

Kristina Holler

## The Representation of U.S. Presidential Candidates in International News Media

### A Eurolinguistic Analysis of Media Bias and Evaluative Language in Newspaper Headlines from the 2016, 2020, and 2024 Elections

#### Abstract

This thesis analyzes how U.S. presidential candidates – especially Donald J. Trump – are represented in international newspaper headlines during the 2016, 2020, and 2024 election cycles. Based on a multilingual corpus of over 10,900 headlines from seventeen countries, it applies a Eurolinguistic and discourse-analytic framework to examine evaluative language, naming practices, and moral framing. Methodologically, the study combines large-scale AI-assisted coding with qualitative analysis. The study formulates twelve hypotheses: headlines are expected to be more evaluative than neutral (H1), with more favorable coverage of Democratic candidates than Trump (H2), and ideological alignment between media outlets and candidate portrayal (H3). These patterns are assumed to remain stable across election cycles (H4). Naming practices are expected to be largely neutral (H5), but more value-laden in reference to Trump (H6). A Eurolinguistic dimension predicts a shared negative orientation toward Trump in European newspapers (H7), stronger than in non-European contexts (H8), alongside cross-linguistic and regional convergence (H9-10), reduced ideological polarization (H11), and consistent neutral naming conventions (H12). Findings show that while neutrality predominates, systematic evaluative asymmetries and cross-national patterns persist.

#### Sommaire

Cette thèse analyse la représentation des candidats à la présidence des États-Unis – en particulier Donald J. Trump – dans les titres de la presse internationale lors des cycles électoraux de 2016, 2020 et 2024. Fondée sur un corpus multilingue de plus de 10,900 titres issus de dix-sept pays, elle mobilise une approche eurolinguistique et d'analyse du discours pour étudier le langage évaluatif, les pratiques de dénomination et le cadrage moral. Méthodologiquement, elle combine un codage assisté par IA à grande échelle avec une analyse qualitative. Douze hypothèses sont formulées : les titres seraient plus souvent évaluatifs que neutres (H1), avec une couverture plus favorable des candidats démocrates que de Trump (H2), et un alignement entre orientation idéologique des médias et représentation des candidats (H3). Ces tendances seraient stables dans le temps (H4). Les pratiques de dénomination seraient globalement neutres (H5), mais plus marquées dans le cas de Trump (H6). Une dimension eurolinguistique prévoit une orientation négative commune envers Trump en Europe (H7), plus forte qu'ailleurs (H8), ainsi qu'une convergence linguistique et régionale (H9-10), une polarisation réduite (H11) et des conventions de dénomination homogènes (H12). Les résultats montrent que, malgré une neutralité dominante, des asymétries évaluatives persistent.

#### Zusammenfassung

Diese Arbeit untersucht, wie US-Präsidentschaftskandidaten – insbesondere Donald J. Trump – in internationalen Zeitungsschlagzeilen während der Wahlzyklen 2016, 2020 und 2024 dargestellt werden. Grundlage ist ein mehrsprachiges Korpus von über 10.900 Schlagzeilen aus siebzehn Ländern. Die Analyse erfolgt im Rahmen eines eurolinguistischen und diskursanalytischen Ansatzes und fokussiert bewertende Sprache, Namensgebung und moralisches Framing. Methodisch kombiniert die Studie KI-gestützte Kodierung mit qualitativer Analyse. Es werden zwölf Hypothesen formuliert: Schlagzeilen sind häufiger wertend als neutral (H1), demokratische Kandidaten werden positiver dargestellt als Trump (H2) und die Berichterstattung folgt der ideologischen Ausrichtung der Medien (H3). Diese Muster bleiben über Wahlzyklen hinweg stabil (H4). Die Namensgebung ist überwiegend neutral (H5), jedoch im Falle Trumps stärker wertend (H6). Eine eurolinguistische Perspektive erwartet eine gemeinsame negative Haltung gegenüber Trump in Europa (H7), stärker als außerhalb Europas (H8), sowie sprach- und regionsübergreifende Konvergenz (H9-10), geringere Polarisierung (H11) und einheitliche Namenskonventionen (H12). Die Ergebnisse zeigen, dass trotz dominierender Neutralität systematische Bewertungsasymmetrien bestehen bleiben.

## Overview: Table of Contents

|                                                                        |     |
|------------------------------------------------------------------------|-----|
| 1. Introduction .....                                                  | 54  |
| 1.1. Media Discourse, Democracy, and the Power of Representation ..... | 54  |
| 1.2. Research Aims and Questions .....                                 | 54  |
| 1.3. Addressing Gaps in the Literature .....                           | 55  |
| 1.4. Theoretical Framework and Scholarly Context .....                 | 55  |
| 1.5. Research Design and Methodology .....                             | 56  |
| 1.6. Limitations and Scope .....                                       | 56  |
| 1.7. Contributions and Significance .....                              | 56  |
| 1.8. Structure of the Thesis .....                                     | 57  |
| 1.9. Conclusion: Reframing the Question .....                          | 57  |
| 2. Background .....                                                    | 57  |
| 2.1. Headlines as Discursive Gateways in Political Communication ..... | 57  |
| 2.2. Media Bias .....                                                  | 61  |
| 2.3. Names and Naming .....                                            | 66  |
| 2.4. Political Communication .....                                     | 70  |
| 2.5. Trump and Media .....                                             | 74  |
| 2.6. The Rationale for Analyzing Trump and Headlines .....             | 75  |
| 2.7. Personal Bias and Motivation .....                                | 77  |
| 2.8. Hypotheses .....                                                  | 79  |
| 3. Literature and Research .....                                       | 81  |
| 3.1. Eurolinguistics .....                                             | 82  |
| 3.2. Discourse Analysis with AI .....                                  | 86  |
| 3.3. Trump, Media Bias, and Political Reporting .....                  | 88  |
| 3.4. Research Gap .....                                                | 93  |
| 4. Data .....                                                          | 94  |
| 4.1. Corpus and Data Selection .....                                   | 95  |
| 4.1.1. Europe North .....                                              | 96  |
| 4.1.2. Europe East .....                                               | 96  |
| 4.1.3. Europe South .....                                              | 96  |
| 4.1.4. Europe West .....                                               | 96  |
| 4.1.5. Europe Center .....                                             | 100 |
| 4.1.6. North America .....                                             | 100 |
| 4.1.7. South America .....                                             | 100 |
| 4.1.8. Asia .....                                                      | 101 |
| 4.1.9. Africa .....                                                    | 101 |
| 4.1.10. Oceania .....                                                  | 101 |
| 4.2. Time Frame .....                                                  | 101 |
| 5. Methodology .....                                                   | 102 |
| 5.1. Coding Scheme .....                                               | 102 |
| 5.1.1. Bias .....                                                      | 102 |
| 5.1.2. Name as Given .....                                             | 104 |
| 5.1.3. Value Judgement .....                                           | 105 |
| 5.1.4. Naming Strategies .....                                         | 106 |
| 5.2. AI Test Runs .....                                                | 107 |
| 5.2.1. Data Collection and Sampling .....                              | 107 |
| 5.2.2. Human Coding .....                                              | 109 |
| 5.2.3. AI Coding .....                                                 | 109 |
| 5.2.4. Comparison of Human and AI Coding .....                         | 112 |
| 5.3. Main Data Analysis .....                                          | 117 |
| 5.3.1. Filtering Procedures .....                                      | 117 |
| 5.3.2. Dataset Size and Composition .....                              | 117 |
| 5.3.3. Reliability and Limitations .....                               | 118 |
| 6. Results .....                                                       | 119 |
| 6.1. H1: Distribution of Bias .....                                    | 119 |
| 6.2. H2: Candidate-Specific Bias .....                                 | 120 |
| 6.3. H3: Bias and Ideological Stance of Newspapers .....               | 121 |
| 6.4. H4: Stability Across Electoral Cycles .....                       | 123 |
| 6.5. H5: Naming Practices .....                                        | 123 |
| 6.6. H6: Value-Judgement Labels .....                                  | 124 |
| 6.7. H7: European Unity of Bias .....                                  | 125 |
| 6.8. H8: Europe Versus the Rest of the World .....                     | 126 |
| 6.9. H9: Cross-linguistic Convergence .....                            | 127 |
| 6.10. H10: Regional Consistency .....                                  | 128 |
| 6.11. H11: European Ideological Distinctiveness .....                  | 131 |

|        |                                                                                                        |     |
|--------|--------------------------------------------------------------------------------------------------------|-----|
| 6.12.  | H12. Unified Naming Conventions.....                                                                   | 132 |
| 6.13.  | Summary of Findings.....                                                                               | 133 |
| 7.     | Discussion.....                                                                                        | 133 |
| 7.1.   | Introduction: Framing the Discussion.....                                                              | 133 |
| 7.2.   | Media Bias and the Myth of Partisan Imbalance .....                                                    | 134 |
| 7.2.1. | Revisiting D'Alessio & Allen: Evidence of Bias or Not? .....                                           | 134 |
| 7.2.2. | Evaluating Kuypers' "Conspiracy of Shared Values" Claim .....                                          | 136 |
| 7.2.3. | Is Trump Treated "Unfairly"? .....                                                                     | 137 |
| 7.3.   | Patterns of Bias and Naming Practices .....                                                            | 138 |
| 7.3.1. | Candidate-Specific Differences and the Nature of Evaluative Bias .....                                 | 138 |
| 7.3.2. | Naming Strategies and Value-Laden Labels.....                                                          | 139 |
| 7.3.3. | Moralization as a Media Strategy .....                                                                 | 140 |
| 7.5.   | Temporal and Structural Dynamics of Coverage.....                                                      | 142 |
| 7.5.1. | Stability and Change Across Electoral Cycles .....                                                     | 142 |
| 7.5.2. | Gatekeeping and Coverage Bias .....                                                                    | 144 |
| 7.6.   | European Perspectives: A Mosaic Rather Than a Monolith.....                                            | 145 |
| 7.7.   | Comparative Perspective: Relating to Existing Scholarship.....                                         | 147 |
| 7.7.1. | Patterson (2016): Visibility and Negativity Revisited .....                                            | 147 |
| 7.7.2. | Nord, Mancini & Gerli (2017): Negative Dominance and the Neutrality Turn.....                          | 150 |
| 7.7.3. | Cazzamatta & Hafez (2021): Moralized Critique and Selective Labeling .....                             | 150 |
| 7.7.4. | Hinck (2024): Global Negativity and Regional Variation .....                                           | 151 |
| 7.7.5. | Bykova et al. (2018) and Rovinskaya & Greydina (2021): Temporal Shifts and the Softening of Tone ..... | 151 |
| 7.7.6. | Synthesis: Confirmation, Complication, and Contribution .....                                          | 152 |
| 7.8.   | Methodological Reflections and AI-Assisted Analysis.....                                               | 152 |
| 7.8.1. | Reliability and Performance of AI Coding.....                                                          | 152 |
| 7.8.2. | Naming, Value Judgement, and the Limits of Prompt Design.....                                          | 153 |
| 7.8.3. | Workflow, Human Oversight, and Practical Considerations .....                                          | 153 |
| 7.8.4. | Expanding Analytical Scope: Scale, Language, and Representativeness .....                              | 154 |
| 7.8.5. | Broader Implications for Discourse Analysis.....                                                       | 154 |
| 8.     | Conclusion .....                                                                                       | 154 |
| 8.1.   | Rethinking Media Discourse in a Changing Democratic Landscape .....                                    | 154 |
| 8.2.   | Reflexive Considerations: The Researcher's Role in Discourse .....                                     | 155 |
| 8.3.   | The Broader Significance of the Findings .....                                                         | 155 |
| 8.4.   | Future Research Directions .....                                                                       | 156 |
| 8.5.   | Concluding Reflections .....                                                                           | 157 |
| 9.     | Reflection on the Use of AI in the Scholarly Process.....                                              | 157 |
|        | References .....                                                                                       | 159 |
|        | Appendixes.....                                                                                        | 163 |
| 1.     | H1. Distribution of bias. ....                                                                         | 163 |
| 2.     | H2. Candidate-specific bias. ....                                                                      | 164 |
| 3.     | H3. Bias and ideological stance of newspapers. ....                                                    | 165 |
| 4.     | H4. Stability across electoral cycles. ....                                                            | 167 |
| 5.     | H5. Naming practices. ....                                                                             | 168 |
| 6.     | H6. Value-judgment labels. ....                                                                        | 169 |
| 7.     | H7. European unity of bias. ....                                                                       | 170 |
| 8.     | H8. Europe versus the rest of the world. ....                                                          | 171 |
| 9.     | H9. Cross-linguistic convergence. ....                                                                 | 172 |
| 10.    | H10. Regional consistency. ....                                                                        | 173 |
| 11.    | H11. European ideological distinctiveness. ....                                                        | 177 |
| 12.    | H12. Unified naming conventions.....                                                                   | 179 |
| 13.    | Counts .....                                                                                           | 180 |
| 14.    | Bias .....                                                                                             | 183 |
| 15.    | Value Judgement .....                                                                                  | 184 |

# 1. Introduction

## 1.1. Media Discourse, Democracy, and the Power of Representation

In democratic societies, news media occupy a pivotal role in shaping how political actors are perceived, how legitimacy is constructed, and how public debates unfold. The headlines that appear in major newspapers do more than simply report events: they frame political reality, influence public opinion, and help define the contours of what is considered legitimate political behavior. As discourse analysts such as Bendel Larcher (2023: 72) have argued, texts do not “represent” people and events as they inherently are; rather, they construct particular images of them through language. The choices that journalists make – which words they use, what they emphasize, what they omit – all contribute to a larger discursive process in which political meaning is negotiated and contested. This process is especially consequential in electoral contexts, where voters’ perceptions of candidates are shaped not only by policy proposals and campaign speeches but also by how those candidates are framed in the press.

