac{p(y = k)}{p(y = 1)} \right) = \hat{L}_k = \overbrace{X\beta_k}^{\text{Fixed}} + \underbrace{\overbrace{Zu_{0k}}^{\text{Subject}} + \overbrace{Av_{0k}}^{\text{Item}}}_{\text{Random effects}}$$

$$\hat{L}_1 = 0$$

These logits are then converted to probability using the softmax function, or equivalently:

$$p(y = k) = \frac{\exp(\hat{L}_k)}{1 + \sum_{k=2}^{K} \exp(\hat{L}_k)}$$

Table 2.1: Multinomial logistic mixed model expected values (and posterior standard deviations). Each row corresponds to a response category. Columns correspond to parameter estimates. The modeled value is the log-odds of choosing the response category, relative to choosing Arabic. “Not” is short for “Not a language”.

| Language   | Intercept     | Chinese    | Random     | $\sigma_{\text{subject}}$ | $\sigma_{\text{Item}}$ |
|------------|---------------|------------|------------|---------------------------|------------------------|
| Arabic     | —             | —          | —          | —                         | —                      |
| Hindi      | -2.27 (.26)   | 3.32 (.23) | 3.38 (.20) | 2.23 (.24)                | .15 (.08)              |
| Chinese    | -3.19 (.20)   | 7.74 (.23) | 3.83 (.22) | 1.54 (.14)                | .25 (.06)              |
| Korean     | -3.37 (.24)   | 6.75 (.22) | 3.90 (.21) | 1.94 (.18)                | .09 (.06)              |
| Vietnamese | -4.46 (.31)   | 6.58 (.21) | 4.20 (.24) | 2.48 (.26)                | .19 (.10)              |
| English    | -10.31 (1.54) | 5.55 (.33) | 3.81 (.30) | 6.45 (1.31)               | .24 (.16)              |
| German     | -7.51 (.89)   | 5.15 (.32) | 3.76 (.29) | 4.06 (.73)                | .24 (.16)              |
| Not        | -5.14 (.34)   | 4.89 (.28) | 9.90 (.29) | 2.62 (.24)                | .45 (.09)              |

The fixed effects are dummy-coded indicators for whether the image is Chinese ( $\beta_1$ ) or Random ( $\beta_2$ ), with Arabic serving as the baseline ( $\beta_0$ ). The random effects therefore represent subject-specific and image-specific effects on expected response probabilities. The coefficients are reported below in Table 2.1, with posterior distributions of indicating each response for each item visualized in Figure 2.1.

Based on this model, Arabic images have a very high probability (.83 [.78,.87]) of being categorized as Arabic. Chinese images were categorized predominantly as Chinese (.68 [.58,.78]), with Korean as the next most probable response (.21 [.14, .31]). Finally, random glyphs were categorized as “Not a language” with a very high probability (.93 [.88, .96]). No images exhibited unique adverse effects on categorizing as intended — The item-specific effects were small and did not change the predicted response category.

### 2.2.2 Explicit Task - Political Icons

Responses were recoded to either correct (1) or incorrect (0). A logistic mixed model was fit with a dummy-coded fixed effect indicating whether the image was of Trump (1) or Sanders (0). Images and subjects were added as random effects, just as in the multinomial model.

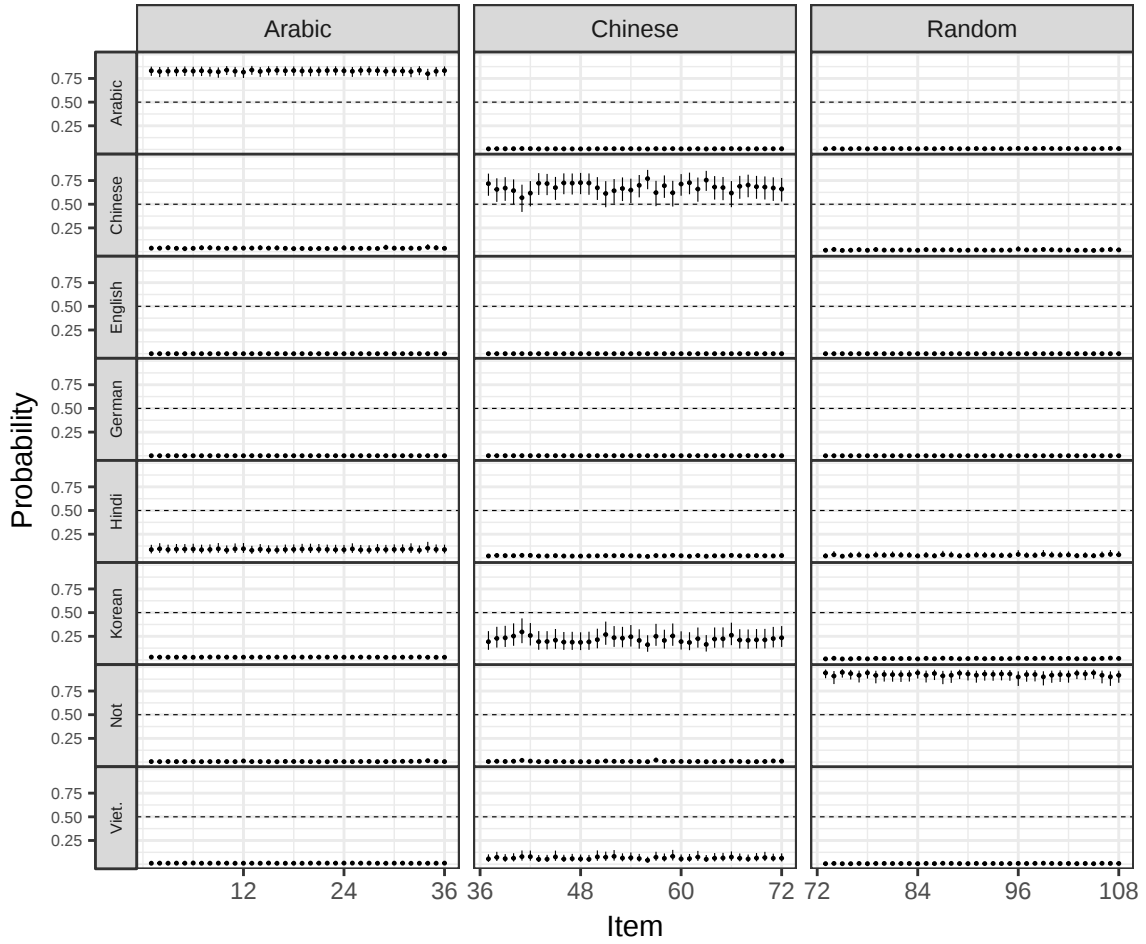


Figure 2.1: Probability of responses for each item, from the multinomial model. Items 1–36 are Arabic images, 37–72 are Chinese, 73–108 are random glyphs. Rows correspond to the response categories.

The intercept ( $\beta_0 = 12.96 [7.92, 22.57]$ ) represents the log-odds of being correct when the image is of Sanders, and indicates that participants were highly likely to be correct on average. The coefficient ( $\beta_1 = .74 [-.19, 1.72]$ ) represents the difference in log-odds when the image is of Trump, and suggests that the difference, should it exist, is minimal. Item-specific effects are minimal ( $\sigma_{\text{Item}} = .30 [.01, .87]$ ), with no images exhibiting a notable item-specific effect. Subject-specific effects ( $\sigma_{\text{Subject}} = 6.16 [2.30]$ ) appear large, but the variance estimate is predominantly affected by six individuals who indicated only one response option for all images (thereby each having an estimated accuracy of approximately .50).

These individuals were not removed from the model, simply because the subject-specific variability was accounted for, and removing such participants merely reduced the variance estimate (i.e., removal of the participants did not affect the expected response probabilities). Based on this model, it is safe to assume that individuals from this population indeed do accurately recognize Trump and Sanders.

### 2.2.3 *Implicit Task - Symbols*

Responses were recoded to either correct (1) or incorrect (0). A logistic mixed model was fit with a dummy-coded fixed effect indicating whether the symbol was Chinese (1) or Arabic (0). Images and subjects were added as random effects.

The intercept ( $\beta_0 = 3.32 [2.93, 3.73]$ ) represents the log-odds of being correct when the image is Arabic, suggesting that participants accurately recognized Arabic script. The coefficient ( $\beta_1 = .07 [-.32, .18]$ ) represents the difference in log-odds when the image is Chinese; sign is ill-determined, and any value within this 95% HDI implies an irrelevant effect size (e.g.,  $< 1\%$  difference in accuracy). Subject-specific effects appear notable ( $\sigma_{\text{Subject}} = 1.39 [1.11, 1.74]$ ), but is largely driven by five participants who only indicated one response, yielding a predicted .5 probability of accuracy. Item-specific effects are minimal ( $\sigma_{\text{Item}} = .18 [.01, .39]$ ), again with no images exhibiting a notable item-specific effect.

### 2.2.4 *Implicit Task - Political Icons*

Responses were recoded to either correct (1) or incorrect (0). A logistic mixed model was fit with a dummy-coded fixed effect indicating whether the symbol was Sanders (0) or Trump (1). Images and subjects were added as random effects.

The intercept ( $\beta_0 = 3.05 [2.40, 3.79]$ ) suggests that participants accurately recognized images of Sanders, and the ability to recognize Trump was not notably different ( $\beta_1 = -.25 [-.63, .13]$ ). Similar to previous models, variability among participants ( $\sigma_{\text{Subject}} = 2.19 [1.67, 2.88]$ ) was largely high due to seven participants who simply indicated identical

Table 2.2: Estimated political leanings of Sanders and Trump on a scale ranging from 1 (Very liberal) to 7 (Very conservative) in four domains.

| Icon    | Foreign Policy    | Economic Issues   | Social Issues     | Religious Issues  |
|---------|-------------------|-------------------|-------------------|-------------------|
| Sanders | 2.59 [2.41, 2.78] | 2.56 [2.39, 2.74] | 2.52 [2.33, 2.70] | 2.56 [2.39, 2.73] |
| Trump   | 5.77 [5.59, 5.96] | 5.74 [5.56, 5.92] | 5.70 [5.51, 5.88] | 5.74 [5.56, 5.92] |

answers for all images, with a predicted .5 probability of accuracy. No notable image-specific effects existed, and variability among images was negligible,  $\sigma_{\text{Item}} = .22$  [.01, .50].

### 2.2.5 Political Leanings

Given that multiple raters (participants) rate two targets on four domains, a mixed model with random effects for the subject and domain is appropriate. This model permits shrunken estimates of expected ratings for each political icon, on each political domain. Subject-specific rating variance was small ( $\sigma_{\text{Subject}} = .58$  [.47, .71]), suggesting that participants indicated similar political leanings. Likewise, domain-specific variance was small ( $\sigma_{\text{Domain}} = .11$  [0, .46]), suggesting that participants rated each political icon consistently. Estimates are available in Table 2.2.

## 2.3 Discussion

The stimuli were all categorized as intended, with negligible variation among subjects and specific stimuli. Unfortunately, subjects were able to accurately categorize the stimuli in the speeded task. This implies that the stimuli are presented in a manner that could yield explicit judgments, rather than purely implicit judgments. This will be discussed in future chapters.

Subjects consistently associated Sanders and Trump with liberal and conservative politics, respectively. Moreover, they could do so using only face recognition. This implies that subjects not only recognized both Sanders and Trump by image, but do associate them with strongly differing political views.

Chinese symbols were secondarily categorized as Korean. Because subjects from this population often fail to distinguish the two, any effects specific to Chinese symbols should be potentially broadened to effects of East Asian societies. No particular stimulus exerted undue influence on categorization, therefore the stimuli were used across all studies.

## CHAPTER THREE

### Study 1

The purpose of study 1 is to examine whether anti-Arab prejudicial attitudes is observably manifest in ratings of Arabic and Chinese symbols in an undergraduate population. To reiterate, Arabic symbols were expected to receive less positive ratings than Chinese symbols (H1). Exposure to Trump (Sanders) was expected to moderate this effect, such that Arabic symbols would receive especially less (more) positive ratings following a Trump (Sanders) exposure (Hypotheses 2 and 3, respectively). SDO and RWA were measured to assess the tentative hypotheses that both variables would strengthen the Arabic effect, and perhaps predict positive responses following Trump images, and negative responses following Sanders images. For the sake of transparency, these latter predictions were primarily exploratory in study 1, but confirmatory in studies 2 and 3. Nevertheless, for the sake of consistency, they are referred to as tested hypotheses here.

Participants were recruited from the Baylor university undergraduate research pool in exchange for participation credit. Only those in the introductory psychology courses were recruited. Due to the partially exploratory nature of study 1, sample size was determined simply by the number of participants obtained by the end of the semester. The within-subject design should provide ample power to detect the directional effects of interest — with 72 observations per person and an estimated 100 persons per semester, this amounts to approximately 7200 observations to assess the effects of the priming conditions.

To clarify this, several simulations were conducted in order to estimate the power of detecting the primary negative Arabic effect across multiple parameter settings. In the most conservative condition, the following parameters were set. The Arabic effect parameter was set to a small value of  $-.2$ , and all parameters were given random effects with a standard deviation of  $.3$ . Under this condition,  $.8$  power is obtained with only 80 participants. If random

effect variance is smaller, or the parameter is larger, then power substantially increases. For instance, if random variation for the Arabic effect remains .3 standard deviations, but the effect is -.4, only 10 subjects are required for .8 power.

### *3.1 Methods*

#### *3.1.1 Experimental Materials*

All study materials were presented to the participant through a series of webpages created through Qualtrics and custom javascript. Participants completed the study on a laptop or desktop computer, and this was enforced by disallowing those on a mobile browser from continuing in the study and by requiring participants to press the space bar on a webpage with no text fields; this procedure effectively disallows those on phones and tablets from continuing in the study. This requirement is chosen because the stimuli presented in the experimental task were presented rapidly and require a fast response by the participant. Because mobiles and tablets do not have adequate hardware for millisecond precision of screen refreshes, and touch responses are slower than hardware responses, those on mobiles and tablets were disallowed from participation.

The experimental task was a modified AMP using the stimuli examined in the pilot study. For each trial, a fixation point was presented for 1000ms, a political icon is displayed for 200ms, followed by a blank for 100ms, a symbol for 200ms, then finally a scrambled black and white visual mask. The timings were based on previous recommendations (Payne & Lundberg, 2014) and from in-lab pilot testing of the stimuli timings and clarity on various Windows, OSX, and Linux desktops and laptops, and across several browsers (e.g., Firefox, Google Chrome, Apple Safari). The priming image was either a gray square (Control), an image of Bernie Sanders, or an image of Donald Trump. The symbols presented after the blank were either Chinese pictographs or Arabic “words.” A visual description of the paradigm is available in Figure 3.1.

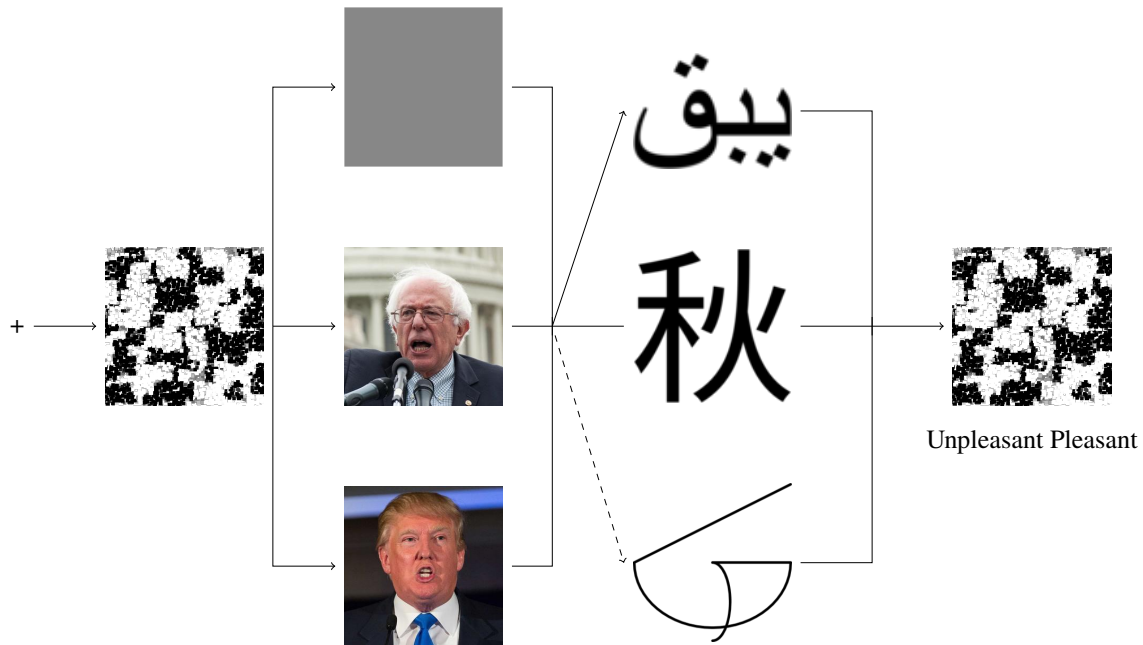


Figure 3.1: Experimental flow of studies 1–3. The diagram represents the order of stimuli seen per trial. The dashed line represents the stimulus only observed during study 3.