Understanding this dynamic is particularly urgent in the case of Donald J. Trump, a figure who has redefined the relationship between political actors, the media, and the public. Since the launch of his first presidential campaign in 2015, Trump has positioned himself in direct opposition to mainstream news outlets, repeatedly accusing them of bias, dishonesty, and fake news. At the same time, his campaigns have exploited new channels of direct communication – particularly social media – to bypass traditional gatekeepers and speak directly to supporters (Ross & Rivers 2018). These twin developments have transformed the political communication landscape. On the one hand, mainstream media’s normative functions – such as informing the public, holding power to account, and fostering democratic deliberation – are increasingly contested (Meeks 2020). On the other, public trust in journalism has become polarized, with support or rejection of Trump often determining whether audiences perceive the media as credible (Mourão et al. 2018). Against this backdrop, studying how mainstream news outlets represent Trump and his political rivals is not just an exercise in media analysis; it is a crucial inquiry into how democratic legitimacy, public discourse, and political meaning are constructed in the contemporary era.

## 1.2. Research Aims and Questions

This thesis addresses two interrelated research questions that are central to contemporary media discourse analysis and political communication research:

1. How have mainstream news outlets across Europe and beyond portrayed Donald Trump and his Democratic rivals during the 2016, 2020, and 2024 U.S. presidential election cycles?
2. What linguistic and discursive strategies have been used to construct political meaning in these portrayals, and how do they vary across ideological, geographical, and linguistic contexts?

At its core, the study is concerned with the ways in which media discourse frames political actors – not only whether coverage is positive, negative, or neutral, but how evaluative stances are embedded in linguistic choices, naming practices, and the moral vocabularies that shape public narratives. These questions respond directly to what (Bendel Larcher 2023: 67) identifies as the guiding inquiries of discourse-analytic research: what image of reality a text seeks to convey, what it wants readers to believe, and how it seeks to influence their perceptions and actions.

### 1.3. Addressing Gaps in the Literature

While media bias in U.S. presidential elections has been extensively studied, much of the existing scholarship suffers from two significant limitations that this thesis aims to address.

First, the vast majority of studies remain confined to national or bilateral contexts. Classic analyses such as D'Alessio and Allen's (2000) meta-study focus exclusively on the United States, while comparative research (e.g., Nord, Mancini & Gerli 2017) typically examines only one or two European countries. Such work has provided important insights into media bias but has left open the question of whether the patterns observed in these contexts hold true globally – particularly in a European press landscape marked by diverse journalistic traditions, political cultures, and linguistic systems. This thesis seeks to fill that gap by adopting a genuinely transnational perspective. It analyses more than 10,900 headlines from newspapers in seventeen countries across six continents, with a particular emphasis on the European press. In doing so, it aims to provide the most comprehensive cross-national account to date of how Trump and his rivals have been represented in mainstream news discourse.

Second, the field has been constrained by methodological limitations. Most discourse-analytic studies are small in scale, relying on manually coded datasets that are necessarily limited in size. This restricts their ability to capture broader trends, subtle cross-linguistic differences, or longitudinal shifts over time. A key innovation of the present study is its use of artificial intelligence (AI) as a tool for large-scale discourse analysis. AI-assisted coding enables the systematic classification of thousands of headlines across multiple languages and election cycles – a task that would be prohibitively time-consuming if performed manually. This methodological approach not only expands the empirical reach of discourse analysis but also offers a model for how computational tools can complement qualitative research in the humanities and social sciences.

### 1.4. Theoretical Framework and Scholarly Context

The study is situated at the intersection of several major strands of scholarship: media bias research and discourse-analytic approaches to political communication, focused through the lens of Eurolinguistics.

The question of whether the media exhibit partisan bias has long been debated. Some scholars argue that, over the long term, news coverage tends to balance out. D'Alessio and Allen (2000), for instance, conclude that presidential campaign coverage shows little net bias once variations are averaged across outlets and election cycles. Others contend that mainstream journalism is structurally liberal, reflecting a “conspiracy of shared values” among journalists and editors (Kuypers 2014, 2020). The evidence for either position remains inconclusive and highly context-dependent. This thesis engages with these debates but does not seek to resolve them definitively. Instead, it offers new empirical evidence that both confirms and complicates existing narratives – showing, for example, that neutrality remains the dominant mode of reporting, yet evaluative asymmetries persist, particularly in coverage of Donald Trump.

The study also draws on discourse theory and critical linguistics, which emphasize the constructive power of language in shaping political reality. As scholars like van Leeuwen (2008) and Würth (2016) have shown, naming practices and evaluative labels are not neutral descriptors but discursive acts that frame how political actors are perceived. Whether a candidate is referred to simply by their surname or by a derogatory epithet can subtly but powerfully shape public understanding. Similarly, the use of moralized language – portraying political actors as virtuous or corrupt, saviors or threats – situates them within broader narratives of legitimacy and illegitimacy (Cazzamatta & Hafez 2021). These theoretical

perspectives inform the analytical lens of this thesis, guiding its examination of how evaluative meaning is encoded in the linguistic fabric of media discourse.

### 1.5. Research Design and Methodology

The empirical foundation of this study is a large, multilingual corpus of 10,939 newspaper headlines published during the 2016, 2020, and 2024 U.S. presidential election cycles. These headlines were drawn from mainstream outlets in seventeen countries across six continents, with a particular focus on European newspapers. Each headline was coded for evaluative stance (*positive*, *negative*, *neutral*), naming strategy (*nomination*, *categorization*, *appraisal*; based on van Leeuwen 2008), and the presence of value-laden labels.

The coding process combined AI-assisted classification with human verification. AI tools were employed to process large volumes of text and identify patterns of bias, while human oversight ensured the reliability, validity, and interpretive depth of the analysis. This hybrid methodology allowed the study to achieve a level of scale and representativeness that would be difficult to attain through manual methods alone, while still retaining the theoretical and contextual sensitivity that discourse analysis requires.

### 1.6. Limitations and Scope

As with any study, this research has limitations that must be acknowledged. Most notably, the analysis is restricted to mainstream news outlets, which were selected for their accessibility, comparability, and representativeness. This means that the results cannot be generalized to the broader media ecosystem, particularly non-traditional platforms such as social media, partisan blogs, or alternative news sites, which may employ different discursive strategies and reach different audiences. Furthermore, the study focuses on three electoral cycles (2016, 2020, and 2024) and on a specific subset of texts (headlines), which, while highly influential, represent only one dimension of media discourse. These limitations do not undermine the study's contributions, but they do highlight areas for future research – particularly the need to examine how Trump and his rivals are represented in the non-mainstream media consumed by many of his supporters (Mourão et al. 2018).

### 1.7. Contributions and Significance

Despite these limitations, the thesis makes several key contributions to the fields of media studies, discourse analysis, and political communication.

- Empirical contribution: By analyzing more than 10,900 headlines across seventeen countries and six continents, the study provides one of the most comprehensive cross-national examinations of media bias in presidential campaign coverage to date.
- Theoretical contribution: The findings challenge simplistic narratives about media partisanship, demonstrating that neutrality remains dominant even as evaluative asymmetries persist. They also highlight the centrality of moralized discourse in shaping political narratives.
- Methodological contribution: The thesis pioneers the integration of AI-assisted discourse analysis into large-scale media research, demonstrating how computational tools can expand the scale, scope, and depth of linguistic inquiry.

Together, these contributions offer a richer, more nuanced understanding of how media discourse constructs political meaning in a rapidly evolving information environment. They also provide a foundation

for future research into how alternative media, social platforms, and emerging technologies further transform the landscape of political communication.

## **1.8. Structure of the Thesis**

The remainder of this thesis is structured as follows. Chapter 2 situates the study with a background on media strategies, media bias, and the political communication environment of the Trump era. Chapter 3 provides the theoretical background with an overview of existing research and literature on Eurolinguistics, AI-assisted discourse analysis, and prior research on media bias related to Trump and his competitors. Chapter 4 offers insight into the composition of the newspaper dataset. Chapter 5 details the methodology, including the construction and testing of the coding procedures, and the use of AI-assisted classification. Chapter 6 presents the results of the analysis, tracing patterns of bias, naming practices, ideological alignment, and cross-national variation. Chapter 7 provides a detailed discussion of these findings, examining their implications for theories of media bias, discourse, and political communication. Finally, Chapter 8 concludes the thesis by reflecting on the broader significance of the results, identifying avenues for future research, and considering how the study contributes to our understanding of media discourse in a changing democratic landscape.

## **1.9. Conclusion: Reframing the Question**

This thesis begins from a deceptively simple question: how do mainstream news outlets portray Donald Trump and his political rivals? But answering that question requires confronting deeper issues about how media discourse shapes – and is shaped by – the democratic process, how language constructs political meaning, and how new technologies can transform the scale and methods of discourse analysis. By integrating a large-scale, cross-national empirical study with theoretical insights from discourse analysis and political communication, this research seeks to contribute to our understanding of one of the most consequential relationships in modern politics: the relationship between the media, political power, and the public imagination.

# **2. Background**

## **2.1. Headlines as Discursive Gateways in Political Communication**

In contemporary political communication, newspapers continue to hold a central position as sources of political information, despite long-term declines in print readership and the rise of digital platforms. Surveys suggest that citizens still perceive newspapers as particularly suitable media for following political events, surpassing television and far exceeding non-news online media in perceived reliability (Bucher 2017: 299). This reflects their established role in providing daily or weekly accounts of ongoing events and offering background analyses that orient readers in complex political landscapes. Newspapers therefore serve not only an informational but also an orientational function, furnishing readers with frameworks through which events and decisions can be interpreted.

At the same time, newspapers are more than neutral conveyors of fact. As Bucher (2017: 299) stresses, print media are themselves actors in political contexts: they bring issues onto the public agenda, frame topics in specific ways, and contribute to the processes of scandalization, personalization, and evaluation. In this sense, newspapers are both mirror and factor of political opinion-formation. Their ability to set agendas has long been recognized in media studies as a decisive power of mass communication, since

the very act of selecting some issues for coverage necessarily entails excluding others. This process of *gatekeeping* (D'Alessio & Allen 2000: 135) underscores the active role of journalists and editors in shaping the horizons of political debate.

This agenda-setting capacity becomes visible in headlines that link candidates to specific policy areas or potential crises. For instance,

- (1) Trump win brings range of possible conflicts Friction likely over trade, climate and immigration (*The Irish Times*, November 6, 2016)<sup>1</sup>

does more than announce an electoral outcome. By immediately connecting Trump's anticipated presidency to conflicts in trade, climate policy, and immigration, the headline channels public attention toward these domains as the primary sites of contestation. In this way, newspapers do not merely report on political developments but actively define what counts as politically salient.

Another central dynamic is *personalization*, whereby newspapers foreground the persona of candidates rather than the issues they advocate. Headlines often frame political actors in terms of their likability, character, or personal qualities, thus shaping the affective orientation of readers. For example,

- (2) Hating Hillary Clinton: The Likability Trap (*The New York Times*: October 10, 2016)

reduces a complex candidacy to a question of personal appeal and emotional response, rather than policy positions. Similarly,

- (3) Joe Biden es un buen hombre y un buen presidente<sup>2</sup> (*El Pais*: June 30, 2024)

highlights personal virtue as a defining trait, implicitly linking moral character to political legitimacy. Such instances demonstrate how newspapers personalize politics, constructing candidates as objects of admiration, dislike, or moral evaluation, and thereby guiding readers' judgements beyond the policy sphere.

Beyond *personalization*, newspapers frequently employ strategies of *scandalization* or *negative framing*, where candidates are presented through lenses of danger, threat, or moral failure. For instance,

- (4) Trump, una amenaza para los mexicanos<sup>3</sup> (*El Universal*: September 29, 2016)

casts Trump not merely as a political figure but as an existential danger to an entire national community. Similarly,

- (5) Política externa de Trump pode aumentar sofrimento mundial<sup>4</sup> (*O Globo*: January 18, 2017)

---

<sup>1</sup> For several newspapers, *LexisNexis* provided solely the headline. In the case of others, the subheadings or beginnings of articles were included. It is to be noted that the results provided are from *LexisNexis*; therefore, there are no symbols separating headings and subheadings. For further information regarding the management of search results post-download, please refer to section 5.3.

<sup>2</sup> Joe Biden is a good man and a good president

<sup>3</sup> Trump, a threat to Mexicans

<sup>4</sup> Trump's foreign policy could increase global suffering

situates Trump's presidency as a source of global harm, attributing to him agency over the potential increase of worldwide suffering. In both cases, the headline condenses complex geopolitical debates into stark evaluative judgements, which not only inform but also predispose readers toward particular interpretations of political actors.

Such evaluative framings, however, are not confined to negative portrayals. Newspapers also construct candidates in highly favorable terms, presenting victories as triumphs and mandates as legitimate endorsements. Headlines such as

(6) Brasil festeja a vitória de Trump<sup>5</sup> (*Folha de São Paulo*: November 6, 2024)

(7) Markets up after Trump election victory (*China Daily*: November 10, 2016)

(8) Après sa réélection triomphale, Trump est armé d'un mandat clair pour mener sa politique<sup>6</sup> (*Le Figaro*: November 7, 2024)

frame electoral success not only as a political outcome but as a celebration or a clear popular mandate. These formulations imbue political events with evaluative force, reinforcing legitimacy and amplifying the symbolic resonance of electoral victories.

Part of the continuing influence of newspapers lies in their recognizable structural conventions. Across national and cultural contexts, newspapers share a broadly comparable architecture of headlines, lead paragraphs, body texts, and accompanying multimodal elements such as images, graphics, and typographic design. Their periodicity – whether daily, weekly, or Sunday editions – ensures a rhythmic flow of communication that allows political events to be sequenced and interpreted over time. These conventions provide a stable format that facilitates both domestic and cross-national comparisons, while also creating the conditions under which discursive strategies can be made visible at the level of the headline. From a discourse-analytic perspective, texts, and thus individual headlines, are not self-contained but fragments of a larger conversation (Bendel Larcher 2023: 55). This view underlines why micro-choices in a headline can carry macro-relevance for how political reality is publicly negotiated.

Within this recognizable format, headlines play a particularly crucial communicative role. They function as entry points to the text, designed to capture attention in a highly competitive media environment and to provide readers with a condensed orientation to the article that follows. In this sense, headlines do not simply summarize but also pre-structure the reading process, highlighting certain aspects of an event while leaving others implicit or excluded. Their strategic brevity compels selectivity: a few words must suffice to establish the newsworthiness of a story, signal its central actors, and hint at its evaluative direction. By doing so, headlines guide readers' expectations and influence how the subsequent article will be received.

As Fowler (1991: 1) notes, headlines are conventionally presented as neutral summaries that allow readers to enter a story in an “unambiguous” and efficient way. Yet even in their most compressed form, they direct attention and set interpretive cues. For instance,

(9) Klarer Sieg, wenig Substanz Kommentar von Bernd Pickert zum TV-Duell Trump gegen Harris<sup>7</sup> (*taz*: September 12, 2024)

---

<sup>5</sup> Brazil celebrates Trump's victory

<sup>6</sup> After his triumphant reelection, Trump is armed with a clear mandate to pursue his policies

<sup>7</sup> Clear victory, little substance Commentary by Bernd Pickert on the TV debate between Trump and Harris

signals not only who prevailed in a televised debate but also how that success should be assessed, offering a ready-made lens for reading the article. Similarly,

(10) Trump, Biden und die “verdammt Deutschen”<sup>8</sup> (*Die Welt*: October 21, 2020)

dramatizes a moment of campaign rhetoric in a way that foregrounds conflict and controversy, drawing the reader into the narrative through a stark phrase. The cases demonstrate what Bucher (2017: 307–308) describes as the *Darstellungsprogramme* of print journalism: conventionalized forms through which complex political realities are made communicatively accessible. Headlines thus operate as gateways that orient the reader toward the article’s framing before a single paragraph of body text has been read.

Beyond their communicative gateway function, headlines are also discursively dense texts in which ideological positions become particularly visible. Their brevity forces a high degree of lexical selectivity: each word carries evaluative weight, whether by presenting actors neutrally (e.g., *Biden announces candidacy*) or by embedding judgement (e.g., *Sleepy Joe loses debate*). As Simpson, Mayr, and Statham (2019: 4) argue, all texts are shaped by ideology and cannot be treated as “neutral or disinterested.” Headlines exemplify this claim, since their condensed form strips away hedging and elaboration, leaving evaluative choices starkly exposed. Rhetorical devices such as metaphor, personalization, or hyperbole can further sharpen this ideological imprint, turning a headline into what Fowler (1991: 4) described as not a value-free reflection of facts but a specific representation among many possible alternatives. In this sense, headlines crystallize the tension between journalism’s aspirations to objectivity and its inevitable role in reproducing social and political values.

This evaluative dimension is particularly clear in headlines that go beyond description and instead mobilize figurative or judgmental language. For example,

(11) IT IS TIME FOR AMERICA TO PUT TRUMP IN HISTORY’S BIN (*The Australian*: November 3, 2020)

not only reports on electoral dynamics but advocates a position, urging readers to consign Trump to a place of political obsolescence. Similarly,

(12) Trumpeinstein showed up when politics failed the working class (*Irish Independent*: November 5, 2016)

invokes a metaphor of monstrosity, situating Trump as the grotesque product of systemic failure. Such rhetorical condensation exemplifies what Simpson, Mayr, and Statham (2019: 6) describe as discourse offering “a constellation of different narrative possibilities,” where linguistic choices construct sharply divergent realities.

Yet the very features that make newspapers and their headlines powerful discursive sites also come with limitations that must be acknowledged. Newspapers are not ideologically neutral spaces; ownership structures, editorial priorities, and political orientations shape what is selected for coverage and how it is represented. As Bucher (2017: 299) observes, newspapers simultaneously act as mirrors and factors of opinion formation, which means that their accounts of political actors are always embedded in broader

---

<sup>8</sup> Trump, Biden, and the “damned Germans”

institutional and cultural contexts. Moreover, the diversity of perspectives available varies considerably across media systems. While some countries sustain a wide spectrum of voices, others operate under stronger constraints, with state-controlled outlets narrowing the range of possible representations. This imbalance implies that cross-national analyses cannot rely on perfectly symmetrical data, yet the contrast itself is analytically revealing. Despite these caveats, the accessibility, structural comparability, and discursive salience of newspapers and their headlines render them particularly well suited for investigating how U.S. presidential candidates are represented in international media discourse.

## 2.2. Media Bias

The notion of media bias has long occupied a central place in debates on journalism and political communication, yet its definition is far from straightforward. At a basic level, bias may be understood as systematic distortion in the selection and presentation of news. D'Alessio and Allen (2000: 135–137) propose a typology that distinguishes between three main types: *gatekeeping bias*, which occurs when certain stories are selected for coverage while others are ignored; *coverage bias*, referring to imbalances in the volume of attention allocated to different actors or issue; and *statement bias*, where evaluative language within coverage favors one side over another. This framework highlights that bias can operate both through omission and through expression, both in what is reported and how it is reported.

Yet this tripartite definition captures only part of the picture. Media scholars have long emphasized that news is not a transparent reflection of reality but the outcome of institutional, cultural, and ideological processes. As Gutsche (2015: 7, 114–116) argues, news production entails practices of construction and control: editors act as “gates” that regulate the flow of information, while the very structures of news institutions align reporting with broader systems of social power. From this perspective, bias cannot be reduced to isolated errors or imbalances but is inherent in the ways journalism interacts with politics, business, and culture. Similarly, Simpson, Mayr, and Statham (2019: 4) underline that all texts are “anything but neutral or disinterested.” Language itself is shaped by ideology, such that journalistic texts inevitably reproduce the assumptions and values of their cultural contexts. Differences in expression – whether in headlines, lead paragraphs, or captions – carry ideological distinctions, meaning that every act of reporting is simultaneously an act of framing. In this sense, evaluation operates along a spectrum from implicit to explicit. Implicit evaluations are included within the words themselves, for instance through connotations, euphemisms, or metaphor; explicit evaluations surface via overt attributes or predicate-bound judgement (Bendel Larcher 2023: 102–105). Connotations are intertwined with the words and thus lexicalized, often unnoticed yet consequential; euphemisms can hide unpleasant facts or even mislead; and metaphors direct our perception and can carry strong value judgements (Bendel Larcher 2023: 102–103). Identifying whether evaluative adjectives merely classify or actually judge is culturally specific and requires contextual interpretation rather than simple counting (Bendel Larcher 2023: 80, 104).

These insights complicate the quantitative focus of earlier bias research. While D'Alessio and Allen's typology provides analytical clarity, it leaves underexplored the discursive dimension whereby evaluative choices embed ideology in news texts. Headlines offer a striking illustration: in

(13) Ein Personal zum Fürchten Regieren Abschieben, Öl bohren, Kohle fördern zählen zu Trumps Programm<sup>9</sup> (*taz*: November 10, 2016)

---

<sup>9</sup> A frightening staff governing deportation, oil drilling, and coal mining are part of Trump's program

the newspaper does not simply report policy priorities but situates them within a frame of fear and threat. Conversely,

(14) Неизчерпаем извор на опит: “Вашингтон пост” официално подкрепи Байдън<sup>10</sup> (*Dnevnik*: September 28, 2020)

conveys explicit endorsement, highlighting experience and credibility as positive traits. Both examples demonstrate statement bias in action, where the ideological position of the outlet becomes legible in condensed linguistic form.

To complicate matters further, perceptions of bias are themselves unstable. Feldman (2014: 549–550) has shown that audiences often perceive identical coverage as hostile to their side, a phenomenon labeled the *hostile media effect*. Even when reporting is balanced, partisans may experience it as biased, undermining trust and reinforcing polarization. This suggests that the study of bias must attend not only to content but also to perception, an issue to which later sections will return (see section 0). It also implies attentiveness to omission as a discursive resource. Because readers cannot know what is missing in unfamiliar domains, leaving out information can be a powerful means of shaping discourse (Bendel Larcher 2023: 88-89). Although the present study cannot systematically reconstruct omissions, recognizing this mechanism helps calibrate claims about bias derived from what headlines include.

D’Alessio and Allen’s (2000) meta-analysis of 59 studies on U.S. presidential campaign coverage between 1948 and 1996 remains a touchstone in the literature. Their synthesis, which included newspapers, television network news, and news magazines, yielded a striking conclusion: across outlets and across decades, no significant net bias could be detected. At the newspaper level, neither *gatekeeping*, *coverage*, nor *statement bias* consistently favored Democrats or Republicans; isolated examples of slant in one direction were offset by examples leaning the other way. In television, they observed a marginal pro-Democratic tendency, especially in the distribution of airtime, but the effect was small. News magazines displayed a slight pro-Republican bias, though again the magnitude was negligible. Taken together, these results led D’Alessio and Allen (2000: 148-149) to argue that “on the whole, across all newspapers and all reporters, there is only negligible, if any, net bias in the coverage of presidential campaigns.”

For media critics accustomed to claims of systematic partisanship, this finding was counterintuitive. Yet it illustrates the value of large-scale quantitative synthesis: by aggregating across outlets and election cycles, D’Alessio and Allen demonstrated that apparent imbalances at the micro level often cancel out in the aggregate. In this view, U.S. journalism may occasionally tilt one way or another, but no monolithic “liberal press” or “conservative press” emerges when data are examined comprehensively.

Nonetheless, the strengths of this study also point to its limitations. Its scope was confined to U.S. presidential campaigns, a context with unusually high visibility, professionalized reporting routines, and strong incentives for balance. Its indicators of bias were primarily quantitative: airtime, column inches, and sentence polarity. Such measures provide clarity but may miss the more subtle discursive processes through which ideology surfaces. For instance, evaluative framing can be conveyed through metaphor or personalization even in the absence of overtly positive or negative statements. A headline such as

(15) Elon Musk reportedly makes surprise appearance on Trump-Zelenskyy call (*The Guardian*: November 8, 2024)

---

<sup>10</sup> An inexhaustible source of experience: The Washington Post officially endorses Biden

exemplifies a neutral, descriptive approach in which the event is reported without evaluative gloss. Similarly,

(16) The Decision to Delay Trump’s Sentencing (*The New York Times*: September 9, 2024)

conveys a negative situation in legal terms but does so with a restrained, institutional language, avoiding more emotive or judgmental phrasing. Such cases remind us that headlines can indeed be balanced and informative. Yet at the same time, the very choice to describe events in legalistic or procedural terms rather than, for example, in moralizing language is itself a decision with implications for how readers interpret a candidate’s legitimacy.

For cross-national research in particular, the limits of purely quantitative measures become clear. While they can reveal aggregate patterns of balance or imbalance, they do not address the culturally specific linguistic strategies through which bias is encoded. Nor do they capture variation in media systems, some of which permit a wide range of partisan voices while others impose strong ideological constraints. Thus, while D’Alessio and Allen’s analysis establishes a valuable baseline, it cannot fully account for the ways in which bias may be discursively constructed in European or other international headlines, nor for the evaluative nuances of naming, framing, and word choice.