Finally, to ensure adequate presentation speeds of the stimuli, a page was created that precaches all images into the browser. Doing so allows the presentation of the stimuli to depend only on the speed of the participants’ computer hardware, rather than depend on the users’ and file servers’ connection speeds. Modern browser animation methods were used to ensure stimuli presentations were millisecond-precise, appeared in the correct order, and were synchronized to the monitor refresh rate. The observed stimulus timings depend on the participants’ hardware capabilities and monitor response time. Nevertheless, to the degree possible, every effort was taken to ensure maximal consistency and speed of the programmed task across participants.

### 3.1.2 Procedure

Data collection occurred across the Spring semester, 2016 (Specifically, from March until May). After obtaining consent and ensuring the participant used an allowed device for the study, participants began the experimental task. Participants were instructed to rate

the *symbol* they saw using the aforementioned response options, as quickly as possible, and to “try not to let the picture affect your rating of the symbol” (Payne & Lundberg, 2014). Participants indicated their responses using the “q” key (Left response) or the “p” key (Right response) to facilitate faster responding. Participants were randomly assigned to one of two response option conditions. In one condition, the response options on each trial were “Unpleasant” or “Pleasant”, and in the other condition the response options were “Peaceful” or “Aggressive”. The former response option represents an intuitive affective response, whereas the latter represents a potentially more intuitive cognitive response.

In total, each participant engaged in the task for 72 trials. Each prime condition and symbol condition was balanced (12 of each prime condition, 36 of each symbol condition). Each trial consisted of a prime image (from one control, 12 Sanders, 12 Trump), and a symbol (from 36 Arabic images, 36 Chinese images). Therefore, each participant saw any given symbol only once and each non-control prime image twice (once paired with an Arabic symbol, once paired with a Chinese symbol). In order to ensure that any one given symbol-prime pairing was not driving the effects of interest, three counterbalancing conditions were utilized that differ only in which symbol is paired with which conditions. Participants were randomly assigned to one of the three counterbalancing conditions.

Following the experimental task, participants completed a modified<sup>1</sup> Anti-Arab prejudice scale ( $\alpha = .95$ ; Echebarria-Echabe & Guede, 2007), a measure of political attitudes and beliefs about the two candidates (See Appendix A;  $\alpha = .94, .96$ ), the SDO scale ( $\alpha = .94$ ; Pratto et al., 1994), the RWA-R scale ( $\alpha = .87$ ; Manganelli Rattazzi et al., 2007), and a demographics questionnaire. Three simple attention check questions were included (e.g., “Are you paying attention? Click 3 for this item”). Participants were then be directed to a page debriefing them on the purpose of the study. The amount of time spent

---

<sup>1</sup> The Anti-Arab scale (Echebarria-Echabe & Guede, 2007) is lengthy and includes items relevant only to Europeans and irrelevant to the American population. Instead, the current studies used 17 items from the scale with high factor loadings without references to Europe or issues relevant only to Europe. See Appendix B for these items.

specifically on the experimental task was not collected, but the mean time to complete the study was approximately 33 minutes.

### 3.2 Results

Ninety-one participants were recruited for participation. Seventeen participants failed any one of the attention checks, and were removed. Ten participants indicated experience with reading or speaking Arabic or Chinese, and were removed. The final sample consisted of 64 participants ( $m_{\text{age}} = 19.91$ ,  $\sigma_{\text{age}} = 2.68$ , 33% male). All continuous variables were standardized to ease interpretation. Stan (Bürkner, 2017; Stan Development Team, 2017) was used for model estimation. The data were be split by outcome response type, and analyzed separately.

Several models were created, each to assess one of the stated hypotheses. However, the models all used the same Bayesian estimation technique, priors, and template. Across all models, the dichotomous symbol ratings were modeled as a logistic outcome, wherein 1 represents a positive rating (Pleasant; Peaceful) and 0 a negative rating (Unpleasant; Aggressive). Bayesian mixed logistic models were fit, with correlated random effects of subjects on the intercepts and slopes. The base model is described more completely as follows:

$$\begin{aligned}
y &\sim \text{Bernoulli}(\text{logit}^{-1}(\hat{y})) \\
\text{logit}(\hat{y}_i) &= \underbrace{X_i}_{n_i \times p} \underbrace{\beta_i}_{p \times 1} \\
\beta_i &= \underbrace{\beta}_{p \times 1}^{\text{Fixed}} + \underbrace{u_i}_{p \times 1}^{\text{Subject-effects}} \\
\beta &\sim \text{Normal}(0, 1) \\
\underbrace{u}_{N \times p} &\sim \text{MVN}(0, \Sigma) \\
\underbrace{\Sigma}_{p \times p} &= LL' = (dc)(dc)' \\
\text{diag}(\underbrace{d}_{p \times p}) &\sim \text{Student-t}^+(\nu = 3, \mu = 0, \sigma = 1) \\
\underbrace{c}_{p \times p} &\sim \text{LKJ}(3)
\end{aligned}$$

The priors are all weakly informative, and merely help regularize the estimate by effectively excluding implausible values and shrinking estimates away from extreme values (Gelman, Jakulin, Pittau, Su, et al., 2008; Gelman, Simpson, & Betancourt, 2017). Instead of placing a prior on the covariance matrix directly, priors are placed on random effect standard deviations and correlations. The “LKJ” prior is a spherical regularizing prior for correlation matrices (Lewandowski, Kurowicka, & Joe, 2009). The (Cholesky) covariance matrix is subsequently constructed for use in the hierarchical prior on  $u$ .

In addition to the base model, a normal-assumptive latent measurement model with fixed variance identification was constructed for any scales (e.g., SDO, RWA-R), and the latent scores were then included into the dichotomous response model as a predictor and interaction term for all  $p - 1$  predictors. All indicator variables were standardized prior to analysis, and any reverse-scored items were transformed to be forward-scored, for the purposes of directional identification. Trump support and Sanders support scales were constructed using items 1–6 and 13 for Trump, and items 14–19 and 26 for Sanders from the political candidate attitude measure (See Appendix A). In Bayesian latent variable

models, latent variables are treated no differently than missing data. The measurement model was estimated alongside the logistic model, and the “missing” latent score for each individual is included as a subject-level predictor.

$$\begin{aligned}
x_j &\sim \text{Normal}(\nu_j + \lambda_j \theta, \sigma_j) \\
\theta &\sim \text{Normal}(0, 1) \\
\nu &\sim \text{Normal}(0, 1) \\
\lambda &\sim \text{Normal}^+(0, 1) \\
\sigma &\sim \text{Student-t}^+(3, 0, 1) \\
\underbrace{\gamma}_{p \times 1} &\sim \text{Normal}(0, 2) \\
\text{logit}(\hat{y}_i) &= X_i \beta_i + X_i \theta_i \gamma
\end{aligned}$$

, where  $x_j$  is the  $j$ th indicator variable of a presumed latent variable. For some trial,  $i$ , this expands to:

$$\begin{aligned}
\text{logit}(\hat{y}_i) &= \beta_{0i} + \beta_{1i} \overbrace{x_{1i}}^{\text{Sanders}} + \beta_{2i} \overbrace{x_{2i}}^{\text{Trump}} + \beta_{3i} \overbrace{x_{3i}}^{\text{Arabic}} \\
&+ \beta_{4i} \overbrace{x_{4i}}^{\text{Sanders:Arabic}} + \beta_{5i} \overbrace{x_{5i}}^{\text{Trump:Arabic}} \\
&+ \gamma_0 \overbrace{\theta_i}^{\text{Latent Score}} \\
&+ \gamma_1 \overbrace{\theta_i x_{1i}}^{\text{Latent:S}} + \gamma_2 \overbrace{\theta_i x_{2i}}^{\text{L:T}} + \gamma_3 \overbrace{\theta_i x_{3i}}^{\text{L:A}} \\
&+ \gamma_4 \overbrace{\theta_i x_{4i}}^{\text{L:S:A}} + \gamma_5 \overbrace{\theta_i x_{5i}}^{\text{L:T:A}}
\end{aligned}$$

To reiterate, the logistic models were fit separately to each response type. However, all participants completed the same scales and corresponding items. Because of this, the logistic parameters were estimated per-response type, but *all* participants were used to estimate the measurement model parameters.

Because all analyses were conducted using Bayesian techniques, the estimates provided include the Bayesian expected value of the posterior distribution (EAP) [and the 95%

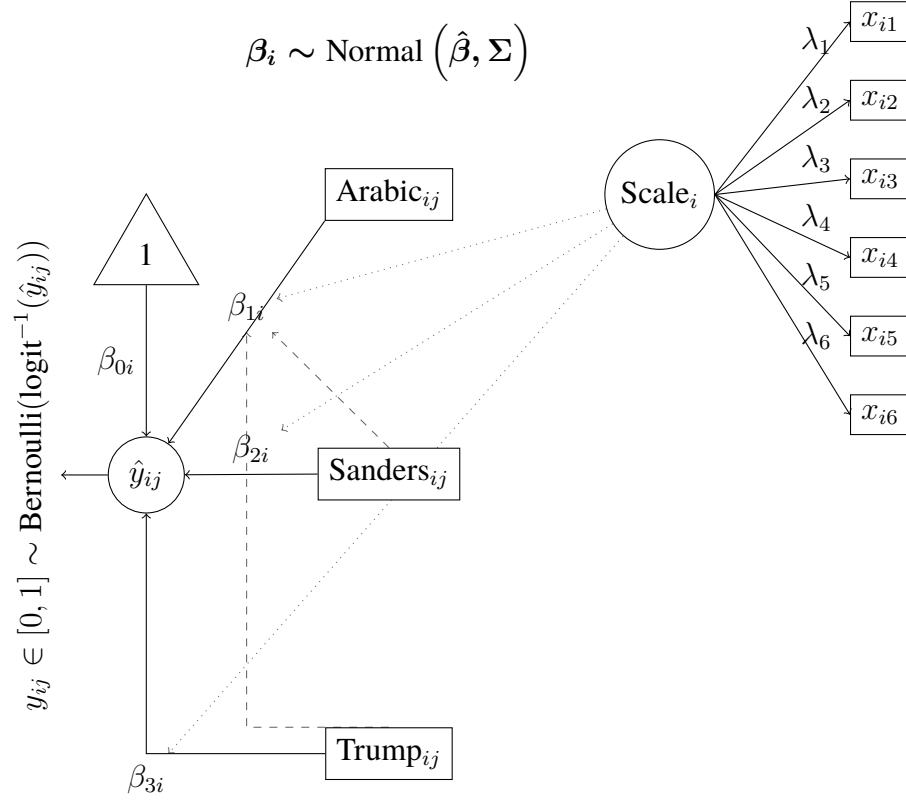


Figure 3.2: Model template graph. Interaction coefficients are not plotted for brevity.

posterior credible interval]. To compromise with convention, and because the hypotheses are directional, decisions will be formed primarily on the basis of whether one can be 95% certain of the sign of an effect, and thus whether the 95% credible interval excludes 0 or posterior probability of sign is sufficiently high; when the interval includes zero, but the reader may benefit from knowing the estimated probability of a particular sign, an additional estimate is be given called *ppp*, which communicates the posterior probability that the parameter is positive. See Figure 3.2 for a graphical depiction of the model template.

No data missingness was planned (e.g., J. W. Graham, Hofer, & MacKinnon, 1996). Less than 1% of all dichotomous observations were missing, and 14% of all responses across all scales were missing. The missing responses across scales appear uniform, such that nine participants did not complete any questions from any scale. After removing these nine participants, .03% of questionnaire responses were missing. Regardless, missing data

were handled within the model itself. Missing data in Bayesian models are formulated as merely another unknown quantity to be integrated over (Rubin, 1976). This permits posterior inferences on the parameters, conditioned on all known observed variables, after marginalizing over the uncertainty in the unobserved data. This approach is akin to full information maximum likelihood in the sense that it integrates over missing variables, and akin to multiple imputation in the sense that the Bayesian MCMC estimator randomly draws missing data from a conditional distribution.

Four chains were used, each with 1000 adaptation and 1000 post-adaptation iterations. No divergent transitions were reported, and all  $\hat{R}$  for all parameters  $\approx 1$ , indicating that the marginal posteriors for all parameters were sampled effectively and consistently between chains. All parameters of interest are reported in Tables 3.1 and 3.2.

The baseline model coefficients indicate the population-level expected log-odds of responding positively (i.e., pleasant, peaceful) given the trial condition. The Sanders effect is the expected change in log-odds due to exposure to Sanders, compared to control. The Trump effect is the expected change in log-odds due to exposure to Trump, compared to control. The Arabic effect is the expected change in log-odds due to the symbol being Arabic (and not Chinese), compared to the *mean* control response;  $2\beta_3$  is the difference between Arabic and Chinese symbols in the log-odds of a positive response. The two interactions are straightforward — S:A is the change in the Arabic symbol effect following a Sanders exposure, and T:A is the change in the Arabic symbol effect following a Trump exposure.

The first hypothesis posited that Arabic symbols should be rated as both less pleasant and less peaceful than Chinese symbols. This hypothesis is only partially supported, given the information available to the model. Although not definitive, the probability that Arabic symbols are rated as less pleasant on average in the population is considerably high (.93). When the response type is aggressive/peaceful, the estimated effect of Arabic symbols is too noisy to discern.

Table 3.1: Parameter estimates for the six models predicting pleasant and unpleasant responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. For example, the fifth column (RWA), fourth row (Arabic) indicates that for each one unit increase in RWA, the effect of Arabic symbol decreases by -.22. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is *positive* (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal).

| Parameter         | Baseline                          | AA                                | SDO                                 | RWA-R                             | TW                                 | SW                                |
|-------------------|-----------------------------------|-----------------------------------|-------------------------------------|-----------------------------------|------------------------------------|-----------------------------------|
| Control $\beta_0$ | <b>.68 (.97)</b><br>[-.01, 1.38]  | <b>-.51 (.07)</b><br>[-1.22, .17] | <b>-.44 (.10)</b><br>[-1.14, .24]   | <b>-.65 (.05)</b><br>[-1.46, .14] | <b>-.40 (.07)</b><br>[-.96, .15]   | .35<br>[-.36, 1.10]               |
| Sanders $\beta_1$ | <b>-.58 (.07)</b><br>[-1.32, .19] | .29<br>[-.45, 1.03]               | -.28<br>[-1.06, .53]                | .05<br>[-.86, .96]                | -.06<br>[-.71, .59]                | .38<br>[-.42, 1.18]               |
| Trump $\beta_2$   | <b>-.98 (.03)</b><br>[-1.97, .02] | <b>1.27 (1)</b><br>[.45, 2.16]    | <b>.79 (.94)</b><br>[-.20, 1.84]    | <b>1.40 (.99)</b><br>[.33, 2.53]  | <b>.94 (.99)</b><br>[.18, 1.73]    | <b>-.66 (.09)</b><br>[-1.74, .35] |
| Arabic $\beta_3$  | <b>-.34 (.07)</b><br>[-.78, .10]  | <b>-.47 (.01)</b><br>[-.93, -.05] | <b>-.65 (.004)</b><br>[-1.13, -.22] | -.22<br>[-.78, .36]               | <b>-.46 (.006)</b><br>[-.85, -.09] | <b>.50 (.98)</b><br>[.038, .98]   |
| S:A $\beta_4$     | <b>.24 (.93)</b><br>[-.08, .60]   | -.03<br>[-.42, .36]               | .04<br>[-.36, .47]                  | <b>-.29 (.10)</b><br>[-.74, .16]  | .08<br>[-.25, .42]                 | .06<br>[-.34, .48]                |
| T:A $\beta_5$     | .19<br>[-.15, .52]                | .07<br>[-.31, .46]                | .13<br>[-.29, .55]                  | .21<br>[-.27, .72]                | -.01<br>[-.34, .30]                | -.17<br>[-.58, .27]               |
| $\beta_0$         | 1.64                              | (.12)                             | (.035)                              | (.87)                             | (.68)                              | (.43)                             |
| $\beta_1$         | -.26                              | 1.82                              | (.99)                               | (.81)                             | (.25)                              | (.56)                             |
| $\beta_2$         | -.38                              | .46                               | 2.46                                | (.28)                             | (.30)                              | (.43)                             |
| $\beta_3$         | .24                               | .18                               | -.12                                | 1.01                              | (.32)                              | (.34)                             |
| $\beta_4$         | .16                               | -.22                              | -.16                                | -.14                              | .38                                | (.52)                             |
| $\beta_5$         | -.25                              | .05                               | -.06                                | -.13                              | .02                                | .38                               |