In sharp contrast to the meta-analytic findings of D’Alessio and Allen, Kuypers (2014; 2020) advances a very different interpretation of the U.S. media landscape. He argues that mainstream newspapers and television outlets are not neutral arbiters of political contestation but systematically biased toward liberal values. For Kuypers, this imbalance stems from what he calls “a conspiracy of shared values” among journalists, rooted in their social backgrounds, educational trajectories, and political leanings. According to this view, reporters and editors overwhelmingly emerge from elite, liberal institutions, enjoy above-average incomes, and disproportionately support Democratic candidates. The result, Kuypers contends, is a news culture that portrays conservative actors in consistently negative ways while reserving positive or neutral labels for liberals.

The rhetorical strategies Kuypers identifies include selective labeling, unbalanced use of sources, and the exclusion of oppositional information. For example, conservative politicians may be described with adjectives such as *ultra-right* or *hardline*, whereas liberal figures are presented more neutrally or even sympathetically (Kuypers 2014: 5). In his later work on Trump-era coverage, Kuypers (2020) argued that this bias reached new intensity: news organizations allegedly distorted or downplayed the content of Trump’s speeches while adopting combative language to delegitimize him. Such claims resonate with strikingly evaluative headlines such as

(17) Trump is a despicable bigot, a bully, a braggart, a sleazy, thin-skinned egoist – but he is a man (*Irish Independent*: November 8, 2016)

(18) Trump a Washington, tapis rouge et peur bleue<sup>11</sup> (*Libération*: January 21, 2025)

(19) Ein Personal zum Fürchten Regieren Abschieben, Öl bohren, Kohle fördern zählen zu Trumps Programm<sup>12</sup> (*taz*: November 10, 2016)

In these cases, Trump is not merely reported on but enveloped in a lexicon of fear and moral condemnation.

---

<sup>11</sup> Trump in Washington, red carpet and extreme fear

<sup>12</sup> A frightening staff governing deportation, oil drilling, and coal mining are part of Trump’s program

Yet Kuypers' account is not without problems. His empirical base often rests on selective case studies, with little consideration of countervailing evidence. Negative coverage of liberal candidates, for instance, is far from absent. Headlines such as

- (20) 51.3 Million Viewers Tuned In for Shaky Biden and Boisterous Trump (*The New York Times*: June 29, 2024)
- (21) If Trump is overconfident, Harris is deeply insecure (*The Irish Times*: November 5, 2024)
- (22) Clinton no es la mejor candidata en términos de 'marketing' político<sup>13</sup> (*El País*: October 10, 2016)

demonstrate that Democrats, too, are subject to critical evaluations that emphasize weakness, insecurity, or strategic inadequacy. Such examples challenge the claim of a unilateral liberal bias and suggest instead a more complex field in which evaluative language is mobilized against candidates across the partisan spectrum.

Moreover, Kuypers' theoretical posture risks circularity: if negative coverage of a conservative figure is automatically seen as evidence of liberal bias, alternative explanations – such as the candidate's own rhetoric, controversies, or political conduct – are left unexamined. Nor does his framework adequately consider the emergence of right-leaning outlets or hyper-partisan online ecosystems, which complicate any simple dichotomy between a "liberal mainstream" and its conservative challengers. For these reasons, while Kuypers' writings capture a powerful perception within U.S. media debates, they fall short of offering a balanced or comprehensive scholarly account. Instead, they underline the need for approaches that systematically examine evaluative language across contexts, including beyond the American setting.

While Kuypers frames bias as a systemic liberal slant, more research emphasizes that partisanship in news media cuts across the political spectrum. Higdon, Marmol, and Huff (2022) describe the current media landscape as structurally "hyper-partisan," in which outlets of both left and right actively cultivate ideological niches. Similarly, Blevens (2022) shows how liberal and conservative news ecosystems mirror one another in their use of selective sourcing, evaluative language, and interpretive framing. The implication is that bias should not be conceptualized as the exclusive property of one political orientation but as a systemic feature of contemporary journalism, particularly in highly polarized political contexts.

This dynamic is visible in the corpus. Headlines that elevate Donald Trump into a heroic figure, such as

- (23) Give Trump the 'deal maker' a chance to make America great (*Irish Independent*: January 23, 2017)
- (24) U.S Election - Ajadi Predicts Global Stability Under Trump (*Vanguard*: November 6, 2024)
- (25) TRUMP, PRESIDENTE PARA DEVOLVER A EEUU A LA 'ERA DORADA'<sup>14</sup> (*El Mundo*: January 21, 2025)

portray him as a uniquely capable leader, restoring national greatness and economic vitality. These formulations echo the aspirational rhetoric of Trump's own campaign and give it a legitimizing frame. Yet the same outlets, despite their broadly conservative editorial stance, also published headlines that are overtly favorable to Trump's opponents.

---

<sup>13</sup> Clinton isn't the best candidate in terms of political marketing

<sup>14</sup> TRUMP, PRESIDENT TO RETURN THE US TO THE 'GOLDEN AGE'

(26) Mrs Clinton's legacy should be of a y [sic!] woman who raised the bar for all (*Irish Independent*: November 11, 2016)

casts Clinton in terms of historic achievement and inspiration, while

(27) El presidente Joe Biden vino a apagar el incendio<sup>15</sup> (*El Mundo*: January 22, 2021)

depicts Biden as a stabilizing figure, arriving at a moment of crisis to restore calm.

Such juxtapositions illustrate precisely the point made by recent scholarship: bias is not a simple question of a liberal or conservative tilt, but a flexible strategy that outlets deploy depending on context, audience, and editorial priorities. In some cases, conservative papers may celebrate a Republican figure as a national savior; in others, the same outlet may elevate a Democrat as a voice of competence and stability. What unites these examples is not their ideological direction but their evaluative intensity. The role of bias here lies less in consistent partisan alignment than in the willingness to frame political figures through highly positive or negative lenses.

This evidence complicates any search for monolithic bias. Instead, it supports the view that bias operates as an inherent feature of political journalism, where headlines condense and dramatize political realities in ways that resonate with audience expectations. The fact that evaluative framings can point in either direction – sometimes favoring Trump, sometimes his opponents, even within the same newspaper – demonstrates the limits of one-sided critiques such as Kuypers'. At the same time, it underscores the need to examine bias at the level of language itself, where evaluative choices and discursive strategies become visible in their most condensed form.

Taken together, these strands of research suggest that media bias is neither a single measurable imbalance nor a purely partisan attribute. It is a composite phenomenon that comprises (i) quantitative asymmetries in selection and prominence (*gatekeeping, coverage*; D'Alessio and Allen 2000), (ii) qualitative-discursive choices that encode evaluation and ideology (*statement bias*; Simpson, Mayr & Statham 2019), and (iii) system-level constraints that condition what counts as sayable across media environments (Gutsche 2015). The corpus examples indicate that evaluative framings are mobilized across the spectrum – sometimes to elevate Trump, sometimes to diminish him, and likewise for Clinton, Biden, or Harris – aligning with recent accounts of a structurally hyper-partisan media ecology rather than a monolithic tilt. As section 2.1 already underscored, these dynamics are further inflected by differences between media systems, which shape the available discursive repertoire and the range of voices admitted to the page. Against this backdrop, the present study proceeds by focusing on statement bias as it materializes in the linguistic forms of headlines – specifically, how naming choices, evaluative lexis, and framing devices crystallize stance in highly condensed text.

Building on this understanding, the next section turns to names and naming to develop one of the linguistic toolkits – *nomination, categorization, and appraisal* (van Leeuwen 2008) – through which headlines encode stance; this provides one of the analytic categories for tracing statement bias in the corpus.

---

<sup>15</sup> President Joe Biden came to put out the fire

### 2.3. Names and Naming

Naming is never a neutral act. In political reporting, the choice of whether to designate a candidate by title, full name, family name, or even by relational or evaluative labels carries implications that go beyond mere identification. Headlines, with their compressed format, sharpen this effect: the difference between

- (28) Mr. Trump (*Die Welt*: October 10, 2016)
- (29) That Other Clinton (*The New York Times*: January 22, 2017)
- (30) La Clinton<sup>16</sup> (*La Stampa*: November, 7, 2016)
- (31) Trump the felon (*Irish Independent*: January 21, 2025)

is not simply stylistic but positions political actors in divergent ways – through respect, comparison, or condemnation. As Fowler (1991: 4) emphasizes, “there are always different ways of saying the same thing,” and these differences carry ideological distinctions rather than neutral alternatives. Similarly, Taylor (2019: 130) highlights that naming fulfills discourse’s central role in both reflecting and shaping perceptions of social realities, while Würth’s (2016) cross-European study demonstrates that even apparently mundane choices – whether *Merkel*, *Angela Merkel*, or *Angela* – can signal familiarity, distance, or evaluation. In light of this, names and naming emerge as one of the key linguistic toolkits for analyzing statement bias in headlines, and the following section turns to van Leeuwen’s (2008) framework to clarify how such strategies can be systematically categorized. In line with this, nomination is the linguistic construction of social actors (Bendel Larcher 2023: 72). Different nominal choices index different stances: sole surname highlights adult agency and responsibility; titles and functions foreground institutional authority; first-name use can signal informality or familiarity but also infantilize – to a surprising extent this is almost exclusively used with women, often with diminutive effect (Bendel Larcher 2023: 73). These patterned effects dovetail with van Leeuwen’s typology and make nomination a sensitive site of statement bias.

Van Leeuwen’s (2008) framework for representing social actors offers a systematic toolkit for identifying how naming encodes stance. He distinguishes three broad strategies: *nomination*, where actors are referred to by proper names or titles; *categorization*, where they are positioned in terms of roles, attributes, or relations; and *appraisalment*, where they are explicitly evaluated through labels. Table 1 summarizes these categories and illustrates them with corpus examples. Each of these strategies can be realized in multiple subtypes, and even subtle shifts carry discursive implications. As Würth’s (2016) study of European headlines demonstrated, the recurrence of particular forms – such as mononymous surnames or given-name familiarity – can mark authority, distance, or solidarity, while deviations from the norm often carry evaluative weight. As Bendel Larcher (2023: 74–75) stresses, labels activate frames and associative stereotypes (*freedom fighter*, *rebel*, *terrorist*, *insurgent*), shaping expectations about roles and legitimacy. Analogously, political-leader designations such as *statesman*, *populist*, or *reformer* cue value-laden scripts that can tilt reception before any propositional content is read.

This taxonomy makes visible the range of options available to headline writers. The same figure may be presented neutrally (*Trump*), formally (*President Trump*), familiarly (*Donald*), or evaluatively (*Trump the felon*). Such choices crystallize ideological stances in condensed form, foregrounding not only who the actor is but how they are to be perceived.

---

<sup>16</sup> The Clinton

Patterns of *nomination* are particularly revealing. The use of surnames alone, common in most European languages, suggests a default, almost bureaucratic form of reference. By contrast, semiformal full-name constructions such as

(32) Donald John Trump (*The Times*: November 10, 2016)

carry a more marked character: they can signal gravitas, solemnity, or occasionally ridicule, depending on context. Informal first-name references, such as

(33) Kamala (*Reforma*: September 12, 2024)

invite a sense of familiarity and intimacy that reduces distance between reader and politician. Würth (2016) shows that such choices vary systematically across national traditions – British newspapers, for instance, employ given-name familiarity more often than their German or Polish counterparts – yet in all contexts, deviations from the surname norm convey subtle evaluations. Titled forms, whether honorifics such as

(34) Le président Trump<sup>17</sup> (*Le Figaro*: November 7, 2024)

or affiliations in the form of for example

(35) Biden's wife, Jill (*Irish Independent*: October 15, 2020)

foreground institutional roles or relational ties and thereby situate the individual in wider networks of authority, kinship, or legitimacy.

*Categorization* operates differently: here, politicians are represented not by who they are but by what they do, what they embody, or with whom they are associated. *Functionalization* such as

(36) Prosecutor Kamala Harris (*The Guardian*: September 11, 2024)

emphasizes activity and professional background, shaping perceptions of competence or aggressiveness. Classificatory identifiers, as in

(37) Джо Байдън- ветеранът от Белия дом<sup>18</sup> (*Dnevnik*: November 2, 2020)

draw on social categories that highlight experience or age, thereby embedding normative judgements about suitability for office. Relational identifiers (*Biden's wife* or *Obama's vice president*) construct political figures indirectly, defining them in terms of more familiar actors. Physical identifiers, although rare in the present corpus, are equally charged: references such as *the woman in a pantsuit* (for Clinton) foreground visual or gendered markers, often with reductive or stereotypical undertones.

The most overtly ideological strategy is *appraisement*, where evaluation is fused into the name itself. Headlines such as

---

<sup>17</sup> President Trump

<sup>18</sup> Joe Biden – the White House veteran

(38) Trump the felon (*Irish Independent*: January 21, 2025)

dispense with descriptive neutrality and collapse legal status into the candidate's identity. In such cases, the headline no longer merely denotes but prescribes an interpretive frame, turning language into an explicit act of judgement. As Fowler (1991: 4) reminds us, "news is not a value-free reflection of facts," and appraisement exemplifies this principle in condensed, unmistakable form.

Taken together, these strategies demonstrate that naming is both a linguistic and ideological act. To call someone *President Trump*, *Donald*, or *Trump the felon* is not merely to label but to orient the reader

| <i>Strategy</i>       | <i>Subtype</i>                                                   | <i>Definition</i>                                           | <i>Examples</i>                                                                                                                                           |
|-----------------------|------------------------------------------------------------------|-------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------|
| <i>Nomination</i>     | Formal (surname only, ± honorific)                               | Family name alone, with or without honorifics               | <b>Trump</b> amplifies personal attacks on Hillary and Bill Clinton ( <i>The Guardian</i> : October 10, 2016)                                             |
|                       | Semiformal (given name + surname)                                | Full given name + surname                                   | The many faces of <b>Donald John Trump</b> ( <i>The Times</i> : November 10, 2016)                                                                        |
|                       | Informal (given name only)                                       | First name only, often signaling familiarity or informality | <b>Kamala</b> , su nueva era <sup>19</sup> ( <i>Reforma</i> : September 12, 2024)                                                                         |
|                       | Non-proper-name nomination (unique role in context)              | Unique role identifiers substituting names                  | US-Klimaschutz Wie gut sind Obamas Massnahmen gegen die Pläne der <b>Trump Administration</b> ? <sup>20</sup> ( <i>Tages-Anzeiger</i> : January 18, 2017) |
|                       | Titulated nomination – honorification (titles/ranks)             | Use of office or title                                      | <b>Le président Trump</b> rêve de s’autoamnistier <sup>21</sup> ( <i>Le Figaro</i> : November 7, 2024)                                                    |
|                       | Titulated nomination – affiliations (kinship/personal/work ties) | Reference via spouse, relative, or colleague                | Democratic Party takes aim at Texas with ‘secret weapon’: <b>Biden’s wife, Jill</b> ( <i>Irish Independent</i> : October 15, 2020)                        |
| <i>Categorization</i> | Functionalization (what they do/roles & activity nouns)          | Actor described through activity or role                    | <b>Prosecutor Kamala Harris</b> put Trump on trial, but the court of public opinion can be fickle ( <i>The Guardian</i> : September 11, 2024)             |
|                       | Identification – classification (social categories)              | Actor defined by social category (e.g. class, ethnicity)    | Джо Байдън- ветеранът от Белия дом <sup>22</sup> ( <i>Dnevnik</i> : November 2, 2020)                                                                     |
|                       | Identification – relational (defined by relations)               | Defined through relation to others                          | <b>Trump-Ally</b> Pleads Guilty in Foreign Lobbying Case Tied to Malaysia and China ( <i>The New York Times</i> : October 21, 2020)                       |
|                       | Identification – physical (salient physical traits)              | Defined by physical features                                | (no example in corpus; e.g. <i>the woman in a pantsuit</i> for Clinton)                                                                                   |
| <i>Appraisal</i>      | Evaluative labels found in discourse                             | Explicit value-laden naming                                 | Inauguration of <b>Trump the felon</b> proves there is no such thing as fair play ( <i>Irish Independent</i> : January 21, 2025)                          |

Table 1 Naming strategies according to van Leeuwen (2008) with examples

<sup>19</sup> Kamala, her new era

<sup>20</sup> US climate protection How effective are Obama’s measures against the Trump administration’s plans?

<sup>21</sup> President Trump dreams of granting himself amnesty

<sup>22</sup> Joe Biden – the White House veteran

toward a stance, whether of respect, intimacy, or condemnation. Naming thus belongs squarely to the domain of statement bias: it is a subtle yet pervasive way in which headlines encode political evaluation.

Building on this foundation, the next section turns to political communication more broadly, situating naming practices within the larger dynamics of framing, agenda-setting, and audience perception that structure how media mediate politics.

## 2.4. Political Communication

Political communication refers to the complex interplay between political actors, the media, and the public in the production, mediation, and reception of political messages. Far from being a transparent channel, communication is shaped at each stage by institutional routines, cultural expectations, and discursive strategies. Newspapers and their headlines form a central element in this process: they not only report on political developments but also orient readers toward particular ways of interpreting them. As such, headlines function not merely as informational devices but as nodal points in the broader circulation of political meaning.

One of the most influential perspectives in this field is *agenda-setting theory*. Weaver and Choi (2014) emphasize that the central claim – that media do not dictate what people think, but strongly influence what people think *about* – retracts considerable empirical support. Applied to election coverage, this means that when newspapers foreground immigration, foreign policy, or a candidate's character in their headlines, they are structuring the public agenda by signaling which issues deserve attention. Crucially, *agenda-setting* is not confined to individual outlets but operates across media systems. As Weaver and Choi note, “the changes in one medium's agenda precede or follow changes in another's” because “journalists tend to closely follow their colleagues' news stories” (2014: 363). This intermedia dynamic is reinforced by the major wire services, such as the *Associated Press*, which exert “important intermedia agenda-setting influence” (2014: 364). The result is that particular topics or frames can cascade rapidly across outlets, shaping the horizon of political debate well beyond their point of origin.

Closely related to *agenda-setting* is the concept of *framing*, which highlights not the selection of issues but the interpretive lenses through which they are presented. As Scheufele and Iyengar (2014: 1) argue, *framing* “refers to subtle alterations in the statement or presentation of judgement and choice problems,” which can profoundly shift how audiences evaluate the same information. Whereas *agenda-setting* tells us *what* to think about, framing shapes *how* those issues are understood. Frames operate by emphasizing certain aspects of a story while omitting others, thereby directing interpretation without explicit instruction. They are especially powerful in ambiguous or contested contexts, where multiple plausible narratives compete for dominance (Scheufele & Iyengar 2014: 2). In headlines, this becomes immediately visible: the difference between describing a debate as *victory*, a *stalemate*, or a *missed opportunity* is not simply stylistic variation, but a discursive cue guiding readers toward a particular evaluative stance. Complementing this, newsroom notions of impartiality and objectivity set different normative horizons across media types. Impartiality is a broadcast convention of even-handedness in controversy, whereas newspapers are “often partisan, and widely assumed to be so,” with objectivity serving as a more general press norm (Montgomery 2019: 59, 88). These institutional expectations shape how far evaluative language in headlines is licensed within a given media system.

A third strand of research centers on the *hostile media effect*, which addresses how news is received rather than how it is produced. Feldman (2014: 151) defines the phenomenon as the tendency of partisans to “perceive neutral news coverage as biased against their side,” even when systematic content analyses

show no such imbalance. The effect has been documented across a wide range of political issues and media contexts, and it implies that the reception of headlines is shaped as much by audience predispositions as by journalistic choices. In highly polarized environments such as U.S. presidential elections, this means that a headline aiming at descriptive neutrality may be interpreted by one group as unfairly critical and by another as unduly favorable. The *hostile media effect* thus complicates the evaluation of bias: it highlights the gap between *perceived* and *demonstrable* imbalance, showing that even the most careful wording cannot escape partisan interpretations.

A fourth strand of scholarship emphasizes the structural and ideological dimensions of political communication. Lichter (2014: 403) observes that media bias is more often invoked as a hypothesis to explain patterns of news coverage than elaborated as a fully developed theory of communication. The term itself is used flexibly, referring variously to distortions of reality, favoritism in the presentation of controversies, or partisan attitudes (Lichter 2014: 404). Within this broad usage, critical approaches situate bias within traditions of ideology critique, treating news as an “ideological product that shapes mass consciousness in a manner that preserves the hegemony of society’s ruling interests” (Lichter 2014: 406). Here, ownership structures and corporate control are viewed as decisive in shaping what is reported and how. By contrast, other perspectives emphasize the ideologies of journalists themselves, a line of argument often associated with popular right-wing critiques of liberal-leaning newsrooms (Lichter 2014: 407).

In a similar vein, Gutsche (2015) underscores the role of media as instruments of social control, highlighting the “common practices of news construction, institutions, and representations” that justify and reinforce dominant values (Gutsche 2015: 7). Central to this account is *gatekeeping*, in which information passes through “gates” managed by journalists and institutions that decide which stories are admitted and which are excluded (2015: 114). From this perspective, journalism is not primarily about reporting reality but about constructing power: “news is power construction that follows methods of cultural indoctrination” (Gutsche 2015: 115). This view is compatible with *agenda-setting* research, which shows how elite actors – including businesses, politicians, and dominant media organizations – shape the flow of information from national outlets to regional ones, funneling what publics should talk and think about (Gutsche 2015: 116). In such approaches, bias is not merely an unintended by-product of individual journalists’ perspectives but a structural feature of the media system itself.

These dynamics become particularly visible in headlines, where *agenda-setting*, *framing*, and perceptions of bias intersect in condensed textual form. The *agenda-setting* function of newspapers can be observed in the way candidates are linked to specific issues.

(39) Подкрепа за добива на петрол и нови мита: Какво ще направи Доналд Тръмп, ако стане президент<sup>23</sup> (*Dnevnik*: November 5, 2024)

positions Trump primarily as a policymaker in the domains of oil drilling and tariffs, implicitly suggesting that these are the areas through which his presidency should be evaluated. Similarly,

(40) An Essential Question: Would Biden’s Tax Plan Help a Weak Economy? (*The New York Times*: October 19, 2020)

---

<sup>23</sup> Support for oil production and new tariffs: What Donald Trump will do if he becomes president

frames Biden's candidacy through the lens of economic policy, raising a pointed question about fiscal competence. Headlines such as

(41) Biden lanza su ofensiva para paliar la crisis<sup>24</sup> (*El Mundo*: January 23, 2021)

likewise highlight the economy as the salient sphere of political action, making readers attend to crisis management as the defining criterion for political leadership. In each case, headlines do not dictate judgments, but they determine which issues are foregrounded and thus structure what is publicly debated.

*Framing* emerges just as clearly in headlines that interpret rather than merely announce political events, as in the following examples:

(42) Donald l'anti-migranti o Clinton la liberal I cattolici d' America restano senza scelta<sup>25</sup> (*La Stampa*: October 20, 2016)

(43) Trump sceglie la rissa e va a caccia dei voti dei bianchi<sup>26</sup> (*La Stampa*: October 11, 2016)

(44) Trump ante portas - Bangen und Hoffnung<sup>27</sup> (*Neue Zürcher Zeitung*: January 20, 2025)

The latter evokes classical imagery of an invader at the gates, encapsulating a dual frame of fear and hope that positions Trump as both a threat and a source of possibility. The headline

(45) Słynny dziennikarz George Packer: Stary Biden będzie potrzebował siły Herkulesa<sup>28</sup> (*Gazeta Wyborcza*: October 31, 2020)

constructs Biden's candidacy through metaphorical framing, simultaneously emphasizing age and invoking mythological strength as a counterbalance. Likewise,

(46) Morning Mail: Trump triumphant as Harris concedes (*The Guardian*: November 6, 2024)

interprets the election outcome not simply as a tally of votes but as a narrative of triumph and defeat. These cases exemplify Scheufele and Iyengar's (2014) observation that frames shape how events are understood by highlighting one interpretive possibility among many.

The *hostile media effect* underscores that even seemingly neutral headlines are subject to divergent partisan readings.

(47) Reportan arresto de un hombre que amenazaba con matar a Biden<sup>29</sup> (*El Universal*, October 22, 2020)

reports an arrest in neutral, descriptive terms, yet it could be received as sensationalist by critics or as evidence of danger and victimhood by sympathizers. Similarly,

---

<sup>24</sup> Biden launches his offensive to alleviate the crisis

<sup>25</sup> Donald the anti-immigrant or Clinton the liberal: American Catholics are left with no choice

<sup>26</sup> Trump opts for confrontation and goes after the white vote

<sup>27</sup> Trump ante portas - Anxiety and hope

<sup>28</sup> Famous journalist George Packer: Old Biden will need the strength of Hercules

<sup>29</sup> Reports of arrest of man who threatened to kill Biden

- (48) The Health Information Voters Need From Biden and Trump (*The New York Times*: June 28, 2024)

appears to provide neutral voter guidance, but partisan readers may see bias in what is emphasized or omitted. Even

(49) Los procesos que inquietan a Biden y Trump<sup>30</sup> (*El Mundo*: January 20, 2021)  
which treats both candidates in parallel terms, may be read as downplaying one figure's culpability while exaggerating the other's. These examples illustrate Feldman's (2014) point that perceptions of bias often diverge from demonstrable content, highlighting the gap between intention and reception.

Finally, headlines also reflect the structural and ideological dimensions highlighted by critical approaches.

- (50) Trump, un mandat pour secouer le système<sup>31</sup> (*Le Figaro*: November 5, 2016)

does more than announce a victory: it constructs Trump as a figure with a legitimate mandate to disrupt established order, thereby reinforcing the symbolic power of electoral success. Similarly,

- (51) Morning Mail: Trump triumphant as Harris concedes (*The Guardian*: November 6, 2024)

frames the event as a triumph, imbuing it with grandeur and finality that extend beyond the factual concession. As Lichter (2014) and Gutsche (2015) argue, such discursive choices reflect not only the perspectives of journalists but also the structural pressures of media systems that incentivize dramatization and ideological reproduction.