Table 3.2: Parameter estimates for the six models predicting peaceful and aggressive responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is *positive* (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal).

| Parameter         | Baseline                          | AA                   | SDO                                  | RWA-R                | TW                   | SW                               |
|-------------------|-----------------------------------|----------------------|--------------------------------------|----------------------|----------------------|----------------------------------|
| Control $\beta_0$ | <b>1.11 (.997)</b><br>[.31, 1.92] | -.59<br>[-1.72, .48] | <b>-1.04 (.015)</b><br>[-2.01, -.11] | -.58<br>[-1.51, .39] | -.42<br>[-1.48, .70] | <b>.50 (.91)</b><br>[-.25, 1.31] |
| Sanders $\beta_1$ | .03<br>[-.44, .54]                | .05<br>[-.68, .82]   | -.06<br>[-.72, .58]                  | .15<br>[-.47, .78]   | .09<br>[-.59, .75]   | .09<br>[-.41, .59]               |
| Trump $\beta_2$   | <b>-.74 (.03)</b><br>[-1.52, .06] | .34<br>[-.85, 1.44]  | .20<br>[-.78, 1.14]                  | .52<br>[-.40, 1.43]  | .85<br>[-.12, 1.89]  | -.30<br>[-1.05, .42]             |
| Arabic $\beta_3$  | .17<br>[-.34, .69]                | .33<br>[-.43, 1.08]  | .31<br>[-.32, .93]                   | .19<br>[-.47, .88]   | -.11<br>[-.85, .62]  | -.18<br>[-.71, .33]              |
| S:A $\beta_4$     | -.14<br>[-.48, .24]               | -.21<br>[-.76, .33]  | .05<br>[-.42, .52]                   | -.18<br>[-.64, .27]  | .06<br>[-.43, .55]   | .04<br>[-.33, .40]               |
| T:A $\beta_5$     | 0.00<br>[-.33, .31]               | .20<br>[-.29, .72]   | .34 (.94)<br>[-.08, .80]             | .20<br>[-.21, .62]   | .18<br>[-.25, .62]   | -.10<br>[-.42, .24]              |
| $\beta_0$         | 1.96                              | (.28)                | (.31)                                | (.13)                | (.72)                | (.35)                            |
| $\beta_1$         | -.17                              | .87                  | (.27)                                | (.45)                | (.21)                | (.49)                            |
| $\beta_2$         | -.11                              | -.17                 | 1.72                                 | (.66)                | (.71)                | (.43)                            |
| $\beta_3$         | -.25                              | -.03                 | .09                                  | 1.17                 | (.69)                | (.49)                            |
| $\beta_4$         | .20                               | -.27                 | .18                                  | .17                  | .39                  | (.49)                            |
| $\beta_5$         | -.14                              | -.01                 | -.06                                 | 0.00                 | .02                  | .23                              |

The second hypothesis posited that the Trump-Arabic interaction would be negative. This was not observed, regardless of response type. The moderating effect of a Trump trial on the Arabic effect is too noisy to discern, and therefore there is not enough information to evaluate the hypothesis.

The third hypothesis posited that the Sanders-Arabic interaction would be positive. The predicted positive moderating effect of Sanders trials on the Arabic effect is partially supported when the response type is pleasant/unpleasant ( $ppp = .93$ ). However, when the response condition is peaceful/aggressive, the estimate is too noisy to discern.

The fourth hypothesis posited that the negative Arabic effect would be strengthened (negatively moderated) by SDO and RWA. Explicit Anti-Arab prejudice, SDO, Trump support, and Sanders support all notably predict the Arabic effect. These effects (Table 3.1, row 4) suggest that those high in anti-Arab prejudice, SDO, and Trump support are especially less likely to rate Arabic symbols as pleasant, relative to Chinese symbols. Conversely, the Arabic effect is effectively negated for those high in Sanders support. Although the effect of RWA was in the expected direction, the sign and magnitude are both too noisy to discern. When the response type was peaceful/aggressive, the estimates are far less precise, and any effects are indiscernible.

The fifth hypothesis posited that judgments following a Trump exposure would be positively moderated by SDO, RWA, and anti-Arab prejudice. The population expected response following a Trump image is unpleasant and aggressive. However, the higher one is in anti-Arab prejudice, SDO, RWA, and Trump support, the more likely one is to respond “pleasant” following a Trump image, effectively negating the otherwise expected negative response. This suggests that those high in these constructs perceive Trump positively. These moderating effects were not discernible when the response condition was aggressive.

The sixth hypothesis posited the opposite following a Sanders exposure. Similarly, the population expected response following a Sanders image is unpleasant. However, no

moderating effects were discernible; the posterior distributions indicated uniformly high uncertainty.

Interestingly, anti-Arab prejudice, SDO, RWA, and Trump support all (to varying degrees of certainty) predict unpleasant responses. It is possible that those high in these factors generally find outgroup symbols unpleasant. This is explored further in study 3.

### 3.3 Discussion

The data were fairly consistent with the hypotheses, but only when the response options were unpleasant or pleasant. The primary hypothesis, that Arabic symbols would be rated as less positive than Chinese symbols, was partially supported. The expected direction of the effect was highly probable (.93), and suggests that Arabic symbols indeed were perceived as unpleasant more often than Chinese symbols. This effect was stronger for those high in explicit anti-Arab prejudice, SDO, and Trump support (and weaker for those high in Sanders support).

The anticipated negative interaction between Trump and Arabic images was not discernible, but the positive interaction between Sanders and Arabic images was highly probable. This suggests in trials with Sanders images, the negative Arabic effect was partially negated. Responses were generally “unpleasant” following either a Sanders or Trump image. However, those high in anti-Arab prejudice, SDO, and RWA were more likely to respond “pleasant” following a Trump image. This perhaps suggests that participants high in such traits perceive Trump and his respective ideology favorably.

The anticipated effects were not observed or discernible when the response options were aggressive or peaceful. It is possible that the response options were not amenable to the speeded task. Whereas one response condition is primarily affective (pleasant, unpleasant), the second is evaluative. Because “aggressive” and “peaceful” are more complex, participants’ responses may reflect predominantly reflective, and not reflexive, responses. The large uncertainty in parameter estimates could support the idea that the task is assessing

automatic evaluations, rather than reflective evaluations, of the symbols and icons, given that the affective response options yielded reasonably consistent estimates.

## CHAPTER FOUR

### Study 2

The purpose of study 2 is similar to that of study 1, but recruited a more representative online sample. Specifically, similar hypotheses are posited for study 2 as that in study 1, with the added goal of confirming any unanticipated results from study 1.

#### 4.1 *Methods*

Two hundred and fifty participants were recruited from the Amazon Mechanical Turk crowdsourcing platform (<http://www.mturk.com>) and paid \$.50 for their work. Based on the power analysis presented in study 1, 250 participants seemed sufficiently powered for detecting even small and variable effects with the current design. The experimental materials and procedures were identical to that of study 1. Data collection occurred across eight days starting June 20, 2016. Again, the time spent to complete the experimental trials was not collected, but the mean time to complete the study was 15 minutes.

#### 4.2 *Results*

Six participants failed one of the attention checks and were removed. Thirteen participants indicated experience with reading or speaking Arabic or Chinese, and were removed. An unknown problem occurred wherein 77 participants have uniformly missing data, and were removed. The final sample consisted of 154 participants ( $m_{\text{age}} = 39.67$ ,  $\sigma_{\text{age}} = 13.45$ , 38.31% male). All data preparation, transformations, and analyses conducted in study 1 were performed in study 2. No data were missing across the scale items, and only .15% (17) of the experimental trials were missing.

Just as in study 1, four chains were used with 1000 post-adaptation iterations. No divergent transitions were reported, and  $\hat{R}$  for all parameters  $\approx 1$ . All parameters of interest are reported in Tables 4.1 and 4.2.

Table 4.1: Parameter estimates for the six models predicting pleasant and unpleasant responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is *positive* (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal).

| Parameter         | Baseline                                 | AA                                  | SDO                                 | RWA-R                                | TW                                    | SW                                   |
|-------------------|------------------------------------------|-------------------------------------|-------------------------------------|--------------------------------------|---------------------------------------|--------------------------------------|
| Control $\beta_0$ | <b>.63 (.98)</b><br>[0, 1.26]            | -0.29<br>[-0.82, 0.27]              | -0.27<br>[-0.86, 0.30]              | -0.27<br>[-0.83, 0.34]               | -0.30<br>[-0.92, 0.29]                | 0.22<br>[-0.34, 0.78]                |
| Sanders $\beta_1$ | .01<br>[-.68, .68]                       | <b>-0.53 (.05)</b><br>[-1.16, 0.10] | <b>-0.49 (.07)</b><br>[-1.17, 0.15] | <b>-0.67 (.02)</b><br>[-1.36, -0.01] | <b>-0.89 (.004)</b><br>[-1.54, -0.26] | <b>1.20 (1)</b><br>[ 0.57, 1.88]     |
| Trump $\beta_2$   | <b>-1.32 (&lt;.001)</b><br>[-2.00, -.64] | <b>0.81 (.99)</b><br>[ 0.20, 1.44]  | 0.22<br>[-0.44, 0.90]               | <b>1.04 (.99)</b><br>[ 0.40, 1.69]   | <b>1.14 (1)</b><br>[ 0.49, 1.80]      | <b>-0.73 (.02)</b><br>[-1.40, -0.06] |
| Arabic $\beta_3$  | <b>-.84 (0)</b><br>[-1.32, -.37]         | <b>-0.29 (.08)</b><br>[-0.72, 0.11] | <b>-0.32 (.07)</b><br>[-0.78, 0.13] | <b>-0.33 (.07)</b><br>[-0.78, 0.12]  | -0.28<br>[-0.77, 0.20]                | <b>0.43 (.98)</b><br>[ 0.01, 0.88]   |
| S:A $\beta_4$     | <b>.27 (.97)</b><br>[-.01, .58]          | 0.07<br>[-0.18, 0.32]               | 0.14<br>[-0.12, 0.40]               | 0.06<br>[-0.23, 0.35]                | 0.01<br>[-0.29, 0.29]                 | -0.07<br>[-0.35, 0.22]               |
| T:A $\beta_5$     | <b>.22 (.93)</b><br>[-.07, .53]          | 0.15<br>[-0.11, 0.41]               | <b>0.30 (.99)</b><br>[ 0.04, 0.57]  | 0.15<br>[-0.13, 0.44]                | <b>0.20 (.92)</b><br>[-0.09, 0.49]    | <b>-0.22 (.06)</b><br>[-0.51, 0.06]  |
| $\beta_0$         | 2.46                                     | (.04)                               | (.02)                               | (.73)                                | (.58)                                 | (.44)                                |
| $\beta_1$         | -.26                                     | 2.66                                | (.65)                               | (.77)                                | (.79)                                 | (.36)                                |
| $\beta_2$         | -.30                                     | .06                                 | 2.64                                | (.23)                                | (.15)                                 | (.54)                                |
| $\beta_3$         | .09                                      | .11                                 | -.10                                | 1.86                                 | (.06)                                 | (.008)                               |
| $\beta_4$         | .06                                      | .21                                 | -.24                                | -.41                                 | .60                                   | (.97)                                |
| $\beta_5$         | -.04                                     | -.09                                | .02                                 | -.60                                 | .61                                   | .65                                  |

Table 4.2: Parameter estimates for the six models predicting peaceful and aggressive responses. The baseline column reports the posterior estimates for the model containing no moderators. Other columns represent the posterior estimates for the effect of the corresponding variable on the row parameter. Brackets contain the 95% credible intervals. Parentheses indicate the posterior probability that the parameter is *positive* (ppp); 1 - ppp is the posterior probability that the parameter is negative. The bottom table provides the random effect standard deviation (diagonals), correlations (below diagonal), and ppp for these correlations (above diagonal).

| Parameter         | Baseline                          | AA                                   | SDO                                 | RWA-R                                | TW                                  | SW                                   |
|-------------------|-----------------------------------|--------------------------------------|-------------------------------------|--------------------------------------|-------------------------------------|--------------------------------------|
| Control $\beta_0$ | <b>.72 (1)</b><br>[.38, 1.06]     | -0.09<br>[-0.46, 0.28]               | -0.10<br>[-0.46, 0.25]              | 0.14<br>[-0.23, 0.50]                | -0.16<br>[-0.47, 0.15]              | 0.02<br>[-0.32, 0.35]                |
| Sanders $\beta_1$ | -.15<br>[-.62, .31]               | <b>-0.70 (.005)</b><br>[-1.23, -.18] | -0.28 (.12)<br>[-0.76, 0.22]        | <b>-0.76 (.001)</b><br>[-1.26, -.27] | <b>-0.36 (.05)</b><br>[-0.78, 0.09] | <b>0.88 (1)</b><br>[.47, 1.31]       |
| Trump $\beta_2$   | <b>-1.33 (0)</b><br>[-1.84, -.83] | <b>0.78 (1)</b><br>[.22, 1.38]       | <b>0.39 (.91)</b><br>[-0.16, 0.95]  | 0.30<br>[-0.27, 0.91]                | <b>0.64 (.99)</b><br>[.18, 1.10]    | <b>-0.67 (.007)</b><br>[-1.17, -.16] |
| Arabic $\beta_3$  | <b>-0.22 (.10)</b><br>[-.61, .13] | <b>-0.27 (.10)</b><br>[-0.69, 0.14]  | -0.05<br>[-0.44, 0.33]              | 0.03<br>[-0.41, 0.46]                | 0.11<br>[-0.25, 0.48]               | -0.11<br>[-0.48, 0.24]               |
| S:A $\beta_4$     | .06<br>[-.13, .27]                | 0.15 (.88)<br>[-0.10, 0.40]          | 0.08<br>[-0.15, 0.30]               | 0.02<br>[-0.23, 0.27]                | 0.04<br>[-0.17, 0.25]               | 0.12<br>[-0.10, 0.33]                |
| T:A $\beta_5$     | .07<br>[-.13, .26]                | -0.13<br>[-0.38, 0.12]               | <b>-0.15 (.08)</b><br>[-0.36, 0.06] | <b>-0.22 (.04)</b><br>[-0.46, 0.02]  | -0.10<br>[-0.30, 0.09]              | 0.09<br>[-0.11, 0.30]                |
| $\beta_0$         | 1.35                              | (.01)                                | (< .001)                            | (.76)                                | (.24)                               | (.71)                                |
| $\beta_1$         | -.32                              | 1.81                                 | (.08)                               | (.49)                                | (.64)                               | (.48)                                |
| $\beta_2$         | -.44                              | -.21                                 | 2.16                                | (.22)                                | (.43)                               | (.30)                                |
| $\beta_3$         | .10                               | -.004                                | -.11                                | 1.58                                 | (.32)                               | (.55)                                |
| $\beta_4$         | -.20                              | .11                                  | -.05                                | -.15                                 | .28                                 | (.64)                                |
| $\beta_5$         | .18                               | -.02                                 | -.18                                | .04                                  | .14                                 | .18                                  |

The first hypothesis was partially supported. When compared to Chinese symbols, Arabic symbols were rated as less pleasant (highly certain) and more aggressive (less certainly so).

The Arabic effect was seemingly positive moderated by the Sanders and Trump images, in the pleasant-unpleasant response condition. Just as in study 1, the negative Arabic effect was attenuated following a Sanders image. Unexpectedly, the effect was also attenuated following a Trump image. The second hypothesis (Negative interaction between Arabic and Trump trials) was not supported, but the third (Positive interaction between Arabic and Sanders trials) was supported in the pleasant-unpleasant response condition.

The fourth hypothesis (RWA and SDO negatively moderate the Arabic effect) was partially supported in the unpleasant-pleasant response condition, and not discerned in the aggressive-peaceful condition. Although not definitive, Arabic symbols were especially more likely to be rated as unpleasant if one is higher in anti-Arab prejudice, SDO, and RWA. The negative Arabic effect was partially negated in those high in Sanders support. Similarly non-definitive, Arabic symbols were rated as especially aggressive if one is higher in anti-Arab prejudice. No other predicted moderational effects were observed for the Arabic effect.

The fifth hypothesis was partially supported in both response conditions. On Trump trials, participants generally responded unpleasant and aggressive. This effect on unpleasant responses is negated as one increases in anti-Arab prejudice, RWA, and Trump support; and is strengthened for those high in Sanders support. Similarly, this effect on aggressive responses is negated as one increases in anti-Arab prejudice, SDO, and Trump support; and strengthens with Sanders support. That is, those who are high in anti-Arab prejudice, RWA, Trump support, and SDO; and who are low in Sanders support; are more likely to respond positively following a Trump image, regardless of symbol.