Taken together, these perspectives illuminate complementary but also limited dimensions of political communication. Agenda-setting research has demonstrated the media's capacity to direct public attention, but it often remains confined to measuring issue salience in quantitative terms, without accounting for the evaluative force of language. Framing theory offers a powerful account of how interpretive lenses guide audience perceptions, yet as Scheufele and Iyengar (2014) themselves acknowledge, it risks tautology: virtually any form of presentation can be described as *frame*, which blurs analytical boundaries. The *hostile media effect* underscores the importance of reception, showing that partisanship powerfully colors how coverage is perceived. Yet this line of work, while crucial in highlighting the perception-reality gap, does little to explain the linguistic mechanisms that produce such perceptions in the first place. Critical approaches, by contrast, highlight the ideological and structural conditions under which news is produced, but they often portray media systems in overly deterministic terms, underestimating the agency of journalists and the variability of discourse across outlets and contexts.

What these traditions share is an acknowledgement that political communication is never neutral: it involves choices about what to report, how to present it, and how it will likely be received. At the same time, each approach leaves unaddressed the micro-level processes through which bias manifests linguistically in individual texts. This is the space in which the present study intervenes. By examining headlines through the lens of discourse analysis, and specifically through the categories of naming strategies and evaluative language, the analysis complements system-level and perception-level accounts with a fine-

---

<sup>30</sup> The trials that concern Biden and Trump

<sup>31</sup> Trump, a mandate to shake up the system

grained investigation of the textual choices that encode bias. Headlines crystallize the broader dynamics described by agenda-setting, framing, and ideological critique, but they do so in miniature, allowing the mechanisms of bias to be traced in their most condensed and visible form.

These theoretical insights provide the backdrop for understanding not only how media shape political discourse in general, but also how individual candidates come to be positioned within it. Few figures illustrate the tensions between *agenda-setting*, *framing*, reception, and ideological critique as vividly as Donald J. Trump. His confrontational style, coupled with his deliberate attacks on the press, transformed the relationship between candidate and media into a central feature of the campaigns themselves. The next section therefore turns to Trump and the Media, examining how this dynamic reconfigured journalistic practices and public debates.

## 2.5. Trump and Media

From the outset of his 2016 campaign, Donald Trump has been establishing a combative relationship with the press, repeatedly accusing legacy media outlets of unfair coverage. His most enduring rhetorical innovations – branding critical outlets as *fake news* and journalists as *enemies of the people* – exemplify what Happer, Hoskins, and Merrin (2019) describe as a media war, in which antagonism toward the press is not incidental but central to Trump’s political communication. The delegitimizing of his rhetoric is evident: by questioning the credibility of mainstream journalism, Trump positions negative reporting as inherently untrustworthy, thereby shielding himself from accountability. Using his direct channels (especially Twitter/X) enables him to bypass “established media institutions” and recast journalists as protagonists within the story (Enli 2017: 53). The same posts then re-enter legacy media as quotable content, intensifying the cycle of contention. At the same time, these attacks have a mobilizing effect, offering his supporters a shared adversary in the figure of the media, whose alleged hostility could be turned into a rallying cry at events.

The consistency of this framing is striking. As Kellner (2019: 47) observes, “Trump’s media bashing and daily attacks on the media via his campaign rallies, Twitter feeds, and off-the-cuff remarks have been a defining feature of both Trump’s presidential campaign and the first 200 days of his presidency.” These attacks are woven into the fabric of Trump’s communication style. By repeatedly foregrounding media bias, Trump has transformed distrust of journalism into a marker of political allegiance. In this way, antagonism toward the press serves not only to delegitimize critical coverage but also to mobilize supporters through a shared sense of grievance against a hostile establishment.

A second defining aspect of Trump’s media strategy is his extensive use of X (formerly Twitter), and more recently Truth Social, as a direct channel of communication. Unlike press releases or official statements, which pass through aides and institutional filters, many of Trump’s tweets are composed and published in his own voice. This gives them an immediacy and authenticity that resonates with supporters, while simultaneously allowing him to bypass traditional journalistic gatekeepers. As Enli (2017: 53) explains, Twitter functioned during the 2016 campaign as a direct source of news, enabling Trump to circumvent the editorial role of mainstream media and to position himself in a running battle with journalists. The platform thus offered Trump a uniquely unmediated means of framing events, attacking opponents, and commenting on coverage in real time.

Trump’s reliance on Twitter was – and remains – not ancillary but central to his political persona. His self-authored posts provide a constant flow of commentary that shapes the news agenda: tweets are quoted in headlines, discussed on television, and amplified across digital networks. In this way, Trump

has been weaponizing social media not only to deliver unfiltered messages but also to force legacy media to cover his statements on his own terms. The result is a paradoxical relationship: even while railing against the press, Trump guarantees that his words dominate the news cycle, ensuring that he remains at the center of global attention.

At the heart of Trump's media strategy lies the recurring accusation that legacy outlets are both "fake news" and systematically negative in their reporting on him. These two motifs are closely intertwined: "fake news" functions as a delegitimizing label, while the claim of relentless negativity personalizes the grievance and frames coverage as a partisan attack. As Meeks (2020: 5) observes, Trump's rhetoric regularly combines these elements in way that not only undermine the credibility of established outlets but also casts them as hostile actors in an "us-versus-them" struggle. By branding the press as "failing" and even as the "enemy of the American People," Trump encourages supporters to distrust critical reporting, effectively inoculating himself against scrutiny.

Trump's own words illustrate the consistency of this strategy. On February 17, 2017, he tweeted: "The FAKE NEWS media (failing @nytimes, @NBCNews, @ABC, @CBS, @CNN) is not my enemy, it is the enemy of the American People!" (Trump 2017). Earlier, he had singled out *The New York Times*, declaring: "The failing @nytimes is truly one of the worst newspapers. They knowingly write lies and never even call to fact check. Really bad people!" (Trump 2016). Such posts, unmediated by aides or institutional filters, exemplify how Trump fused delegitimization with a personalized narrative of victimhood, turning attacks on journalism into proof of media hostility. The sheer frequency of these accusations ensured that "fake news" became a central motif of his political identity, one that blurred the line between criticism of coverage and rejection of journalism's normative role.

The paradox of Trump's relationship with the press is that even as he denounces mainstream outlets as biased and untrustworthy, he depends on them to sustain his visibility. His constant presence in headlines and broadcasts ensures that he remains at the center of global attention, regardless of whether the coverage is favorable or critical. As Morini (2020) emphasizes, Trump's communication strategy is based on dominating the media environment, a process in which controversy itself becomes a resource for commanding attention. By generating a steady flow of provocative statements – often through tweets that are immediately picked up by journalists – Trump guarantees that his perspective shapes the news agenda.

This dynamic positions the media simultaneously as amplifier and adversary. On the one hand, legacy outlets relay Trump's statements widely, often reproducing his phrases verbatim in headlines or lower thirds, thereby amplifying the very rhetoric that disparaged them. On the other hand, their critical coverage supplies the evidence for Trump's claims of hostility, reinforcing the narrative of persecution that energizes his base. Several scholars stress the mutually reinforcing yet adversarial character of Trump's media conflict: Symons (2019: 183) argues that Trump's interactions with television comedians shifted "from symbiotic coexistence to an effective state of war," while Jarvis (2019) characterizes the broader Trump-press dynamic as a "murder-suicide pact." In other words, media organizations and Trump sustain each other's visibility by engaging in a constant cycle of attack and counterattack. The effect is that Trump not only survived but thrived in a hostile media environment, turning negative coverage into proof of his centrality and resilience.

## 2.6. The Rationale for Analyzing Trump and Headlines

Here, the threads from the preceding sections converge into the precise rationale for the present study. Section 2.1 established why headlines are a privileged unit of analysis: they are the first and often only

contact point between readers and a story, and their condensed form makes lexical choices and evaluative cues unusually visible. Section 2.2 then clarified that media bias is not a single variable but a bundle of mechanisms – *gatekeeping*, *coverage*, and especially *statement bias* – through which news can tilt the representation of political actors. Section 2.3 demonstrated that one of the most sensitive conduits of statement bias is naming: whether a candidate is designated by title, surname, relational tie, or evaluative label is never merely referential but inherently stance-bearing. Section 0 situated these textual choices within broader theories of political communication – *agenda-setting*, *framing*, *hostile media* perceptions, and structural constraints – reminding us that headlines both reflect and shape the agenda, and that audience predispositions modulate how any given formulation is received.

Against that theoretical backdrop, section 2.5 identified a distinctive feature of the Trump era: a sustained, high-salience metadiscourse about the press itself. Trump’s repeated characterizations of legacy outlets as “fake news” and “enemies of the people,” most prominently on Twitter/X, do two things at once. First, they pre-emptively delegitimize critical reporting by portraying it as inherently untrustworthy. Second, they personalize negativity claims – asserting that the press is not just unreliable in general, but specifically and systematically hostile to him. This dual move is exactly what earlier sections equip us to examine empirically. If headlines are where evaluative language is distilled (2.1), if bias includes how – not only whether – issues and actors are expressed (2.2), and if naming choices crystallize stance into micro-form (2.3), then headlines are the most suitable site to test whether Trump’s allegation of relentless negativity is borne out in practice. And because agenda-setting and intermedia dynamics (0) can rapidly propagate frames across outlets and borders, a cross-national, Eurolinguistic comparison can reveal whether putative hostility is a U.S.-specific pattern, a transatlantic echo, or a more variegated, system-dependent phenomenon. At the same time, the bounds of discourse analysis apply. No study can capture a discourse in its entirety; many fragments are lost or practically unprocessable (Bendel Larcher 2023: 59). The claims advanced here are therefore explicitly limited to accessible headline corpora in defined windows, with inferences calibrated accordingly.

This logic yields a focused, falsifiable research task. The claim to be scrutinized is not whether Trump was controversial or frequently in the news – those are background conditions – but whether headlines systematically encoded negative stance toward him relative to contemporaneous competitors (Clinton, Biden, Harris) and whether that negativity displayed stable patterns across different media systems. This study therefore concentrates on statement bias as realized in headline language, with two principal levers: (i) naming strategies (*nomination*, *categorization*, *appraisement*) that cue respect, familiarity, distance, or condemnation; and (ii) evaluative and framing lexis that push readers toward particular judgements before they encounter the body text. Because headlines are standardized, widely archived, and structurally comparable across languages, they enable both within-country contrasts (between outlets of different ideological leanings) and between-country contrasts (across European and non-European press). The Eurolinguistic lens is not an afterthought; it addresses whether there are shared European tendencies in naming and evaluation – or whether national traditions, political cultures, and media-system constraints yield distinct patterns.

At the same time, the earlier sections caution against overreading any single lens. Agenda-setting reminds us that high visibility need not equal hostility; framing theory warns that multiple readings are often available and that ambiguity can be dispositive; hostile media research indicates that perceptions of bias can outstrip textual evidence; and critical accounts of media systems suggest that ownership, editorial lines, and commercial incentives shape the discursive space in which headline writers work. The present design therefore does not adjudicate factual truth claims (as section 2.7 will emphasize), nor does it attempt to infer intent from outcomes. Instead, it restricts itself to what can be shown directly in text:

how frequently and in what ways headlines encode stance toward each candidate through naming and evaluative choices, and how these distributions vary by outlet, language, and electoral cycle.

Bringing the argument back to Trump’s own metadiscourse, this approach has two payoffs. Substantively, it addresses the core of his allegation – persistent, exceptional negativity – at the precise discursive location where such an allegation would be most legible. Methodologically, it complements past work that has either aggregated imbalances (e.g. coverage time) or theorized bias at the level of institutions and perceptions by zooming in on the linguistic micro-mechanisms that make a headline read as respectful, skeptical, or condemnatory. In doing so, it builds a bridge between quantitative traditions and discourse-analytic traditions, using a coding scheme (see section 5) anchored in van Leeuwen’s typology as well as statement bias theory to keep categories transparent and replicable.

The cross-national scope further strengthens the rationale. As section 2.1 noted, newspapers share enough structural commonalities to permit comparison, but media systems differ in press-party alignments, norms of personalization, and tolerance for overt appraisal. If Trump’s charges of hostility are correct in a strong sense, we would expect similar negative naming and evaluative patterns across a wide range of outlets and countries. If the patterns instead vary – by language, outlet ideology, or electoral moment – then the data will point to a more contingent landscape, in line with the hyper-partisan and heterogeneous picture sketched in section 2.2. Either outcome is informative: convergence would corroborate claims of broad hostility; divergence would suggest that “fake news” rhetoric collapses diverse practices into a single accusatory frame.

Finally, this bridge sets up this study’s stance and safeguards. Section 2.7 immediately following articulates positionality and the steps taken to prevent values from steering results: preregistration, fixed label definitions, and transparent prompt design. That declaration is not ancillary; it operationalizes the normative commitments stated earlier – namely, that claims about media bias should be auditable, modest in scope, and calibrated to what the data can support. In short, the background has established (i) why headlines matter, (ii) what “bias” can mean at the level of language; (iii) how naming functions as a core vehicle of statement bias, (iv) how political communication dynamics and audience perceptions complicate interpretation, and (v) why Trump’s persistent accusations make this inquiry timely and relevant. The analysis that follows takes these premises seriously by testing them where they are most visible: in the naming strategies and evaluative language of international newspaper headlines about U.S. presidential candidates.

## 2.7. Personal Bias and Motivation

This section responds directly to Sylvia Bendel Larcher’s (2023) argument that a completely objective, ideology-free science is impossible, and that research ceases to be legitimate unless its political, ethical, or religious vantage point is declared and accompanied by systematic self-critique from the formulation of the research question to the conclusions (Bendel Larcher 2023: 45, 255). For Bendel Larcher, discourse analysis carries a particular responsibility: it must reject the fiction of neutrality, reflect on its own choices of topic, categories, and interpretations, and justify its societal relevance not by description alone but by *Aufklärung* – clarifying the discursive entanglements that shape public understanding (Bendel Larcher 2023: 255–256).

I adopt this stance. The following pages outline (i) my motivation for the study, (ii) my normative standpoint and prior expectations, and (iii) the safeguards I have built into the design to ensure that my values

do not silently steer the analysis. This declaration is not an apology for bias; it is an intentional design feature that renders the study auditable.

I write as a European observer of U.S. presidential campaigns whose consequences extend well beyond U.S. borders. My interest crystallized during the 2024 Biden/Harris-Trump contest, when debates over “fake news” and media partisanship once again became highly salient. Having previously analyzed newspaper headlines on other topics (Holler 2023), I was struck by the intensity with which Donald J. Trump framed the press as biased against him and sympathetic to his opponents. Presidential elections are particularly valuable for analysis because competing candidates are simultaneously salient and often co-framed – for example, in debate coverage – making direct contrasts possible. My international background—residences in Germany, the United States, Italy, China, and Egypt—has shaped a plural perspective. It encourages skepticism towards single-country “common sense”, heightens sensitivity to translation issues, and underscores the fact that elections vary widely in democratic quality. These experiences inform the study’s Eurolinguistic approach and its comparative focus on global media discourse across the 2016, 2020, and 2024 races. I approach this research from a clearly articulated normative standpoint. I regard constitutional democracy and the rule of law as superior constraints on raw majorities: durable rights and institutional checks should temper fleeting partisan will. I am strongly opposed to illiberal campaign tactics, such as intimidating officials or journalists or calling for opponents’ imprisonment without due process, and I consider rhetoric that labels adversaries or the press as “enemies” unacceptable within democratic competition. In terms of media practice, I hold a baseline trust in major European, UK, and U.S. outlets, though I do not assume uniform bias: my expectation is varied by outlet and national context. I value above all accuracy and verification, transparency about sourcing and methods, and accountability through corrections. “Both-sides” framing I consider appropriate only where evidence and expert consensus are genuinely divided, since forced symmetry can otherwise obscure real asymmetries.

Because discourse analysis cannot – and should not – pretend to be value-free, I also declare my priors regarding the objects of analysis. I expect heterogeneous behavior across outlets – by country, format, and language – rather than a single left-right pattern. My baseline attitude toward Donald J. Trump as a political actor is negative, while my priors for Hillary Clinton, Joe Biden, and Kamala Harris are mixed. I acknowledge, however, that such attitudes may reflect disproportionately the most audible segments of public debate rather than a balanced exposure to coverage. Linguistically, I expect naming choices (for example, *President Trump* versus *Trump*), metaphorical frames of conflict (*battle*, *attack*), and headline architecture to carry much of the framing weight. Equally important are the boundaries of this study. It is not a fact-checking project: I do not adjudicate the truth of the claims reported in the headlines, nor do I supply systematic corrective context for contested statements. The object of analysis is the form and polarity of news language, not the verification of claims. Contextual details are added only when required to avoid misclassifying of discourse function – for example, distinguishing between an allegation and a verdict when the headline itself sets up that contrast. Such additions serve categorization only, not evaluation.

To guard against my values shaping the outcomes, I delegate the coding of headlines to an AI model. Before any modeling is conducted, I fix the analysis plan in detail: prompt wording, label definitions, model and decoding parameters, and corpus inclusion and exclusion criteria. Prompts and parameters are finalized on the training subset and validated by a human coder; the preregistered plan is then applied once to the unseen corpus. To enable audit, I publish the full prompt templates, parameter settings, and a sample of headlines with predicted labels. Prompts are deliberately minimal and non-reactive, avoiding ideological or outlet-specific cues, so that my priors cannot enter indirectly. Corpus selection is likewise

fixed ex ante: outlets are chosen on the basis of reach, language, and ideological reputation, and time windows are anchored to predefined campaign events.

This project does not pursue advocacy goals. Its aim is descriptive and explanatory clarity in analyzing discursive patterns, not the promotion of political outcomes. At the same time, consistent with Bendel Larcher's argument that the humanities are most relevant when they contribute to *Aufklärung*, I take the societal relevance of this study to lie in clarifying how linguistic choices shape perceptions of political actors during high-stakes elections. Several limitations follow from these commitments. The findings address framing and relative polarity, not the factual accuracy of political claims. Results depend on the chosen models and prompts; while stability checks and sensitivity analyses reduce model dependence, they cannot eliminate it. Finally, the focus on presidential campaigns privileges periods of intense candidate co-framing and may not capture the dynamics of routine governance coverage. These are not afterthoughts but deliberate boundaries, which delineate the claims the study can and cannot make, and thereby uphold its accountability to readers.

Having clarified these commitments, the background chapter now turns from conceptual foundations and reflexive positioning to the concrete research task. The final section formulates the study's hypotheses: specific, testable expectations about how statement bias manifests in headline language across candidates, outlets, and national contexts. These hypotheses operationalize the rationale developed so far and provide the framework against which the empirical analysis will be conducted.

## 2.8. Hypotheses

The preceding sections have established the rationale for investigating bias in international newspaper headlines about U.S. presidential candidates. Headlines crystallize evaluative choices in condensed form (section 2.1), bias is expressed not only through selection but through statement-level language (section 2.2), naming constitutes a sensitive discursive site for stance (section 2.3), and political communication theory explains how *agenda-setting*, *framing*, and audience perceptions complicate interpretation (section 0). Against this backdrop, Donald J. Trump's repeated denunciations of the press as "fake news" and "enemies of the people" (section 2.5) provide both the justification and the focal point of this study. His allegation of systemic hostility requires empirical scrutiny: do headlines, across languages and national contexts, systematically encode negative stance toward Trump relative to his opponents?

This study advances the following hypotheses:

### H1. Distribution of bias.

Headlines are more often evaluative than neutral. Specifically, the combined share of positive and negative headlines outweighs the neutral category. This reflects Trump's own framing of the media as adversarial: if he perceived relentless negativity, then a baseline assumption is that evaluative stance (whether hostile or laudatory) dominates headline reporting.

### H2. Candidate-specific bias.

Across the international corpus, Hillary Clinton, Joe Biden, and Kamala Harris receive a more positive balance of coverage than Donald J. Trump, who is disproportionately framed in negative terms. This hypothesis is directly aligned with Trump's claim that mainstream journalism is structurally biased against him.

### H3. Bias and ideological stance of newspapers.

The orientation of bias in headlines correlates with the ideological leaning of outlets. Right-leaning newspapers are more likely to frame Trump positively and Democratic candidates negatively, whereas left-leaning outlets display the reverse tendency. As Higdon, Marmol, and Huff (2022: 265) note, legacy media tend to reproduce their audiences' ideological expectations, presenting partisan actors either as "heroes or archvillains."

### H4. Stability across electoral cycles.

The distribution of bias remains consistent across the 2016, 2020, and 2024 campaigns. If Trump's characterization of legacy media as persistently "bad" were accurate, negative bias would not be episodic but sustained across cycles.

### H5. Naming practices.

Legacy media adhere primarily to neutral naming conventions (surname only or first name + surname) rather than overtly evaluative forms (nicknames, epithets, or derogatory designations). This builds on van Leeuwen's (2008) typology of naming strategies and Würth's (2016) findings that newspapers tend to rely on conventional nomination practices.

### H6. Value-judgment labels.

When evaluative labels are attached to candidates' names, they occur more frequently with Donald Trump than with his Democratic counterparts. This hypothesis echoes Kuypers' (2014: 5) observation that value-laden descriptors are disproportionately applied to Republicans.

To integrate the explicitly Eurolinguistic dimension of the study, the following hypotheses extend the framework:

### H7. European unity of bias.

European newspapers display a shared discursive pattern of negativity toward Donald Trump, suggesting a common transnational orientation. This draws on Eurolinguistic research (Grzega 2025; Giessen & Szwed 2025) which emphasizes the identification of pan-European discursive features.

### H8. Europe versus the rest of the world.

European outlets frame Trump more negatively than newspapers from non-European regions, reflecting structural and cultural differences in political communication.

### H9. Cross-linguistic convergence.

Across Europe's major linguistic families (Germanic, Romance, Slavic, etc.), newspapers exhibit comparable patterns of bias in reporting on U.S. presidential candidates. This reflects a shared European discursive orientation rather than language-family-specific traditions.

#### H10. Regional consistency.

Despite cultural and geopolitical diversity, European newspapers converge on a broadly similar evaluative stance toward Trump and his opponents. Bias does not systematically differ across Northern, Southern, Western, Eastern, or Central Europe, indicating pan-European alignment in headline discourse.

#### H11. European ideological distinctiveness.

Within Europe, the ideological stance of newspapers (left-leaning vs. right-leaning) has less impact on their portrayal of Donald Trump than in the North American media context. This suggests a shared European distance from Trump that overrides traditional partisan divides.

#### H12. Unified naming conventions.

European newspapers converge on neutral naming practices (surname, or first + surname), with minimal use of evaluative nicknames or derogatory labels. This uniformity reflects a European journalistic convention distinct from U.S. patterns.

### 3. Literature and Research

This chapter surveys the scholarship that anchors and motivates the thesis, moving from the conceptual foundations of Eurolinguistics to recent advances in AI-assisted discourse analysis and, finally, to media-bias research on Donald J. Trump. The aim is twofold: first, to clarify what counts as a credible, Europe-wide comparison of discourse – why newspapers (and especially headlines) are a productive site for such work and how prior studies have operationalized European commonalities and divergences; second, to position the present design against the most influential accounts of Trump’s coverage across regions, outlets, and genres, drawing out where existing findings align with – and where they challenge – the hypotheses tested here.

The discussion begins by establishing the Eurolinguistic frame. It synthesizes recent work that problematizes definitions of *Europe*, *European language*, and *European feature*, and then narrows to discourse-analytic approaches that treat newspapers as revealing intersections of shared pragmatic tendencies and national specificities. From there, the chapter turns to methodological developments in discourse annotation with large language models. It reviews what current evaluations say about baseline capability, task adaptation, and reproducibility controls, and it extracts best practices for hybrid AI-human workflows suited to large, multilingual headline corpora. The third block consolidates comparative findings on Trump’s media portrayal within and beyond Europe, emphasizing patterns that bear directly on the thesis hypotheses: the distribution and direction of evaluative tone, outlet ideology effects, temporal stability across election cycles, naming practices, and value-judgement labels. The chapter closes by identifying the research gap that motivates the study’s design: the need for a standardized, multilingual, headline-level comparison aligned to identical event windows across 2016/2020/2024, integrating neutrality alongside positive and negative tone and leveraging scalable, controlled LLM coding with targeted human oversight.

Read in sequence, these sections provide a continuous argumentative arc: Eurolinguistics supplies the comparative rationale and sampling discipline; AI methods supply scalable, transparent tools for coding;

the Trump/media-bias literature supplies the empirical expectations the study seeks to confirm, qualify, or contest. Together, they motivate the thesis's core contribution: a cross-European, cross-cycle analysis of headline discourse that tests clearly specified hypotheses about bias, ideology, naming, and regional convergence with methodological rigor and comparative breadth.

### 3.1. Eurolinguistics

Eurolinguistics has emerged as a field concerned with identifying and analyzing linguistic commonalities across Europe. As Grzega (2025: 3) notes, "Eurolinguistics is the study of commonalities among European languages." This broad definition highlights the comparative ambition of the discipline, which ranges from structural linguistics to sociolinguistics and pragmatic perspectives. Yet the apparent simplicity of the term conceals considerable conceptual and methodological challenges. What counts as *Europe*, which languages qualify as *European*, and how one should define a *European feature* are all questions without stable answers. The openness of the field is therefore both an asset – allowing for wide-ranging comparative investigations – and a difficulty, as the absence of consensus can limit the precision of Eurolinguistic claims.