The sixth hypothesis was partially supported across both response conditions. On Sanders trials, participants were not generally more likely to respond positively or negatively,

but the effect of Sanders images is moderated in a manner nearly opposite to that of the Trump effect. Participants were more likely to rate the symbol as pleasant and peaceful following a Sanders image if the person is low in anti-Arab prejudice, SDO, RWA, and Trump support, and high in Sanders support. In essence, those who are low in anti-Arab prejudice, RWA, Trump support, and SDO; and who are high in Sanders support; are *more* likely to respond positively following a Sanders image.

### 4.3 Discussion

The data produced clearer inferences compared to study 1. Hypotheses one, two, three, four, five and six all gained at least partial support in the unpleasant-pleasant response condition. Specifically, the negative Arabic effect was observed with high certainty ( $\approx 1$ ), and was strengthened by anti-Arab prejudice, SDO, and RWA, albeit with less certainty ( $> .92$ ). Following a Trump exposure, those high in anti-Arab prejudice and RWA were more likely to indicate a positive response. Following a Sanders exposure, those low in anti-Arab prejudice, SDO, and RWA were more likely to indicate a positive response, and the Arabic effect was attenuated. Finally, unlike in study 1, the scales were not discernibly related to pleasant responses on their own. Nevertheless, study 3 examines whether these constructs are related to evaluation of known group symbols, as opposed to random unknown symbols.

The negative Arabic effect was unexpectedly attenuated following a Trump exposure. It is possible that exposure to both Sanders and Trump promote more equal evaluations of Arabic and Chinese symbols, albeit due to different ideological reasons — One emphasizing equality of groups, the other emphasizing outgroups are equally unpleasant. Another explanation is that a floor effect exists. Trump conditions predict a strongly negative response probability, response probabilities which are bounded by 0 and 1. If the predicted response probability is near the zero-boundary, then Arabic symbol effects may be too weak to further decrease the response probability; this consequently causes the difference between Arabic and Chinese symbols to decrease, and an interaction to emerge. A third

explanation is suggested by the third-order moderation by SDO and Trump support. For those high in SDO and Trump support, the negative Arabic effect is strengthened; however, the Trump-Arabic interaction also increases, which may suggest that for those who respond especially negatively following an Arabic symbol, a Trump exposure partially negates the negative effect of Arabic symbols.

The negative Arabic effect was less convincing in the aggressive-peaceful response condition ( $pp = .10$ ), and was only strengthened by anti-Arab prejudice ( $pp = .10$ ); other moderators' signs were indiscernible. The signs of the interactions between political icon and symbol were indiscernible. Nevertheless, peaceful responses were more probable following a Trump exposure for those high in anti-Arab prejudice and SDO, and following a Sanders exposure for those low in anti-Arab prejudice, RWA, and to a less certain degree, SDO.

## CHAPTER FIVE

### Study 3

Study three differs from the first two studies in two important factors. First, only the “Unpleasant vs Pleasant” response condition were utilized. Second, a third condition was added to the symbols factor.

A limitation of the first two studies was that the responses to Arabic symbols are compared to the responses to Chinese symbols. Consequently, any hypothesized differences between the two symbols could occur because the Arabic symbols are viewed as unpleasant (as hypothesized), or because the Chinese symbols are viewed as particularly pleasant. In order to rectify this, several hundred glyphs were randomly generated based on a modified version of Glyph Generator (<https://github.com/MadEqua/glyphgenerator>; modified with permission of author). From these glyphs, 36 were selected on the premise that they appear to be plausible linguistic symbols. Because these glyphs were randomly generated without any association to any language, culture, or group, they act as a control group against which the Chinese and Arabic symbols may be compared. Whether Arabic symbols are perceived to be unpleasant, or Chinese symbols as pleasant, can be determined by the use of these glyphs as a control group.

#### *5.1 Methods*

Two hundred and fifty participants were recruited from the Amazon Mechanical Turk crowdsourcing platform and paid \$.50 for their work. Given the power analysis from study 1 and the results from studies 1 and 2, 250 participants should be sufficient for a power of .8. Those who participated in study 2 were excluded from participation. Data collection occurred for nine days starting July 1, 2018.

All stimuli and procedures from studies 1 and 2 were employed for study 3 with minor modifications outlined here. An additional set of non-linguistic random symbols were used as control stimuli in addition to the Arabic and Chinese symbols. Using a free tool (Lourenço, 2015), hundreds of random pictographs were generated, from which 36 were selected based on desired complexity and uniqueness from known languages and cultures. Therefore, the stimuli included the twelve Trump images, twelve Sanders images, a control prime image, a fixation point, 36 Arabic images, 36 Chinese images, and 36 control target images, and finally the black and white reverse visual mask. Each participant completed 108 trials, across which the three prime and three symbol images are balanced. Instead of randomly assigning subjects to a response condition, all participants completed the “Unpleasant vs Pleasant” response condition. However, participants were assigned at random to the reverse ordering, “Pleasant vs Unpleasant”, in order to balance out the effects of hand dominance and response option order effects on the effects of interest.

In addition to the measures from studies 1 and 2, participants completed a series of thermometer items as a simple (reversed) measure of generalized prejudice (Cecil, Kate, & Sibley, 2017); the items are available in the Appendix. In order to assess generalized prejudice without overlapping with the groups and stimuli employed, only items 4, 5, 6, 7, 10, and 11 were used to inform the latent variable ( $\alpha = .92$ ). The attitudes toward political candidates scale was revised due to the completion of the primary and presidential elections. Specifically, the two items on the questionnaire about supporting the respective candidates in the upcoming primaries were removed. The mean time spent to complete the study was approximately 15 minutes.

## 5.2 Results

Two hundred and ninety six individuals participated. After removing those who failed any attention check, 242 remained. Of those, 215 remained after removing those who

indicated experience with Chinese or Arabic languages. The final sample consisted of 215 subjects ( $m_{\text{age}} = 39.79$ ,  $\sigma_{\text{age}} = 12.49$ , 43.26% male).

A Bayesian logistic GLMM was employed for modeling the effects of interest, similar to studies 1 and 2. The symbol conditions were Helmert-coded, in which Arabic and Chinese symbols were both compared to Control, and Arabic was compared to Chinese. All analyses conducted in studies 1 and 2 were performed in study 3, with the additional Helmert-coded contrasts as both a main effect and an effect to be moderated by SDO, RWA, and the political icons. All data preparation, transformations, and analyses conducted in studies 1 and 2 were performed in study 3. Only one response from all questionnaires was missing, and no experimental trial data were missing.

All model priors were the same as before. Four chains were used with 1000 post-adaptation iterations. No divergences were reported, and  $\hat{R}$  for all parameters  $\approx 1$ . All parameters of interest are reported in 5.1 and 5.2. For clarity, Figure 5.1 presents the population-predicted probability of responding “pleasant” for each experimental condition, across each moderator.

The first hypothesis was supported. Responses following an Arabic symbol were likely unpleasant. The hypothesized negative effect of Trump images on the Arabic effect was only partially supported; the probability that the moderation coefficient is negative is not satisfactorily certain (.91). Moreover, the most probable interactive effect sizes are small. The predicted positive effect of Sanders images on the Arabic effect was not supported. The sign is not discernible, but even the most probable effects are too small to consider supportive of the hypothesis.

Hypothesis four was not supported. The negative Arabic effect was only notably moderated by explicit anti-Arab prejudice, and to a lesser degree, Sanders support and generalized prejudice. The signs of the moderating coefficients were not discernible, and the most probable effects were fairly small.

Table 5.1: Parameter estimates for all models. Expected a-posteriori values provided, with [95% credible intervals] and (ppp). AA = anti-Arab prejudice, SDO = social dominance orientation, RWA-R = right-wing authoritarianism, TW = Trump warmth, SW = Sanders Warmth, Therm = Warmth toward other groups. First column represents the base model, with no moderators. Subsequent columns represent the moderating effect of the column construct on the corresponding row's effect.

| Parameter | Baseline                          | AA                                    | SDO                                 | RWA-R                                 | TW                                    | SW                                    | Therm                              |
|-----------|-----------------------------------|---------------------------------------|-------------------------------------|---------------------------------------|---------------------------------------|---------------------------------------|------------------------------------|
| Control   | <b>.42 (1)</b><br>[.20, .62]      | <b>-0.45 (0)</b><br>[-0.67, -0.24]    | <b>-0.46 (0)</b><br>[-0.70, -0.23]  | <b>-0.37 (.001)</b><br>[-0.60, -0.15] | <b>-0.28 (.005)</b><br>[-0.52, -0.07] | <b>0.20 (.96)</b><br>[-0.02, 0.42]    | <b>0.54 (1)</b><br>[0.32, 0.76]    |
| Sanders   | -.04<br>[-.22, .14]               | <b>-0.17 (.03)</b><br>[-0.37, 0.01]   | -0.08<br>[-0.27, 0.12]              | <b>-0.14 (.09)</b><br>[-0.34, 0.06]   | <b>-0.14 (.08)</b><br>[-0.33, 0.05]   | <b>0.45 (1)</b><br>[0.27, 0.64]       | 0.02<br>[-0.17, 0.21]              |
| Trump     | <b>-49 (.002)</b><br>[-.83, -.16] | <b>1.02 (1)</b><br>[0.70, 1.38]       | <b>0.95 (1)</b><br>[0.62, 1.28]     | <b>0.90 (1)</b><br>[0.55, 1.26]       | <b>1.18 (1)</b><br>[0.87, 1.53]       | <b>-0.77 (0)</b><br>[-1.12, -0.42]    | <b>-0.74 (0)</b><br>[-1.10, -0.39] |
| Arabic    | <b>-35 (0)</b><br>[-.52, -.18]    | <b>-0.24 (.002)</b><br>[-0.42, -0.06] | -0.07<br>[-0.25, 0.12]              | -0.02<br>[-0.21, 0.17]                | 0.08<br>[-0.10, 0.26]                 | <b>0.17 (.96)</b><br>[-0.02, 0.35]    | <b>0.19 (.97)</b><br>[0.00, 0.38]  |
| Group     | <b>.21 (1)</b><br>[.09, .34]      | <b>-0.12 (.04)</b><br>[-0.25, 0.01]   | <b>-0.12 (.05)</b><br>[-0.25, 0.02] | <b>-0.10 (.09)</b><br>[-0.24, 0.05]   | -0.04<br>[-0.17, 0.10]                | -0.02<br>[-0.16, 0.12]                | <b>0.11 (.94)</b><br>[-0.03, 0.25] |
| S:A       | 0<br>[-.11, .11]                  | 0.04<br>[-0.07, 0.15]                 | 0.02<br>[-0.10, 0.13]               | -0.04<br>[-0.16, 0.07]                | -0.06<br>[-0.17, 0.06]                | -0.06<br>[-0.19, 0.06]                | -0.03<br>[-0.14, 0.08]             |
| T:A       | <b>-08 (.09)</b><br>[-.19, .03]   | 0.00<br>[-0.12, 0.12]                 | 0.00<br>[-0.11, 0.12]               | 0.00<br>[-0.12, 0.13]                 | -0.02<br>[-0.14, 0.10]                | <b>-0.15 (.009)</b><br>[-0.28, -0.02] | -0.02<br>[-0.14, 0.09]             |
| S:G       | <b>-05 (.05)</b><br>[-.11, .01]   | 0.03<br>[-0.03, 0.10]                 | 0.02<br>[-0.04, 0.09]               | 0.02<br>[-0.05, 0.08]                 | 0.03<br>[-0.03, 0.09]                 | 0.00<br>[-0.06, 0.07]                 | 0.00<br>[-0.07, 0.06]              |
| T:G       | <b>-05 (.08)</b><br>[-.11, .02]   | 0.01<br>[-0.06, 0.08]                 | 0.00<br>[-0.06, 0.07]               | 0.01<br>[-0.07, 0.08]                 | 0.00<br>[-0.07, 0.07]                 | -0.04<br>[-0.11, 0.03]                | -0.02<br>[-0.09, 0.05]             |

Table 5.2: Correlation matrix of random effects (lower triangle), standard deviation of random effects (diagonal), and ppp (above diagonal) for study 3.

|         |      |       |       |       |       |       |       |       |       |
|---------|------|-------|-------|-------|-------|-------|-------|-------|-------|
| Control | 1.50 | (.06) | (0)   | (.97) | (.99) | (.48) | (.33) | (.48) | (.48) |
| Sanders | -.14 | 1.16  | (.72) | (.60) | (.51) | (.34) | (.40) | (.54) | (.34) |
| Trump   | -.50 | .05   | 2.36  | (.27) | (.40) | (.51) | (.63) | (.47) | (.69) |
| Arabic  | .15  | .02   | -.05  | 1.14  | (1)   | (.45) | (.48) | (.59) | (.30) |
| Group   | .19  | 0.00  | -.02  | .26   | .91   | (.69) | (.46) | (.55) | (.27) |
| S:A     | -.02 | -.10  | .01   | -.03  | .13   | .10   | (.54) | (.54) | (.43) |
| T:A     | -.12 | -.06  | .09   | -.01  | -.02  | .04   | .09   | (.50) | (.57) |
| S:G     | -.01 | .03   | -.01  | .05   | .03   | .03   | 0.00  | .05   | (.51) |
| T:G     | -.01 | -.09  | .11   | -.11  | -.14  | -.04  | .04   | 0.00  | .09   |

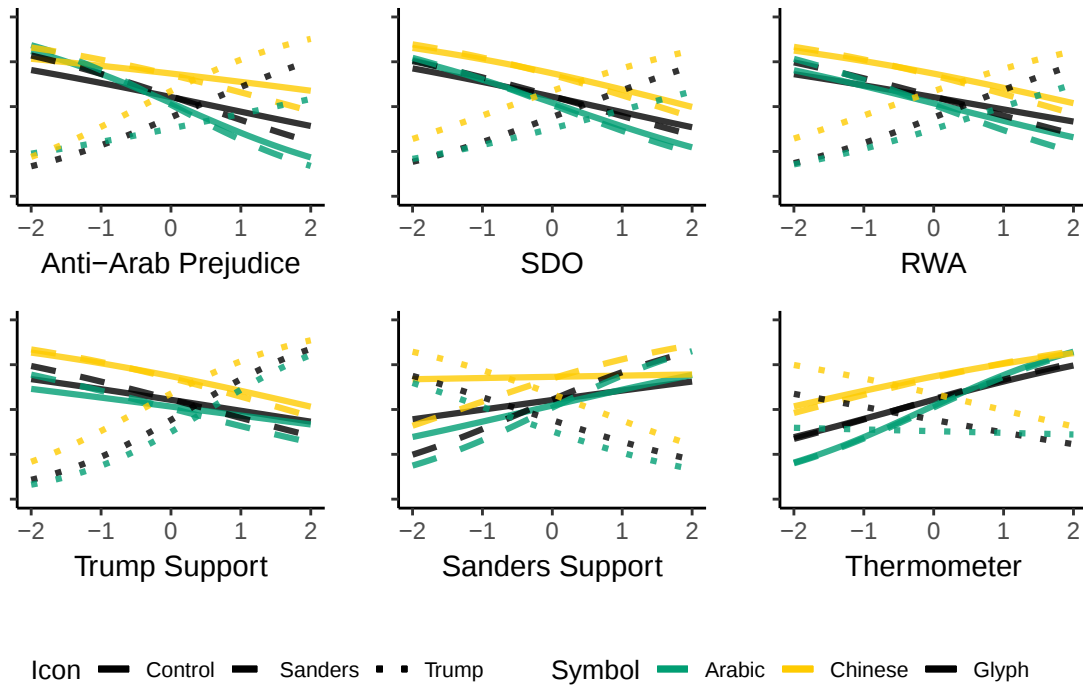


Figure 5.1: Fixed effect predictions for each condition, across each latent moderator. The y-axis is the probability, ranging from 0 to 1, of responding "pleasant".

Hypothesis five was well supported. Although responses tended to be unpleasant following a Trump exposure, anti-Arab prejudice, SDO, RWA, Trump support, and generalized prejudice all predicted more pleasant responses. Sanders support predicted unpleasant responses following a Trump exposure.

Hypothesis six was only partially supported. Responses following a Sanders prime were not expected to be either unpleasant or pleasant. Those high in anti-Arab prejudice, and to a lesser degree, RWA and Trump support, were less likely to indicate pleasant following a Sanders image. Sanders support predicted pleasant responses following a Sanders image.

Hypotheses seven and eight, that SDO, RWA, and generalized prejudice would predict positive judgments of random symbols compared to both Arabic and Chinese symbols, was partially supported. On average, the group symbols (Chinese and Arabic) were rated as pleasant more often than random symbols. However, this effect is weakened for those high in anti-Arab prejudice, SDO, and to a less certain degree, RWA, and generalized prejudice. Interestingly, those high in anti-Arab prejudice, SDO, RWA, Trump support, generalized prejudice, and low in Sanders support, were again more likely to respond unpleasant across all symbols.

### 5.2.1 Additional Models

In addition to the models above, two others were fit. The first model examines whether a fatigue effect is present, such that effects weaken or strengthen as the subject completes more trials. The second model is similar to the above models, except all latent moderators are simultaneously entered, in order to estimate partial effects of each factor.

In order to assess the presence of a fatigue effect, the randomized trial order data were saved for each participant and entered into the baseline model as a fixed predictor and interaction. Table 5.3 presents the estimated effects of order on each coefficient, *assuming the trial is the final trial* (i.e.,  $108 \times \beta$ ). This rescaling was performed to present the estimates in a more intuitive manner, as a worst case scenario. For the most part, the signs of order effects

Table 5.3: Parameter estimates for the baseline model with trial order interacting with each predictor. Because the coefficients for Trial Order are  $< .002$ , these coefficients are multiplied by 108. Therefore, the Trial Order column indicates how much the respective row's effect would change by the final trial.

| Parameter | Baseline                | Trial Order $\times 108$                   |
|-----------|-------------------------|--------------------------------------------|
| Control   | 0.45<br>[ 0.22, 0.67]   | -0.02<br>[-0.22, 0.19]                     |
| Sanders   | 0.07<br>[-0.17, 0.32]   | <b>-0.24 (.05)</b><br><b>[-0.53, 0.05]</b> |
| Trump     | -0.47<br>[-0.82, -0.12] | -0.06<br>[-0.36, 0.23]                     |
| Arabic    | -0.31<br>[-0.53, -0.10] | -0.06<br>[-0.31, 0.21]                     |
| Group     | 0.27<br>[ 0.12, 0.42]   | <b>-0.10 (.08)</b><br><b>[-0.25, 0.05]</b> |
| S:A       | 0.04<br>[-0.19, 0.25]   | -0.07<br>[-0.44, 0.30]                     |
| T:A       | -0.02<br>[-0.24, 0.20]  | -0.13<br>[-0.50, 0.25]                     |
| S:G       | -0.11<br>[-0.23, 0.01]  | 0.12<br>[-0.10, 0.32]                      |
| T:G       | -0.14<br>[-0.26, -0.01] | 0.18<br>[-0.03, 0.38]                      |

on the baseline model parameters are uncertain. The effect of seeing Sanders appears to decrease as trials continue, although the effect of Sanders itself is uncertain. Additionally, the effect of seeing group symbols (as opposed to random glyphs) appears to decrease as trials continue, but the interaction effect is too small to negate the effect.

The second model included all latent factors as simultaneous, correlated predictors and moderators. For computational efficiency, latent score distributions were estimated separately from the predictive model. Specifically, a measurement-model was fit for all subjects, yielding posterior distributions of each subjects' respective factor scores. These posteriors were then entered as priors in the predictive model. Importantly, *point estimates* were not extracted from the measurement model, but entire distributions. This permits uncertainty in factor scores to be considered and preserved. Algebraically, this approach is

Table 5.4: Correlations (lower diagonal) of latent factors with all other latent factors (with posterior probability that the correlation is positive) on the upper diagonal. Several 1s and 0s are present; this is because to .0001 precision, the value would be rounded to 0 or 1.

| Factor | AA   | SDO  | RWA  | TW   | SW  | Therm  |
|--------|------|------|------|------|-----|--------|
| AA     | 1    | (1)  | (1)  | (1)  | (0) | (0)    |
| SDO    | .71  | 1    | (1)  | (1)  | (0) | (0)    |
| RWA    | .65  | .69  | 1    | (1)  | (0) | (0)    |
| TW     | .64  | .66  | .69  | 1    | (0) | (.001) |
| SW     | -.31 | -.25 | -.35 | -.40 | 1   | (1)    |
| Therm  | -.45 | -.36 | -.20 | -.20 | .30 | 1      |

equivalent to simultaneously estimating factor scores and the predictive model, but is far more computationally efficient.

The latent factor correlations are presented in Table 5.4, and the simultaneous model coefficients are presented in Table 5.5. The factors are all strongly and certainly correlated as expected. Trump support was strongly positively correlated with anti-Arab prejudice, SDO, RWA, and, to a lesser degree, generalized prejudice. Sanders support was negatively correlated with anti-Arab prejudice, SDO, RWA, and generalized prejudice. Due to high correlations among these factors producing high posterior covariance of the predictors, most of the previously observed moderational effects were highly uncertain when all factors were simultaneously entered.

Only two effects emerged wherein sign was discernible. Those high in RWA were likely to respond unpleasant, after controlling for all other factors, regardless of symbol type. Likewise, those high in SDO were likely to respond unpleasant to group symbols compared to random glyphs.

### 5.3 Discussion

The negative Arabic effect was once again observed. However, this effect was not moderated by SDO nor RWA, but was moderated by Sanders support and generalized

Table 5.5: Coefficients for model including all predictors and moderators simultaneously.

| Parameter | AA                     | SDO                                 | RWA                                 | TW                     | SW                     | Therm                  | $\frac{\sigma_b^2 - \sigma_m^2}{\sigma_b^2}$ |
|-----------|------------------------|-------------------------------------|-------------------------------------|------------------------|------------------------|------------------------|----------------------------------------------|
| Control   | -0.19<br>[-1.04, 0.72] | -0.14<br>[-1.03, 0.72]              | <b>-0.75 (.03)</b><br>[-1.55, 0.05] | 0.46<br>[-0.44, 1.40]  | -0.14<br>[-0.61, 0.31] | 0.57<br>[ 0.11, 1.06]  | .23                                          |
| Sanders   | -0.36<br>[-1.18, 0.48] | 0.08<br>[-0.77, 0.88]               | -0.15<br>[-0.90, 0.60]              | 0.32<br>[-0.56, 1.17]  | 0.53<br>[ 0.11, 0.97]  | -0.16<br>[-0.58, 0.28] | .07                                          |
| Trump     | 0.77<br>[-0.37, 1.98]  | -0.28<br>[-1.42, 0.92]              | 0.06<br>[-1.04, 1.14]               | 0.84<br>[-0.41, 2.05]  | -0.09<br>[-0.78, 0.60] | -0.45<br>[-1.17, 0.26] | .17                                          |
| Arabic    | -0.26<br>[-0.85, 0.35] | -0.18<br>[-0.82, 0.40]              | 0.37<br>[-0.21, 0.93]               | 0.10<br>[-0.58, 0.78]  | 0.09<br>[-0.19, 0.39]  | 0.01<br>[-0.29, 0.33]  | .62                                          |
| Group     | 0.14<br>[-0.44, 0.73]  | <b>-0.47 (.06)</b><br>[-1.12, 0.13] | -0.24<br>[-0.78, 0.31]              | 0.40<br>[-0.25, 1.04]  | -0.04<br>[-0.33, 0.25] | -0.02<br>[-0.32, 0.28] | .20                                          |
| S:A       | 0.30<br>[-0.24, 0.85]  | -0.15<br>[-0.71, 0.39]              | 0.10<br>[-0.44, 0.63]               | -0.33<br>[-0.94, 0.27] | 0.15<br>[-0.11, 0.42]  | -0.18<br>[-0.44, 0.07] | 0                                            |
| T:A       | 0.15<br>[-0.42, 0.70]  | -0.28<br>[-0.86, 0.28]              | -0.05<br>[-0.61, 0.51]              | 0.18<br>[-0.46, 0.83]  | 0.09<br>[-0.19, 0.38]  | -0.22<br>[-0.49, 0.05] | 0                                            |
| S:G       | -0.08<br>[-0.40, 0.24] | 0.15<br>[-0.21, 0.50]               | 0.02<br>[-0.30, 0.35]               | -0.12<br>[-0.47, 0.26] | 0.09<br>[-0.08, 0.26]  | -0.02<br>[-0.19, 0.13] | 0                                            |
| T:G       | -0.22<br>[-0.61, 0.17] | 0.15<br>[-0.28, 0.55]               | -0.16<br>[-0.55, 0.22]              | 0.13<br>[-0.30, 0.58]  | -0.07<br>[-0.27, 0.12] | 0.04<br>[-0.14, 0.23]  | 0                                            |

prejudice. Furthermore, the hypothesized effect of Trump exposure on the negative Arabic effect was uncertain, and small. The moderation by Sanders exposures was unsupported.

Responses following Trump images were positively moderated by anti-Arab prejudice, SDO, RWA, Trump support, and generalized prejudice. This once again suggests that those high in these factors tend to respond favorably following an image of Trump. Conversely, responses following Sanders images were negatively moderated by anti-Arab prejudice and to a lesser certainty, RWA, and Trump support.

Those higher in anti-Arab prejudice, SDO, generalized prejudice, and, to a lesser certainty, RWA were more likely to respond pleasant to random glyphs than to Arabic or Chinese symbols, compared to those lower. According to Figure 5.1, this effect was not precisely as predicted. The predicted effect was that those high in SDO and RWA would respond more positively to random glyphs than to group symbols, on average. In contrast, the figure suggests that those high in SDO and RWA do respond less positively to group

symbols, on average, but this effect seems predominantly due to responding unpleasant to Arabic symbols in particular.

It is puzzling that the moderational effect of SDO on the Arabic effect did not emerge in study 3, given the relatively consistent effect in studies 1 and 2. It is possible, although unverifiable with the present data, that any normalization of anti-Arab sentiment and attitudes that occurred since 2016 attenuates the effect of SDO on the Arabic effect. That is, those low or average in SDO in previous years may have inhibited negative responses to Arabic symbols due to social pressures. However, if anti-Arab attitudes are normalized and anti-Arab information is highly accessible, then the covariance between SDO and ease of negative responses to Arabic symbols may be attenuated. Of course, the failure to replicate prior studies' moderational effects may simply be due sampling variation, as prior information about these parameters was not taken into account. The next chapter summarizes an integrative analysis in order to synthesize results across studies.

## CHAPTER SIX

### Integrative Data Analysis

An integrative data analysis was planned and attempted. The planned model hierarchically models all shared components and parameters across studies. This approach is akin to a random effects meta-analysis on the estimated linear components, with the addition of random effects on each study's random effect variances and correlation matrices.

Multiple versions of this approach were attempted. The full model included hierarchical correlation matrices using the convex combination property of correlation matrices to pool individual study matrices toward a global correlation matrix. This model additionally included all latent scores (with corresponding uncertainties and imputed values) as predictors in separate submodels. Studies were assumed to have random effects on the linear coefficients, the random effect correlations, and the random effect variances. Unfortunately, this model was unable to converge after several days of sampling. The model experienced several hundred divergences, and the chains failed to converge to a single typical set. This implies that the data are not informative enough to produce a unimodal posterior solution. Although stronger priors may have helped identify the model, the only objective prior to use would be from the studies themselves, which would make double use of the data, and overestimate the certainty of the respective parameters.

Instead, a simpler model was attempted in which latent scores were treated as known, but this too was unable to converge. A further simplified model was attempted in which studies had random effects on only the linear components; this is essentially a full-information random effects model with random effects of study on the linear coefficients, and is similar to a random effects meta-analysis on raw data. Unfortunately, these models similarly failed to converge.

Finally, a simpler integrative technique — Random effects meta-analysis — using only sufficient statistics (e.g., posterior means and standard deviations from each previous study) was attempted. However, with only three studies, this method failed to converge.

Nevertheless, a fixed-effects meta-analysis was estimated in order to synthesize across studies. Fixed-effects meta-analyses assume that uncertainty in a parameter value is solely due to sampling variability, and all estimates are manifestations of a respective singular fixed value. A meta-analysis was performed for the base model and all moderational models from studies 1–3 that employed the “Pleasant/Unpleasant” response options, excluding the coefficients unique to study 3 (e.g., generalized prejudice, random glyph effects). A similar analysis was performed for the base model and all moderational models from studies 1–2 that employed the “Aggressive/Peaceful” response option.

Tables 6.1 and 6.2 includes the fixed-effects meta-analytic estimates. Arabic symbols are less likely to produce a pleasant response, and this effect is moderated by explicit anti-Arab prejudice and SDO; support for Sanders weakens the effect.

Following a Trump image, responses were more likely to be unpleasant. However, this effect is reversed for those high in anti-Arab prejudice, SDO, RWA, and Trump support; and is strengthened for those high in Sanders support. Conversely, responses following a Sanders image were more likely to be pleasant for those low in anti-Arab prejudice, RWA, Trump support, and to a less certain degree, SDO.

The interactions between political icon and symbol were essentially zero. The only discernible third-order interaction is that of Sanders support on the Trump-Arabic interaction, such that for those high in Sanders support, the effect of seeing an Arabic symbol is more negative following a Trump image.

Anti-Arab prejudice, SDO, RWA, and Trump support all predict unpleasant responses across symbol conditions. Conversely, Sanders support predicts pleasant responses.

The second response option meta-analysis yielded nothing certain about the predicted Arabic effect or interactions with it. Similar to the previous meta-analytic model, “peaceful”

Table 6.1: Fixed effects meta-analysis on shared linear components across studies 1–3 [with 95% credible intervals] and (posterior probability of positivity, where notable) for the pleasant vs unpleasant response option data.

| Parameter | Baseline                           | AA                                    | SDO                                  | RWA                                 | TW                                    | SW                                    |
|-----------|------------------------------------|---------------------------------------|--------------------------------------|-------------------------------------|---------------------------------------|---------------------------------------|
| Control   | <b>0.46 (1)</b><br>[ 0.27, 0.65]   | <b>-0.43 (0)</b><br>[-0.63, -0.24]    | <b>-0.44 (0)</b><br>[-0.65, -0.23]   | <b>-0.38 (0)</b><br>[-0.58, -0.18]  | <b>-0.30 (.001)</b><br>[-0.49, -0.10] | <b>0.21 (.98)</b><br>[ 0.02, 0.41]    |
| Sanders   | -0.07<br>[-0.23, 0.10]             | <b>-0.18 (.03)</b><br>[-0.36, 0.00]   | <b>-0.12 (.10)</b><br>[-0.29, 0.06]  | <b>-0.17 (.04)</b><br>[-0.35, 0.02] | <b>-0.19 (.02)</b><br>[-0.37, -0.01]  | <b>0.50 (1)</b><br>[ 0.33, 0.67]      |
| Trump     | <b>-0.69 (0)</b><br>[-0.96, -0.40] | <b>1.00 (1)</b><br>[ 0.73, 1.28]      | <b>0.81 (1)</b><br>[ 0.51, 1.09]     | <b>0.97 (1)</b><br>[ 0.66, 1.26]    | <b>1.14 (1)</b><br>[ 0.87, 1.42]      | <b>-0.75 (0)</b><br>[-1.05, -0.46]    |
| Arabic    | <b>-0.40 (0)</b><br>[-0.55, -0.24] | <b>-0.27 (.001)</b><br>[-0.43, -0.12] | <b>-0.17 (.02)</b><br>[-0.33, -0.01] | -0.08<br>[-0.24, 0.08]              | -0.05<br>[-0.20, 0.10]                | <b>0.24 (1)</b><br>[ 0.08, 0.39]      |
| S:A       | 0.05<br>[-0.05, 0.15]              | 0.04<br>[-0.06, 0.14]                 | 0.04<br>[-0.06, 0.14]                | -0.04<br>[-0.15, 0.06]              | -0.04<br>[-0.13, 0.06]                | -0.05<br>[-0.16, 0.05]                |
| T:A       | -0.02<br>[-0.12, 0.08]             | 0.03<br>[-0.08, 0.14]                 | 0.06<br>[-0.05, 0.16]                | 0.04<br>[-0.07, 0.15]               | 0.01<br>[-0.10, 0.11]                 | <b>-0.16 (.003)</b><br>[-0.27, -0.04] |

Table 6.2: Fixed effects meta-analysis on shared linear components across studies 1–2 [with 95% credible intervals] and (ppp), for peaceful vs aggressive response option data.

| Parameter | Baseline                           | AA                                   | SDO                                 | RWA                                  | TW                                  | SW                                   |
|-----------|------------------------------------|--------------------------------------|-------------------------------------|--------------------------------------|-------------------------------------|--------------------------------------|
| Control   | <b>0.78 (1)</b><br>[ 0.48, 1.08]   | -0.14<br>[-0.49, 0.21]               | <b>-0.21 (.10)</b><br>[-0.53, 0.10] | 0.05<br>[-0.29, 0.39]                | -0.18<br>[-0.47, 0.13]              | 0.09<br>[-0.21, 0.39]                |
| Sanders   | -0.06<br>[-0.38, 0.27]             | <b>-0.44 (.02)</b><br>[-0.86, -0.04] | -0.20<br>[-0.59, 0.18]              | <b>-0.41 (.02)</b><br>[-0.80, -0.03] | <b>-0.22 (.10)</b><br>[-0.58, 0.12] | <b>0.54 (1)</b><br>[ 0.23, 0.84]     |
| Trump     | <b>-1.16 (0)</b><br>[-1.58, -0.74] | <b>0.69 (.99)</b><br>[ 0.18, 1.19]   | <b>0.34 (.92)</b><br>[-0.13, 0.83]  | <b>0.37 (.93)</b><br>[-0.12, 0.86]   | <b>0.68 (1)</b><br>[ 0.27, 1.08]    | <b>-0.55 (.01)</b><br>[-0.96, -0.12] |
| Arabic    | -0.09<br>[-0.39, 0.22]             | -0.13<br>[-0.50, 0.24]               | 0.05<br>[-0.28, 0.38]               | 0.08<br>[-0.27, 0.43]                | 0.07<br>[-0.24, 0.38]               | -0.13<br>[-0.43, 0.16]               |
| S:A       | 0.02<br>[-0.16, 0.19]              | 0.09<br>[-0.14, 0.32]                | 0.07<br>[-0.12, 0.28]               | -0.03<br>[-0.24, 0.19]               | 0.04<br>[-0.15, 0.23]               | 0.10<br>[-0.09, 0.29]                |
| T:A       | 0.05<br>[-0.12, 0.22]              | -0.06<br>[-0.29, 0.15]               | -0.05<br>[-0.24, 0.13]              | -0.11<br>[-0.32, 0.09]               | -0.06<br>[-0.23, 0.12]              | 0.04<br>[-0.14, 0.22]                |

responses were more likely following a Trump exposure for those high in anti-Arab prejudice, Trump support, and to a lesser certainty, SDO and RWA. Conversely, “aggressive” responses were more likely following a Sanders exposure for those high in anti-Arab prejudice, RWA, and to a lesser certainty, Trump support. Sanders support predicted “peaceful” and “aggressive” responses following a Sanders and Trump exposure, respectively. Unlike the previous meta-analytic model, only SDO predicts negative responses in the political icon control trials, across symbol types.

### 6.1 *Reanalysis of RWA Models*

The RWA-R scale is composed of two facets: Authoritarian aggression/submission (AAS) and Conservatism (C). In order to assess whether one facet or another is responsible for any effects of RWA-R in studies 1–3, all pleasant response option data from each study was reanalyzed using the AAS and C facets of the scale. Table 6.3 presents the parameter estimates for each study and facet. The two facets were highly correlated across studies ( $r_s = .49, .59, .50$ ; posterior  $\sigma_{r_s} = .13, .06, .05$ ).

The model for study 1 is difficult to interpret due to convergence concerns. With so few individuals in study 1, and fewer items for each factor, the uncertainty in the AAS facet scores resulted in a problematic fit (divergent transitions were detected,  $\hat{R} < 1.11$ ). Due to this, studies 2 and 3 are instead discussed.

In study 2, both AAS and C appear to strongly and consistently moderate the Sanders and Trump effects in a manner similar to the general factor model. However, the negative Arabic effect observed in study 2 appears to be influenced by the C facet, and not the AAS facet. The C facet appeared to have a positive third-order interaction with the Trump-Arabic interaction, such that those high in the C facet are more likely to respond positively following an Arabic image, if it followed a Trump image. This effect was not observed in study 2 using the general factor model.

Table 6.3: Re-analysis of the moderating effect of RWA split by facets Authoritarian aggression/submission (AAS) and Conservatism (C).

| Parameter | Authoritarian aggression & submission |                                     |                                      | Conservatism                         |                                      |                                     |
|-----------|---------------------------------------|-------------------------------------|--------------------------------------|--------------------------------------|--------------------------------------|-------------------------------------|
|           | Study 1                               | Study 2                             | Study 3                              | Study 1                              | Study 2                              | Study 3                             |
| Control   | 0.09<br>[-1.85, 4.69]                 | -0.27<br>[-0.93, 0.38]              | <b>-0.36 (0)</b><br>[-0.59, -0.15]   | <b>-0.99 (.032)</b><br>[-2.59, 0.03] | -0.21<br>[-0.83, 0.47]               | <b>-0.44 (0)</b><br>[-0.69, -0.19]  |
| Sanders   | 0.14<br>[-2.48, 8.82]                 | <b>-0.71 (.05)</b><br>[-1.55, 0.10] | -0.12<br>[-0.32, 0.08]               | -0.04<br>[-1.47, 1.27]               | <b>-0.81 (.02)</b><br>[-1.65, -0.06] | <b>-0.21 (.03)</b><br>[-0.41, 0.00] |
| Trump     | <b>2.85 (.93)</b><br>[-0.80, 17.71]   | <b>0.98 (.99)</b><br>[ 0.21, 1.89]  | <b>0.90 (1)</b><br>[ 0.51, 1.25]     | <b>1.75 (.99)</b><br>[ 0.25, 3.52]   | <b>1.25 (1)</b><br>[ 0.51, 2.00]     | <b>0.83 (1)</b><br>[ 0.45, 1.22]    |
| Arabic    | -0.18<br>[-1.18, 1.06]                | -0.24<br>[-0.81, 0.39]              | 0.01<br>[-0.18, 0.21]                | -0.21<br>[-0.94, 0.54]               | <b>-0.43 (.08)</b><br>[-1.05, 0.15]  | -0.10<br>[-0.31, 0.11]              |
| S:A       | <b>-0.44 (.09)</b><br>[-1.42, 0.20]   | -0.04<br>[-0.38, 0.30]              | -0.05<br>[-0.17, 0.06]               | -0.31<br>[-0.84, 0.23]               | 0.11<br>[-0.23, 0.42]                | -0.01<br>[-0.15, 0.12]              |
| T:A       | 0.12<br>[-1.10, 1.11]                 | 0.07<br>[-0.27, 0.41]               | -0.02<br>[-0.15, 0.10]               | 0.23<br>[-0.41, 0.87]                | <b>0.26 (.95)</b><br>[-0.09, 0.59]   | <b>0.09 (.90)</b><br>[-0.04, 0.22]  |
| Group     |                                       |                                     | <b>-0.12 (.058)</b><br>[-0.26, 0.02] |                                      |                                      | -0.02<br>[-0.19, 0.15]              |
| S:G       |                                       |                                     | 0.02<br>[-0.06, 0.09]                |                                      |                                      | 0.01<br>[-0.06, 0.08]               |
| T:G       |                                       |                                     | 0.01<br>[-0.07, 0.09]                |                                      |                                      | 0.00<br>[-0.08, 0.09]               |

In study 3, both AAS and C appear to predict unpleasant responses as a main effect, and exhibit an interaction with the Trump images, similar to the general factor model. The C facet is responsible for the negative interactive effect with Sanders images, and an uncertain positive third-order interaction with the Trump-Arabic interaction. That is, those higher in the C facet are expected to respond more positively after an Arabic symbol, if a Trump exposure was present. This third-order interaction, albeit uncertain, was not observed when using the general RWA factor in study 3. The AAS facet is responsible for the (albeit, less certain) interaction with group symbols, such that those high in AAS are less prone to respond pleasant when a symbol is Arabic or Chinese, compared to random glyphs.

Across studies, both AAS and C facets strongly predict pleasant responses following Trump images. A novel effect emerged consistent, that those high in the C facet are more likely to respond positively after an Arabic symbol if it follows a Trump image.

Table 6.4: Fixed effects meta-analytic estimates of the two RWA-R facets' effects.

| Parameter | AAS                              | C                                 |
|-----------|----------------------------------|-----------------------------------|
| Control   | <b>-.34 (0)</b><br>[-.56, -.14]  | <b>-.43 (0)</b><br>[-.66, -.20]   |
| Sanders   | <b>-.15 (.06)</b><br>[-.34, .04] | <b>-.24 (.01)</b><br>[-.43, -.04] |
| Trump     | <b>.92 (1)</b><br>[.58, 1.27]    | <b>.95 (1)</b><br>[.62, 1.29]     |
| Arabic    | -.02<br>[-.21, .15]              | <b>-.14 (.07)</b><br>[-.33, .05]  |
| S:A       | -.06<br>[-.17, .05]              | -.01<br>[-.13, .11]               |
| T:A       | -.01<br>[-.13, .11]              | <b>.12 (.97)</b><br>[0, .24]      |

A fixed effects meta-analysis across the studies' pleasant-unpleasant data for the separate RWA facets is presented in Table 6.4. Both facets predict unpleasant responses, especially following a Sanders exposure, and is negated following a Trump exposure. Although uncertain, the conservatism facet predicts stronger Arabic effects; this effect is attenuated if the Arabic symbol follows a Trump exposure.

It is worth noting that different effects of AAS and C facets are conflated by a method effect. The RWA-R's C facet consists of all negatively worded items, whereas the AAS items are all not. Therefore, it is difficult to conclude whether conservatism or negatively worded items are responsible for the differential effects of AAs and C scales.

## CHAPTER SEVEN

### General Discussion

A pilot and three experimental studies were performed. The pilot study suggested stimuli categories were perceived as intended, and no particular stimulus demonstrated unexpected categorizations. Unfortunately, accuracy in the “implicit” task was unexpectedly high, suggesting that subjects were fully aware of all stimuli. This raises inferential concerns for studies 1–3. Notably, any detected effects may not be reflecting *implicit* evaluations, but explicit evaluations. However, the detected Arab effects differed strongly between the two response option conditions in studies 1 and 2. Whereas one response option condition evokes affective evaluations (Pleasant-Unpleasant), the other may require cognitive evaluations (Peaceful-Aggressive). The predicted negative Arabic effect consistently appeared in the affective response option data, but not in the cognitive response option data. This *may* suggest that the modified AMP task did reflect implicit evaluations when possible, but the pilot study and Peaceful-Aggressive variants did not permit reflexive evaluations. Nevertheless, the data do suggest that when affective judgments are made about Arabic symbols, they are expected to be less positive than Chinese symbols.

Table 7.1: Simple summary of hypothesis support across studies. ✓: Full support, /: Partial support (substantively, or probability > .9), ?: Uncertain due to noise, X: Unsupported.

| Hypothesis | Study 1 | Study 2 | Study 3 | Meta |
|------------|---------|---------|---------|------|
| 1          | /?      | ✓/      | ✓       | ✓?   |
| 2          | ??      | X?      | /       | XX   |
| 3          | /?      | ✓?      | X       | XX   |
| 4          | /?      | /?      | X       | /?   |
| 5          | ✓?      | //      | ✓       | ✓/   |
| 6          | ??      | //      | /       | //   |
| 7          | —       | —       | /       | —    |
| 8          | —       | —       | /       | —    |

Support for the hypotheses are discussed in turn. To clarify, superscripts denote the study (and response option, where applicable) that supports the given statement. Numbers refer to the study (m denotes meta-analysis), letters to the response option (“PI” = pleasant, “Pe” = peaceful), and italics for whether the claim is only partially supported due to uncertainty (i.e., the direction is only known with .90 probability).

The first hypothesis, that Arabic symbols receive less positive ratings than Chinese symbols, was generally supported.<sup>*1(PI),2(PI),2(Pe),3,M(PI)*</sup> Although consistent across “Pleasant-Unpleasant” data, this effect was weak and uncertain in “Aggressive-Peaceful” data.

The second hypothesis, that Trump exposure would strengthen the negative Arabic effect, was unsupported. Although a small negative interaction occurred once<sup>3</sup>, the effect was otherwise absent<sup>*1,2(Pe),M*</sup> or even positive.<sup>*2(PI)*</sup>

Support for third hypothesis, that Sanders exposure would attenuate the negative Arabic effect, is less clear. Positive interactions emerged<sup>*1(PI),2(PI)*</sup>, but otherwise were too variable to discern<sup>*1(Pe),2(Pe)*</sup> or were too small to consider.<sup>*3,M*</sup>

The fourth hypothesis, that the negative Arabic effect would be stronger for those high in SDO and RWA gained mixed support. The negative Arabic effect was moderated by SDO,<sup>*1(PI),2(PI),M(PI)*</sup> but this was not observed in study 3. Unexpectedly, RWA (as a unidimensional factor), did not consistently nor strongly influence the Arabic effect.<sup>*2(Pe)*</sup> The two facets of RWA demonstrated differential effects. The “authoritarian aggression and submission” facet did not discernibly influence the Arabic effect.<sup>*1,2,3,M*</sup> The “conservatism” facet demonstrated the predicted effect, albeit uncertainly.<sup>*2,M*</sup>

The fifth hypothesis; that those high in SDO<sup>*1(PI),2(Pe),3,M(PI),M(Pe)*</sup>, RWA<sup>*1(PI),2(PI),3,M(PI),M(Pe)*</sup>, and anti-Arab prejudice<sup>*1(PI),2,3,M*</sup> would respond positively following a Trump exposure; was strongly supported. Both facets of RWA exhibited the predicted positive interaction with Trump exposures<sup>*1(PI),2(PI),3(PI),M*</sup>. Moreover, the effect sizes were strikingly high. Although on average, responses following a Trump image were highly likely to be negative, these factors predict positive responses within their domains.

The sixth hypothesis, that those *low* in  $SDO^{2,M(Pl)}$ ,  $RWA^{2,3,M}$ , and anti-Arab prejudice $^{2,3,M}$  would respond positively following a Sanders exposure, was not as consistent, but generally supported. Both the  $AAS^{2(Pl),M}$  and C facets $^{2(Pl),3,M}$  exhibited the predicted negative interaction with Sanders exposures.

The seventh hypothesis, that  $SDO^3$  and  $RWA^3$  negatively moderate the group symbol effect, gained mixed support. The moderating effect of RWA appears to be due to the AAS facet, and not the C facet (see Table 6.3). Finally, the eighth hypothesis, that generalized prejudice would predict positive judgments of random symbols, gained partial support<sup>3</sup>. However, when plotted, the effect is not exactly as predicted. Specifically, those low in generalized prejudice are expected to rate group symbols preferably to random glyphs; those high in generalized prejudice tend not to prefer group symbols over either Arabic or Chinese symbols, on average. They do, however, prefer random symbols over Arabic symbols (See Figure 5.1).

Perhaps obviously, both Trump and Sanders support predicted negative responses on Sanders $^{2,3,M(Pl),M(Pe)}$  and Trump trials $^{1(Pl),2,3,M}$ , respectively. Whether Trump or Sanders support strengthens the Arabic effect is tenuous. Trump support only notably strengthened the negative Arabic effect once $^{1(Pl)}$ , and this effect was otherwise indiscernible or absent. Sanders support, however, consistently weakened the negative Arabic effect $^{1(Pl),2(Pl),3,M(Pl)}$ . Interestingly, Sanders support moderated the Trump-Arabic interaction $^{2(Pl),3,M(Pl)}$ , such that those high in Sanders support respond more negatively following an Arabic image if preceded by a Trump image. Third-order interactions are notoriously difficult to interpret. It is possible that those high in Sanders support have a strong association between Trump and anti-Arab prejudice, and by seeing Trump, an unpleasant response to Arabic symbols is prepotent. Alternatively, that same association may simply produce an unpleasant affect when seeing the two stimuli paired with one another, or Arabic symbols especially cause Sanders supporters to respond unpleasant when Trump is observed.

Explicit Anti-Arab prejudice <sup>$I(Pl),3,M(Pl)$</sup> , SDO <sup>$I(Pl),1(Pe),3,M(Pl),M(Pe)$</sup> , RWA <sup>$1(Pl),3,M(Pl)$</sup> , Trump support <sup>$I(Pl),3,M(Pl)$</sup> , and generalized prejudice<sup>3</sup> all predict negative responses regardless of the symbol — Sanders support predicts positive responses <sup>$I(Pe),3,M(Pl)$</sup> . The reason for this is unknown. Given that these traits are all strongly correlated with positive responses following Trump exposures, it is possible that those high in such traits simply responded negatively to any stimulus not involving Trump. This would explain the positive effect of Trump warmth on the Trump-Arabic interaction observed in study 2 — The more one supports Trump, the negative Arabic effect is further attenuated if the image follows a Trump image.

The most consistent results are the most troubling. Namely, Arabic symbols are rated as unpleasant compared to Chinese symbols, and this effect is largest for those high in SDO. Moreover, Arabic symbols are rated as unpleasant, compared to random glyphs, for those high in RWA, SDO, and generalized prejudice. These results are consistent with the past literature on anti-Arab prejudice (Altemeyer & Hunsberger, 1992; Cohrs & Asbrock, 2009; Echebarria-Echabe & Guede, 2007; Henry et al., 2005; Johnson et al., 2012; McFarland, 2010; Nosek et al., 2007; Pedersen & Hartley, 2012; Pratto et al., 1994; Rowatt et al., 2005; Uenal, 2016). In this case, however, the manifestation of anti-Arab prejudice appears apparent even in a minimal paradigm, using models with maximal uncertainty. That is — Anti-Arab prejudice is observable in the affective ratings of randomly generated Arabic script, analyzed with models that permit heterogeneity in effects and uncertainty in scores on the relevant factors. The mean response following an Arabic symbol is relatively negative in the sampled population. Variability, of course, exists within the population, but a person chosen at random has a 62–67% probability of exhibiting the negative Arabic effect, assuming the expected values of the models across studies 1–3.

Those high in anti-Arab prejudice, SDO, RWA, and generalized prejudice were consistently more likely to respond positively after seeing Trump. This suggests that Trump and his ideology are well-received by those high in such traits. Such individuals are likely

to respond negatively to any group symbol employed across studies, and respond more negatively to Arabic symbols than randomly generated glyphs.

Conversely, those low in such traits are likely to respond positively following Sanders images, supporting the notion that the ideologies of Sanders and Trump are fairly opposite of one another. Moreover, those supporting Sanders tended to not exhibit the negative Arabic effect.

Thankfully, the notion that a Trump image would predict a stronger negative Arabic effect was unsupported. However, this does not suggest a lack of normalization of anti-Arab prejudice, nor that such an ideology is unrelated to anti-Arab prejudice. Likewise, brief Sanders exposures did not notably alter the negative Arabic effect. This suggests that a boundary condition was found, wherein brief exposures to such political icons do not appear to notably alter reactions to Arabic symbols. Ideology may be a potent predictor for anti-Arab attitudes, but it is not so strong that brief reminders of them or corresponding icons alter immediate attitudes toward groups.

### *7.1 Limitations*

The studies were purposefully minimal. By utilizing a minimal design in which political icons precede ratings of random group symbols, we extended the boundary for how anti-Arab prejudice is manifested. I.e., if individuals are biased against symbols that merely appear Arabic, then anti-Arab prejudice must be fairly popular and potent. Indeed, explicit anti-Arab prejudice consistently predicted comparatively negative ratings of Arabic symbols.<sup>1(PI),2,3,M(PI)</sup> However, by doing so there was ample uncertainty in the estimates themselves. For the purposes of assessing directional hypotheses, this is less of a concern, but readers should attend to the credible intervals before assessing whether the effects are uniformly small or large. Nevertheless, any consistent results in the context of a minimal design for assessing prejudice are concerning.

The studies only employed Chinese, Arabic, and random glyphs. Although many predictions were supported, the generalizability is constrained by the stimuli used. It appears that Chinese symbols are preferred over both Arabic and random glyphs, but would presumably not be preferred over ingroup symbols. It remains possible that Chinese symbols elicit fewer negative responses due to frequent exposure in popular culture and art. Future studies could compare responses to Arabic or other outgroup symbols with ingroup symbols. Romantic and Cyrillic character sets can appear similar to those used in the English language, and thus the symbols may be semantically meaningful to American monolinguals. The benefit of choosing Arabic and Chinese symbols is that both languages' character sets are semantically unrecognizable to monolingual American speakers. One would either need to find known ingroup symbols that are semantically meaningless to a monolingual American, or use a separate methodology altogether in order to study expressions of prejudice toward benign cultural associates, relative to ingroup symbols.

All studies assume that Trump and Sanders represent cohesive ideologies, and that exposure to such political icons is an indirect method of signaling an ideology. Although pilot data support this to some degree, individuals may simply be reacting to the individuals, per se, and not to their respective ideologies. Those who are high in SDO, RWA, and anti-Arab attitudes likely select themselves into groups and consume media that are ideologically similar to them. These groups and the selected media may purposefully associate Trump with positive characteristics and Sanders with negative characteristics. The observation that individuals evaluate the politicians as a function of these traits does not necessarily imply that the politicians are in fact high in these traits, only that individuals high in these traits evaluate the candidates differently. Moreover, the studies employed Trump and Sanders due to their contemporary influence and popularity in politics. Obviously future studies or replication efforts must assess who acts as ideological icons in future zeitgeists.

Finally, the pilot study indicated that individuals are able to explicitly, accurately categorize stimuli, despite the masks on the speeded task. This raises concerns that responses

to the stimuli reflect explicit, and not implicit evaluations. However, effects were relatively consistent when the response options were affective (Unpleasant-Pleasant) compared to when the response options were not (Peaceful-Aggressive). If explicit prejudicial beliefs about Arabs include the propensity for violence and terror, and the task is assessing explicit prejudicial attitudes, then the negative Arabic effect should have been relatively similar between response conditions. It is therefore possible that the task does evaluate reflexive evaluations of group symbols when the response option permits it. Subtle, affective response options may better facilitate the manifestation of implicit prejudicial behaviors than harder, reflective response options. Conversely, it is possible that the task does evaluate explicit attitudes toward Arabs, and individuals simply have a stronger affective response than a belief that any given symbol is peaceful or aggressive. Finally, it is possible that the task does reflect implicit prejudices when the response option is subtle and affective, but the Peaceful-Aggressive response option may promote socially desirable responding; however, if this were true, those high in explicit anti-Arab prejudice should not have this concern (given their positive endorsement of the items), and a moderation effect should have emerged.

Nevertheless, the studies suggest that individuals are likely to perceive random, meaningless Arabic symbols as less pleasant than Chinese symbols, and random glyphs when high in explicit anti-Arab prejudice, SDO, RWA (Conservatism facet), and generalized prejudice. Presumably, this implies that anti-Arab prejudicial attitudes are potent enough to manifest in speeded evaluations of benign words typeset in an Arabic font, and common enough that most individuals in the sampled population are expected to demonstrate a bias against these Arabic symbols. Although prejudice toward Arabs is a well-studied and well-established phenomenon, these studies demonstrate how influential such prejudices may be.

## APPENDICES

## APPENDIX A

### Political Candidate Attitude Measure

*Instructions* On a scale ranging from 1 (Not at all) to 5 (Extremely), answer the following questions.

- (1) Are your feelings toward Donald Trump generally warm and favorable?
- (2) Overall, would you say you like Donald Trump?
- (3) To what extent is Donald Trump a competent leader?
- (4) To what extent do you think Donald Trump is a moral leader?
- (5) To what extent do you share Donald Trumps positions on political issues?
- (6) To what extent do you support Donald Trump in the upcoming elections?
- (7) To what extent do you think Donald Trump is anti-Arab?
- (8) To what extent do you think Donald Trump is anti-Muslim?
- (9) To what extent do you think Donald Trump is anti-Chinese?
- (10) To what extent do you think Donald Trump believes Arabic individuals are aggressive?
- (11) To what extent do you think Donald Trump believes Muslim individuals are aggressive?
- (12) To what extent do you think Donald Trump believes Chinese individuals are aggressive?
- (13) To what extent do you think Donald Trump is pleasant?
- (14) Are your feelings toward Bernie Sanders generally warm and favorable?
- (15) Overall, would you say you like Bernie Sanders?

- (16) To what extent is Bernie Sanders a competent leader?
- (17) To what extent do you think Bernie Sanders is a moral leader?
- (18) To what extent do you share Bernie Sanderss positions on political issues?
- (19) To what extent do you support Bernie Sanders in the upcoming elections?
- (20) To what extent do you think Bernie Sanders is anti-Arab?
- (21) To what extent do you think Bernie Sanders is anti-Muslim?
- (22) To what extent do you think Bernie Sanders is anti-Chinese?
- (23) To what extent do you think Bernie Sanders believes Arabic individuals are aggressive?
- (24) To what extent do you think Bernie Sanders believes Muslim individuals are aggressive?
- (25) To what extent do you think Bernie Sanders believes Chinese individuals are aggressive?
- (26) To what extent do you think Bernie Sanders is pleasant?
- (27) Are your feelings toward Arabic individuals warm and favorable?
- (28) Are your feelings toward Muslim individuals warm and favorable?
- (29) Are your feelings toward Chinese individuals warm and favorable?
- (30) To what extent are you familiar with Donald Trump?
- (31) To what extent are you familiar with Bernie Sanders?

## APPENDIX B

### Revised Anti-Arab Prejudice Scale

- (1) Islam is an archaic religion, unable to adapt to the present.
- (2) Islam respects human rights.
- (3) Separation between religion and state is impossible in the Muslim culture.
- (4) Islam is a threat for women.
- (5) Hatred against the West is common among Arabs.
- (6) Most Arab countries are fanatics.
- (7) Islam is radical and intolerant.
- (8) Arabs are all the same. They are resentful against the west.
- (9) Arabs have positively contributed to culture and science.
- (10) The Western culture is superior to Arabic culture.
- (11) Most Arabs are glad about terrorism against Western interests.
- (12) Islam is not strictly a religion, but a terrorist movement.
- (13) Arabs are a future threat for my country.
- (14) Arabs love peace and coexistence.
- (15) Islam is a religion and culture deserving of respect.
- (16) Police should pay special attention to Arabic residents.
- (17) Arabs who reject our culture and traditions should return to their countries.

## APPENDIX C

### Prime Images and Mask



APPENDIX D

Chinese Symbols

|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 秋 | 媛 | 拳 | 枝 | 本 | 空 |
| 立 | 輪 | 社 | 車 | 对 | 英 |
| 定 | 井 | 報 | 村 | 採 | 壳 |
| 流 | 枚 | 著 | 馬 | 科 | 雷 |
| 会 | 裏 | 希 | 揭 | 路 | 重 |
| 房 | 言 | 的 | 全 | 客 | 秒 |

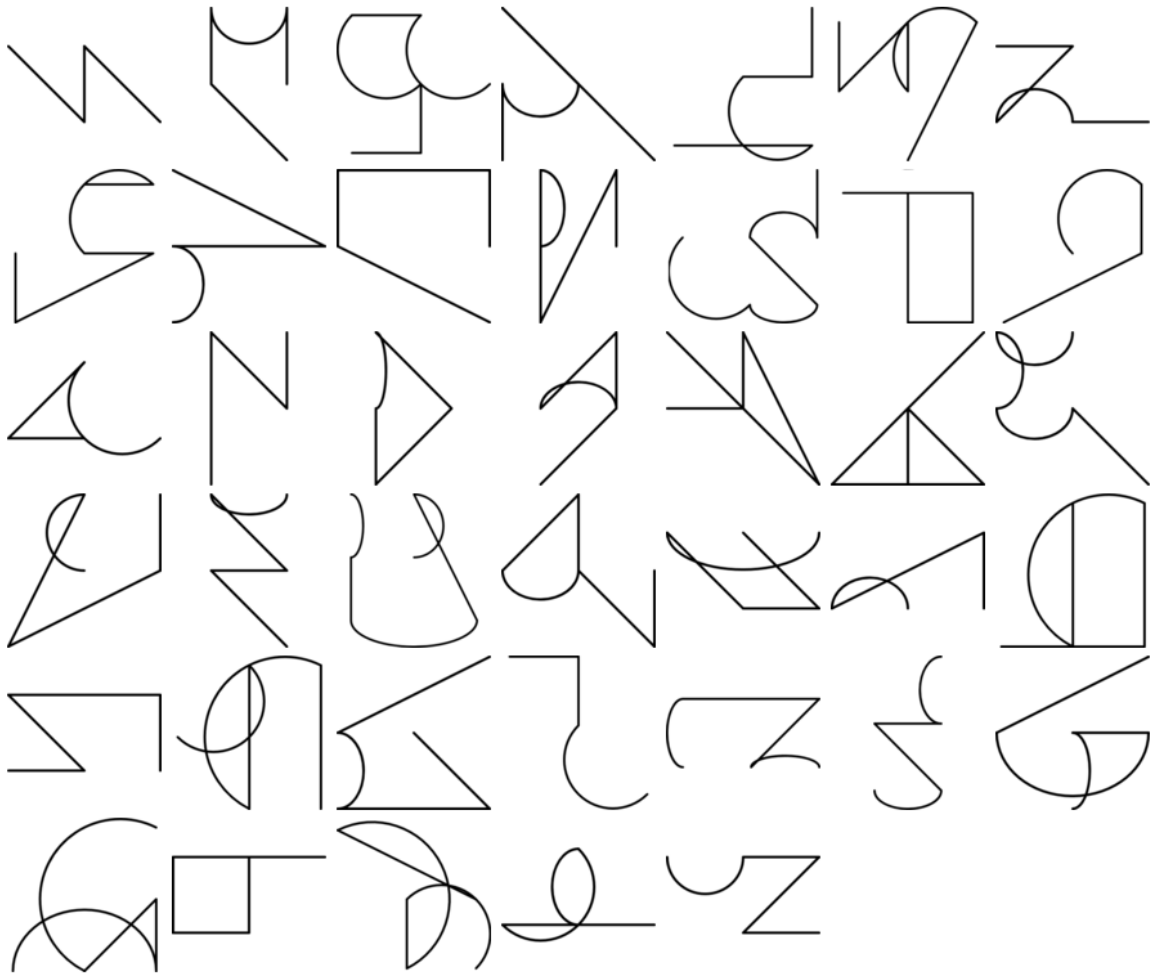
## APPENDIX E

### Arabic Symbols

|         |         |          |             |          |          |
|---------|---------|----------|-------------|----------|----------|
| بقيادة  | وتتامت. | السادس,  | الى         | الشرقية. | ببعض     |
| تاريخ   | إعلان   | وانتهاءً | يرتبط       | سقوط     | العمليات |
| الإتحاد | أفريقيا | الواقعة  | العظمى      | فشكّل    | أضف      |
| قد      | نفس,    | الضغوط   | واندونيسيا، | بأضرار   | العالم   |
| يبيق    | تكاليف  | هامش     | مع          | العدّ    | تعد      |
| وصل     | الثانية | قتيل،    | ليتسنّى,    | شيء      | مايو     |

## APPENDIX F

### Random Glyph Symbols



## APPENDIX G

### Thermometer Items (Reversed Generalized Prejudice)

*Instructions* Indicate how warm your feelings are toward each listed group on a scale from 1 (Feel least warm) to 7 (Feel most warm)

- (1) Caucasian
- (2) Arabic
- (3) Chinese
- (4) African American
- (5) Hispanic
- (6) Native American
- (7) European
- (8) East Asian
- (9) Middle Eastern
- (10) South Asian
- (11) African

## REFERENCES

- A fair and humane immigration policy. (2017, February 7). Retrieved February 7, 2017, from <https://berniesanders.com/issues/a-fair-and-humane-immigration-policy/>
- Agerström, J., & Rooth, D.-O. (2009). Implicit prejudice and ethnic minorities: Arab-Muslims in Sweden. *International Journal of Manpower*, 30(1/2), 43–55. doi:10.1108/01437720910948384
- Ahn, T. K., Janssen, M. A., & Ostrom, E. (2004). Signals, symbols, and human cooperation. In R. W. Sussman & A. R. Chapman (Eds.), *The origins and nature of sociality* (Chap. 6, pp. 122–139). New York: Aldine De Gruyter.
- Alcorta, C. S., & Sosis, R. (2005). Ritual, emotion, and sacred symbols. *Human Nature*, 16(4), 323–359. doi:10.1007/s12110-005-1014-3
- Altemeyer, B., & Hunsberger, B. (1992). Authoritarianism, religious fundamentalism, quest, and prejudice. *International Journal for the Psychology of Religion*, 2(2), 113–133. doi:10.1207/s15327582ijpr0202\_5
- Baldassarri, D., & Gelman, A. (2008). Partisans without constraint: Political polarization and trends in american public opinion. *American Journal of Sociology*, 114(2), 408–447. doi:10.2139/ssrn.1010098
- Becker, J. C., Enders-Comberg, A., Wagner, U., Christ, O., & Butz, D. A. (2012). Beware of national symbols: How flags can threaten intergroup relations. *Social Psychology*, 43(1), 3–6. doi:10.1027/1864-9335/a000073
- Bilewicz, M., Soral, W., Marchlewska, M., & Winiewski, M. (2017). When authoritarians confront prejudice. Differential effects of SDO and RWA on support for hate-speech prohibition. *Political Psychology*, 38(1), 87–99. doi:10.1111/pops.12313
- Bolin, J. L., & Hamilton, L. C. (2018). The news you choose: News media preferences amplify views on climate change. *Environmental Politics*, 27(3), 455–476. doi:10.1080/09644016.2018.1423909
- Bulbulia, J. (2012). Spreading order: Religion, cooperative niche construction, and risky coordination problems. *Biology and Philosophy*, 27(1), 1–27. doi:10.1007/s10539-011-9295-x

- Bürkner, P.-C. (2017). brms: An R package for bayesian multilevel models using stan. *Journal of Statistical Software*. doi:10.18637/jss.v080.i01
- Bushman, B. J. (1984). Perceived symbols of authority and their influence on compliance. *Journal of Applied Social Psychology*, 14(6), 501–508. doi:10.1111/j.1559-1816.1984.tb02255.x
- Butz, D. A. (2009). National symbols as agents of psychological and social change. *Political Psychology*, 30(5), 779–804. doi:10.1111/j.1467-9221.2009.00725.x
- Butz, D. A., Plant, E. A., & Doerr, C. E. (2007). Liberty and justice for all? Implications of exposure to the U.S. flag for intergroup relations. *Personality and Social Psychology Bulletin*, 33(3), 396–408. doi:10.1177/0146167206296299
- Callahan, S. P., & Ledgerwood, A. (2016). On the psychological function of flags and logos: Group identity symbols increase perceived entitativity. *Journal of Personality and Social Psychology*, 110(4), 528–550. doi:10.1037/pspi0000047
- Cecil, M., Kate, B. F., & Sibley, C. G. (2017). Generalized and specific components of prejudice: The decomposition of intergroup context effects. *European Journal of Social Psychology*, 47(4), 443–456. doi:10.1002/ejsp.2252
- Clifford, S., & Jerit, J. (2013). How words do the work of politics: Moral foundations theory and the debate over stem cell research. *The Journal of Politics*, 75(3), 659–671.
- Cohrs, J. C., & Asbrock, F. (2009). Right-wing authoritarianism, social dominance orientation and prejudice against threatening and competitive ethnic groups. *European Journal of Social Psychology*, 39(2), 270–289. doi:10.1002/ejsp.545
- Coryn, C. L., Beale, J. M., & Myers, K. M. (2004). Response to September 11: Anxiety, patriotism, and prejudice in the aftermath of terror. *Current Research in Social Psychology*, 9(12), 1–28.
- Crandall, C. S., Miller, J. M., & White, M. H., II. (2018). Changing norms following the 2016 U.S. presidential election: The trump effect on prejudice. *Social Psychological and Personality Science*, 9(2), 186–192. doi:10.1177/1948550617750735
- Crandall, C. S., & White, M. H., II. (2016). Trump and the Social Psychology of Prejudice. Retrieved April 19, 2017, from <https://undark.org/article/trump-social-psychology-prejudice-unleashed/>

- Crawford, J. T., Jussim, L., Cain, T. R., & Cohen, F. (2013). Right-wing authoritarianism and social dominance orientation differentially predict biased evaluations of media reports. *Journal of Applied Social Psychology*, 43(1), 163–174. doi:10.1111/j.1559-1816.2012.00990.x
- Day, M. V., Fiske, S. T., Downing, E. L., & Trail, T. E. (2014). Shifting liberal and conservative attitudes using moral foundations theory. *Personality and Social Psychology Bulletin*, 40(12), 1559–1573. doi:10.1177/0146167214551152
- Derous, E., Nguyen, H.-H., & Ryan, A. M. (2009). Hiring discrimination against Arab minorities: Interactions between prejudice and job characteristics. *Human Performance*, 22(4), 297–320. doi:10.1080/08959280903120261
- Derous, E., Ryan, A. M., & Nguyen, H. H. D. (2012). Multiple categorization in resume screening: Examining effects on hiring discrimination against Arab applicants in field and lab settings. *Journal of Organizational Behavior*, 33(4), 544–570. doi:10.1002/job.769. arXiv: 0803973233
- Dunwoody, P. T., & McFarland, S. G. (2018). Support for anti-Muslim policies: The role of political traits and threat perception. *Political Psychology*, 39(1), 86–106. doi:10.1111/pops.12405
- Echebarria-Echabe, A., & Guede, E. F. (2007). A new measure of anti-Arab prejudice: Reliability and validity evidence. *Journal of Applied Social Psychology*, 37(5), 1077–1091. doi:10.1111/j.1559-1816.2007.00200.x
- Ehrlinger, J., Plant, E. A., Eibach, R. P., Columb, C. J., Goplen, J. L., Kunstman, J. W., & Butz, D. A. (2011). How exposure to the confederate flag affects willingness to vote for Barack Obama. *Political Psychology*, 32(1), 131–146. doi:10.1111/j.1467-9221.2010.00797.x
- Feinberg, M., & Willer, R. (2013). The moral roots of environmental attitudes. *Psychological Science*, 24(1), 56–62. doi:10.1177/0956797612449177
- Gawronski, B., & Ye, Y. (2013). What drives priming effects in the affect misattribution procedure? doi:10.1177/0146167213502548
- Gelman, A., Jakulin, A., Pittau, M. G., Su, Y.-S., et al. (2008). A weakly informative default prior distribution for logistic and other regression models. *The Annals of Applied Statistics*, 2(4), 1360–1383. doi:10.1214/08-AOAS191
- Gelman, A., Simpson, D., & Betancourt, M. (2017). The prior can generally only be understood in the context of the likelihood. arXiv: 1708.07487

- Graham, J., Haidt, J., & Nosek, B. A. (2009). Liberals and conservatives rely on different sets of moral foundations. *Journal of personality and social psychology*, 96(5), 1029–1045. doi:10.1037/a0015141
- Graham, J., Nosek, B. A., Haidt, J., Iyer, R., Koleva, S., & Ditto, P. H. (2011). Mapping the moral domain. *Journal of personality and social psychology*, 101(2), 366. doi:10.1037/a0021847
- Graham, J. W., Hofer, S. M., & MacKinnon, D. P. (1996). Maximizing the usefulness of data obtained with planned missing value patterns: An application of maximum likelihood procedures. *Multivariate Behavioral Research*, 31(2), 197–218. doi:10.1207/s15327906mbr3102\_3
- Gray, K., & Keeney, J. E. (2015). Disconfirming moral foundations theory on its own terms: Reply to Graham (2015). *Social Psychological and Personality Science*, 6(8), 874–877. doi:10.1177/1948550615592243
- Haidt, J. (2007). The new synthesis in moral psychology. *Science*, 316(5827), 998–1002. doi:10.1126/science.1137651
- Haidt, J., & Graham, J. (2007). When morality opposes justice: Conservatives have moral intuitions that liberals may not recognize. *Social Justice Research*, 20(1), 98–116. doi:10.1007/s11211-007-0034-z
- Hall, R. E. (2003). A note on September eleventh: The Arabization of terrorism. *The Social Science Journal*, 40(3), 459–464. doi:10.1016/S0362-3319(03)00042-9
- Henry, P. J., Sidanius, J., Levin, S., & Pratto, F. (2005). Social dominance orientation, authoritarianism, and support for intergroup violence between the middle east and America. *Political Psychology*, 26(4), 569–583. doi:10.1111/j.1467-9221.2005.00432.x
- Immigration reform that will make America great again. (2018, February 7). Retrieved February 7, 2018, from <https://assets.donaldjtrump.com/Immigration-Reform-Trump.pdf>
- Iyengar, S., & Hahn, K. S. (2009). Red media, blue media: Evidence of ideological selectivity in media use. *Journal of Communication*, 59, 19–39. doi:10.1111/j.1460-2466.2008.01402.x

- Johnson, M. K., Labouff, J. P., Rowatt, W. C., Patock-Peckham, J. A., & Carlisle, R. D. (2012). Facets of right-wing authoritarianism mediate the relationship between religious fundamentalism and attitudes toward arabs and african americans. *Journal for the Scientific Study of Religion*, 51(1), 128–142. doi:10.1111/j.1468-5906.2011.01622.x
- Kugler, M., Jost, J. T., & Noorbaloochi, S. (2014). Another look at moral foundations theory: Do authoritarianism and social dominance orientation explain liberal-conservative differences in moral intuitions? *Social Justice Research*, 27(4), 413–431. doi:10.1007/s11211-014-0223-5
- Lawrence, E., Sides, J., & Farrell, H. (2010). Self-segregation or deliberation? Blog readership, participation, and polarization in American politics. *Perspectives on Politics*, 8(1), 141–157. doi:10.1017/S1537592709992714
- Layman, G. C., Carsey, T. M., & Horowitz, J. M. (2006). Party polarization in American politics: Characteristics, causes, and consequences. *Annual Review of Political Science*, 9, 83–110. doi:10.1146/annurev.polisci.9070204.105138
- Levedusky, M. S. (2013). Why do partisan media polarize viewers? *American Journal of Political Science*, 57(3), 611–623. doi:10.1111/ajps.12008
- Lewandowski, D., Kurowicka, D., & Joe, H. (2009). Generating random correlation matrices based on vines and extended onion method. *Journal of Multivariate Analysis*, 100(9), 1989–2001. doi:10.1016/j.jmva.2009.04.008
- Lourenço, B. (2015). Glyph generator [computer software]. Retrieved April 17, 2017, from <https://github.com/MadEqua/glyph-generator>
- Manganelli Rattazzi, A. M., Bobbio, A., & Canova, L. (2007). A short version of the Right-Wing Authoritarianism (RWA) Scale. *Personality and Individual Differences*, 43(5), 1223–1234. doi:10.1016/j.paid.2007.03.013
- Mange, J., Chun, W. Y., Sharvit, K., & Belanger, J. J. (2012). Thinking about Arabs and Muslims makes Americans shoot faster: Effects of category accessibility on aggressive responses in a shooter paradigm. *European Journal of Social Psychology*, 42(5), 552–556. doi:10.1002/ejsp.1883
- McFarland, S. (2010). Authoritarianism, social dominance, and other roots of generalized prejudice. *Political Psychology*, 31(3), 453–477. doi:10.1111/j.1467-9221.2010.00765.x

- Milojev, P., Osborne, D., Greaves, L. M., Bulbulia, J., Wilson, M. S., Davies, C. L., . . . Sibley, C. G. (2014). Right-wing authoritarianism and social dominance orientation predict different moral signatures. *Social Justice Research*, 27(2), 149–174. doi:10.1007/s11211-014-0213-7
- Nosek, B. A., Smyth, F. L., Hansen, J. J., Devos, T., Lindner, N. M., Ranganath, K. A., . . . Banaji, M. R. (2007). Pervasiveness and correlates of implicit attitudes and stereotypes. *European Review of Social Psychology*, 18(1), 36–88. doi:10.1080/10463280701489053
- Obama: Why I won't say 'Islamic terrorism'. (2016, September 29). Retrieved September 29, 2016, from <https://www.cnn.com/2016/09/28/politics/obama-radical-islamic-terrorism-cnn-town-hall/index.html>
- Oswald, D. L. (2005). Understanding anti-Arab reactions post-9/11: The role of threats, social categories, and personal ideologies. *Journal of Applied Social Psychology*, 35(9), 1775–1799. doi:10.1111/j.1559-1816.2005.tb02195.x
- Payne, K., Cheng, C. M., Govorun, O., & Stewart, B. D. (2005). An inkblot for attitudes: Affect misattribution as implicit measurement. *Journal of Personality and Social Psychology*, 89(3), 277–293. doi:10.1037/0022-3514.89.3.277
- Payne, K., Hall, D. L., Cameron, C. D., & Bishara, A. J. (2010). A process model of affect misattribution. *Personality and social psychology bulletin*, 36(10), 1397–1408. doi:10.1177/0146167210383440
- Payne, K., & Lundberg, K. (2014). The affect misattribution procedure: Ten years of evidence on reliability, validity, and mechanisms. *Social and Personality Psychology Compass*, 8(12), 672–686. doi:10.1111/spc3.12148
- Pedersen, A., & Hartley, L. K. (2012). Prejudice against Muslim Australians: The role of values, gender and consensus. *Journal of Community and Applied Social Psychology*, 22(3), 239–255. doi:10.1002/casp.1110
- Pereira, C., Vala, J., & Costa-Lopes, R. (2010). From prejudice to discrimination: The legitimizing role of perceived threat in discrimination against immigrants. *European Journal of Social Psychology*, 40(7), 1231–1250. doi:10.1002/ejsp.718. arXiv:1311.3475
- Perry, R., Sibley, C. G., & Duckitt, J. (2013). Dangerous and competitive worldviews: A meta-analysis of their associations with social dominance orientation and right-wing authoritarianism. *Journal of Research in Personality*, 47(1), 116–127. doi:10.1016/j.jrp.2012.10.004

- Poole, K. T., & Rosenthal, H. (1984). The polarization of American politics. *The Journal of Politics*, 46(4), 1061–1079. Retrieved from <http://www.jstor.org/stable/2131242>
- Pratto, F., Sidanius, J., Stallworth, L. M., Malle, B. F., Clements, N., Escobar, M., . . . Pasupathi, M. (1994). Social dominance orientation : A personality variable predicting social and political attitudes. *Journal of Personality and Social Psychology*, 67(4), 741–763. doi:10.1037/0022-3514.67.4.741
- Prior, M. (2013). Media and political polarization. *Annual Review of Political Science*, 16, 101–127. doi:10.1146/annurev-polisci-100711-135242
- Rowatt, W. C., Franklin, L. M., & Cotton, M. (2005). Patterns and personality correlates of implicit and explicit attitudes toward Christian and Muslims. *Journal for the Scientific Study of Religion*, 44(1), 29–43. doi:10.1111/j.1468-5906.2005.00263.x
- Rubin, D. B. (1976). Inference and missing data. *Biometrika*, 63(3), 581–592. doi:10.2307/2335739
- Saleem, M., & Anderson, C. A. (2013). Arabs as terrorists: Effects of stereotypes within violent contexts on attitudes, perceptions, and affect. *Psychology of Violence*, 3(1), 84–99. doi:10.1037/a0030038
- Schein, C., & Gray, K. (2015). The unifying moral dyad: Liberals and conservatives share the same harm-based moral template. *Personality and Social Psychology Bulletin*, 41(8), 1147–1163. doi:10.1177/0146167215591501
- Schein, C., & Gray, K. (2018). The theory of dyadic morality: Reinventing moral judgment by redefining harm. *Personality and Social Psychology Review*, 22(1), 32–70. doi:10.1177/1088868317698288
- Shaver, J. H., Sibley, C. G., Osborne, D., & Bulbulia, J. (2017). News exposure predicts anti-Muslim prejudice. *PLOS ONE*, 12(3). doi:10.1371/journal.pone.0174606
- Sibley, C. G., & Duckitt, J. (2008). Personality and prejudice: A meta-analysis and theoretical review. *Personality and Social Psychology Review*, 12(3), 248–279. doi:10.1177/1088868308319226
- Sibley, C. G., Hovard, W. J., & Duckitt, J. (2011). What's in a flag? Subliminal exposure to New Zealand national symbols and the automatic activation of egalitarian versus dominance values. *The Journal of social psychology*, 151(4), 494–516. doi:10.1080/00224545.2010.503717

- Sinn, J. S., & Hayes, M. W. (2017). Replacing the moral foundations: An evolutionary-coalitional theory of liberal-conservative differences. *Political Psychology*, 38(6), 1043–1064. doi:10.1111/pops.12361
- Stan Development Team. (2017). RStan: The R interface to Stan. R package version 2.17.2. Retrieved from <http://mc-stan.org/>
- Strabac, Z., & Listhaug, O. (2008). Anti-Muslim prejudice in Europe: A multilevel analysis of survey data from 30 countries. *Social Science Research*, 37(1), 268–286. doi:10.1016/j.ssresearch.2007.02.004
- Suhler, C. L., & Churchland, P. (2011). Can innate, modular foundations explain morality? Challenges for haidt's moral foundations theory. *Journal of Cognitive Neuroscience*, 23(9), 2103–2116. doi:10.1162/jocn.2011.21637
- Talaska, C. A., Fiske, S. T., & Chaiken, S. (2008). Legitimizing racial discrimination: Emotions, not beliefs, best predict discrimination in a meta-analysis. *Social Justice Research*, 21(3), 263–296. doi:10.1007/s11211-008-0071-2. arXiv: NIHMS150003
- Terrizzi, J. A., Shook, N. J., & McDaniel, M. A. (2013). The behavioral immune system and social conservatism: A meta-analysis. *Evolution and Human Behavior*, 34(2), 99–108.
- The latest: Trump says 'radical Islamic terrorism' must end. (2017, August 18). Retrieved August 18, 2017, from <http://www.foxnews.com/us/2017/08/18/latest-trump-says-radical-islamic-terrorism-must-end.html>
- Trump ends DACA but gives Congress window to save it. (2017, September 5). Retrieved September 5, 2017, from <http://www.cnn.com/2017/09/05/politics/daca-trump-congress/index.html>
- Trump travel ban: Timeline of a legal journey. (2018, January 20). Retrieved January 20, 2018, from <http://www.foxnews.com/politics/2018/01/20/trump-travel-ban-timeline-legal-journey.html>
- Trump unveils legislation limiting legal immigration. (2017, August 2). Retrieved August 2, 2017, from <https://www.npr.org/2017/08/02/541104795/trump-to-unveil-legislation-limiting-legal-immigration>
- Uenal, F. (2016). Disentangling Islamophobia: The differential effects of symbolic, realistic, and terroristic threat perceptions as mediators between social dominance orientation and Islamophobia. *Journal of Social and Political Psychology*, 4(1), 66–90. doi:10.5964/jspp.v4i1.463

- Widner, D., & Chicoine, S. (2011). It's all in the name: Employment discrimination against Arab Americans. *Sociological Forum*, 26(4), 806–823. doi:10.1111/j.1573-7861.2011.01285.x
- Winterich, K. P., Zhang, Y., & Mittal, V. (2012). How political identity and charity positioning increase donations: Insights from moral foundations theory. *International Journal of Research in Marketing*, 29(4), 346–354. doi:10.1016/j.ijresmar.2012.05.002
- Wong, R. Y.-m., & Hong, Y.-y. (2005). Dynamic influences of culture on cooperation in the prisoner's dilemma. *Psychological Science*, 16(6), 429–434. doi:10.1111/j.0956-7976.2005.01552.x