One source of conceptual instability is the very definition of *Europe*. As Grzega (2025: 3–4) points out, the continent can be approached geographically, politically, or culturally, and none of these perspectives provides a definitive boundary. Geographical accounts typically delimit *Europe* between the Atlantic and the Ural Mountains, though even this naturalized definition occasionally excludes the British Isles or Iceland. Political definitions link *Europe* to institutions such as the European Union or the Council of Europe, but these are shifting entities, whose memberships can vary over time. Cultural-anthropological definitions are even less clear-cut, relying on shared but unevenly distributed traditions in philosophy, religion, or the arts, many of which extend beyond Europe itself (Grzega 2025: 4). As Grzega concludes, "*European* may not even form a separate category of its own, but only a subcategory of an entity in some cases call *western*" (2025: 5, emphasis in the original). This definitional fluidity does not invalidate the project of Eurolinguistics, but it underscores the importance of situating any comparative claims within explicitly stated criteria for what counts as *European*.

Similar uncertainties arise when defining what qualifies as a *European language*. On the most inclusive reading, this would mean all languages currently spoken within the widest geographical understanding of Europe, encompassing hundreds of autochthonous and migrant languages, as well as standard and non-standard varieties. On the narrowest definition, only "the indigenous standard varieties of official languages in the European Union" would qualify (Grzega 2025: 6). Between these poles, further ambiguities arise: should colonial varieties such as American English or Brazilian Portuguese be considered *European* in certain contexts, given their historical roots and ongoing discursive relevance? How should multiple standard varieties of the same language, such as British and American English, be treated? As Grzega emphasizes, "the relevance of each of the four questions depends upon the research question formulated and also the period investigated" (2025: 6). For the present thesis, which focuses on mainstream newspapers, this entails concentrating on the most widely used autochthonous standard varieties within Europe, while also considering extraterritorial Englishes where contextual relevance demands it. What emerges is not a universally applicable definition of *European language*, but rather the need for methodological clarity: researchers must explicitly delimit which languages and varieties they consider *European* in light of their empirical aims.

The question of what constitutes a specifically *European feature* raises further methodological complications. In principle, a European commonality would require a feature to be shared across all European

languages. Yet as Grzega notes, “most studies ... have investigated language-internal features” (2025: 8), and few traits are in fact universal. To address this, scholars have proposed threshold models: a feature may be considered *European* if it is present in two-thirds to three quarters of the languages or countries examined (Grzega 2013: 36; Grzega 2025: 9). The precise percentage varies according to the scope of the study: the fewer the languages analyzed, the higher the threshold should be to justify claims of Europeaness. This criterion underscores the delicate balance between inclusiveness and analytic rigor. Too narrow a definition risks excluding relevant variation, while too broad an approach may dilute the notion of Europeaness to the point of triviality. The problem is not merely quantitative, however, but also qualitative: some features may cluster in geo-cultural subregions without extending across the continent, producing patterns of regional convergence and divergence. Such dynamics call for careful interpretation rather than essentialist claims.

Closely tied to definitional concerns about languages and features is the question of what constitutes a valid Eurolinguistic study. Here, the issue is one of representativeness. As Grzega argues, two or three languages are not enough (2025: 10). Instead, the selection of languages and regions must be tailored to the type of study undertaken. For historical-anthropological and cultural analyses, this involves a balanced sample spanning culturally central areas (such as German-, French-, Dutch-, and Italian-speaking countries), peripheral regions (e.g. Finland, the Baltics, Hungary), intermediate cases, and ideally at least one borderline case such as Russia (Grzega 2025: 10). Geographical approaches require coverage across northern, western, southern, and eastern Europe, though terminological variation between UN and EU categorizations complicates these assignments. For historical-linguistic studies, representatives of all major Indo-European branches, as well as the Finno-Ugric languages, are indispensable. Synchronic investigations, by contrast, should include the Charlemagne Sprachbund, East-Central Europe, the Balkans, and Russia (Grzega 2025: 10). The underlying principle is clear: the credibility of Eurolinguistic claims hinges on a sufficiently broad and balanced empirical foundation. Narrowly focused studies can yield valuable insights, but without sampling, they cannot substantiate generalizations about Europe as a whole.

While much of Eurolinguistics has traditionally concentrated on structural or lexical features, recent scholarship has emphasized the importance of discourse-analytic perspectives. Giessen and Szwed argue that studying discourses across Europe is crucial because they are “of normative and political significance” (2025: 247). Their approach highlights that linguistic features cannot be detached from the communicative and cultural contexts in which they occur: “it is always important to consider the structures in which a text is found” (2025: 248). To illustrate, they compare TV crime dramas from five countries, finding both shared tropes – such as the older, male detective wrestling with existential dilemmas – and profound divergences rooted in national cultural codes. A second study of newspaper headlines around the European games similarly revealed cross-national similarities, such as predominantly negative framings of host countries Azerbaijan and Belarus, alongside striking national differences, with Luxembourgish newspapers offering comparatively positive coverage. From these cases, Giessen and Szwed conclude that discursive patterns can reveal both convergence and divergence, sometimes inverting expectations: whereas crime dramas displayed divergent symbols but convergent values, newspaper headlines converged on superficial rhetorical structures while diverging in underlying values (2025: 258). For present purposes, their contribution is twofold: it demonstrates the viability of discourse as an object of Eurolinguistic study, and it points to newspapers as particularly revealing sites where shared patterns and national specificities intersect.

Building on these theoretical and methodological foundations, a number of Eurolinguistic studies have turned explicitly to newspapers as empirical sites for investigating commonalities and divergences.

Headlines in particular have attracted scholarly interest because of their dual function: they condense complex events into compressed linguistic forms, and they exert what Grzega terms “enormous manipulative power” (2016: 19). This makes them especially suitable for probing whether European media share pragmatic tendencies in framing political and social issues.

One strand of such research has focused on the presence or absence of shared rhetorical values in newspaper discourse. In his comparative study of quality newspapers from eight countries, Grzega (2016) examined how concepts such as *competitiveness* and *solidarity* are framed, as well as the evaluative polarity of terms like *privatization*, *socialization*, and *welfare state*. His findings revealed both pan-European tendencies and geo-cultural splits. For instance, *competitiveness* appeared more prominently in Mediterranean countries (France, Spain, Italy) than in northern and Germanophone ones, suggesting that while neoliberal vocabulary was gaining ground across Europe, its salience varied regionally. *Solidarity*, by contrast, never outweighed *competitiveness*, highlighting a broader shift in press discourse away from traditional European values. The study thus not only demonstrated that headlines can reveal shared pragmatic orientations but also underscored the unevenness of such orientations across Europe.

Another approach has been to examine naming practices in headlines, which offer insights into how political figures are framed across national contexts. Würth’s (2016) analysis of designations for six heads of state and government in ten newspapers across five countries showed that, despite variation, the family name alone predominated across Europe. This tendency was partly pragmatic – driven by the brevity required in headlines – but it also revealed subtle evaluative strategies. For example, British newspapers displayed a stronger preference for combining first and last names, while in other contexts, figures like Vladimir Putin were more often presented mononymously, underscoring familiarity or notoriety. The study found that biased designations were relatively rare, yet deviations from the default surname pattern could function as implicit framing devices, signaling distance, familiarity, or value judgments. As discussed in section 2.3, Würth’s findings provide a methodological precedent for the present thesis: they demonstrate that seemingly simple naming practices can encode significant pragmatic and cultural variation, making them a valuable site for examining European commonalities and differences.

Further work has highlighted how newspapers employ metaphorical and conflict-related framings in ways that both converge and diverge across Europe. Hanusch’s (2014) study of headlines on the eve of the Iraq War demonstrated how five newspapers from Germany, Italy, the United Kingdom, Poland, and the United States drew on metaphorical domains such as natural catastrophe, theater, and games. These framings, while not simply reproducing U.S. government rhetoric, nonetheless downplayed the political and violent realities of the conflict by casting it as spectacle or inevitability. Hippler’s (2016) study of coverage of the Syrian war similarly revealed how European print media tended to construct binary oppositions between the West (often including allies like Qatar and Saudi Arabia) and actors such as Syria, Russia, and Iran. Her combination of qualitative discourse analysis and quantitative headline study showed a pronounced preference for *war* over *peace* terminology, as well as emotionally charged language in German and British reporting compared to more neutral tones elsewhere. Together, these studies demonstrate how newspapers across Europe both reflect and shape public perceptions of international conflicts, relying on discursive strategies that reveal common tendencies toward dramatization while also reproducing national differences in tone and framing.

Beyond individual case studies, Grzega’s work has systematically expanded Eurolinguistic inquiry into the domains of peace linguistics and bias analysis, showing how newspaper and magazine discourse both reflects and shapes collective orientations. A recurring theme across his research is the prevalence of violence-related vocabulary in contexts far removed from military conflict (2019), ranging from business

and sport to political campaigns. His analysis of news magazine covers across different European regions (2017) demonstrated that violence metaphors are not only entrenched in everyday discourse but also unevenly distributed, with some outlets normalizing conflict language more than others. This strand of work also revealed how European media framed global political actors differently: Barack Obama's military actions were often downplayed, while Vladimir Putin's were depicted in starkly negative terms, and Donald Trump was frequently cast as a threat compared to the more favorable treatment of Hillary Clinton. Such contrasts point to a pattern in which evaluative asymmetries are embedded within ostensibly descriptive discourse, illustrating how linguistic choices can legitimize or delegitimize political actors.

Grzega has since deepened his line of research with large-scale comparative analyses that span European and non-European contexts. His studies of the concepts *limits*, *liberty*, and *security* (2016b), the framing of *fake news* during the COVID-19 pandemic (2020), and the coverage of Russia's 2022 invasion of Ukraine (2023) collectively highlight how lexical and rhetorical framing interacts with broader cultural and political orientations. Across these works, discursive strategies such as naming strategies, the deployment of war rhetoric, and the inflationary use of stigmatizing labels emerge as central mechanisms through which media shape public perception. At the same time, his findings underscore that European newspapers do not present a monolithic discourse but rather exhibit both shared tendencies – such as the privileging of neoliberal vocabulary like *competitiveness* over *solidarity* – and regional or national divergences that complicate claims of a unified *European* voice. Taken together, this body of work demonstrates that newspapers provide a particularly fertile ground for investigating the pragmatolinguistic commonalities and divergences that define Europragmatics, while also revealing how linguistic framing contributes to wider processes of bias and cultural positioning.

Holler's (2023) cross-national headline study on China further anchors discourse-analytic Eurolinguistics by demonstrating how evaluative tendencies travel across languages, outlets, and media systems while remaining measurably heterogeneous. Working with 2,958 headlines from eight leading newspapers in Germany, Austria, the UK, and Ireland, she shows that neutrality dominates overall but that negativity framing consistently outweighs positive, with negativity rising from 17% (2010) to 21% (2020). Importantly, structural factors help explain this bias: newspapers from larger countries display more evaluative language than those from smaller ones; English-language outlets use more judgmental framings than German-language ones; and conservative outlets are reliably more negative than their center-left counterparts. Holler concludes that European press discourse exhibits a structural tilt toward negative portrayals of China, conditioned by political orientation and, to a lesser extent, language and country size. Her findings dovetail with the broader Eurolinguistic observation that European newspapers simultaneously display shared pragmatic orientations and regionally inflected divergences, while also prefiguring the patterns of negativity, neutrality, and ideological asymmetry that this thesis investigates in the context of Trump coverage.

Taken as a whole, the Eurolinguistic scholarship surveyed here highlights both the promise and the difficulty of identifying European commonalities in language use. Definitions of *Europe*, *European language*, and *European feature* remain contested, but what emerges clearly is the need for broad, balanced, and methodological explicit studies. Within this landscape, discourse-analytic approaches have proven especially valuable, as they shift attention from structural features to the ways in which language mediates social and political realities. Newspapers, and headlines in particular, stand out as productive focus: they condense complex issues into highly visible linguistic forms, and they exert considerable influence over public opinion. Prior work by Grzega, Giessen and Szwed, Würth, Hanusch, Hippler, Holler, and others has shown that newspaper discourse can reveal both pan-European tendencies – such as recurring metaphors, lexical asymmetries, or shared rhetorical strategies – and significant regional or national

divergences. For these reasons, newspapers provide an ideal corpus for testing whether pragmatic commonalities exist across European public spheres, and discourse analysis offers the methodological lens for uncovering them. At the same time, the growing scale and complexity of such corpora necessitate tools that go beyond traditional manual analysis. This is where advances in artificial intelligence, and large language models in particular, offer new possibilities for extending Eurolinguistic and discourse-analytic research.

### 3.2. Discourse Analysis with AI

The turn to artificial intelligence in discourse analysis reflects both a methodological opportunity and a challenge. For Eurolinguistic and media-bias research, the availability of thousands of headlines across multiple election cycles presents a scale that would be prohibitive for manual annotation alone. Large language models (LLM) offer a potential solution: they can process vast corpora rapidly, apply consistent coding schemes, and uncover patterns that might otherwise remain hidden. At the same time, their use raises critical questions of reliability, reproducibility, and interpretability. The following review considers recent evaluations of LLMs in discourse annotation, focusing on what they reveal about baseline capabilities, task-specific adaptations, and the methodological safeguards required to integrate such tools into scholarly research.

Early evaluations of LLMs suggest that even without fine-tuning, they are capable of handling core discourse tasks with surprising competence. Fan et al. (2024) tested GPT-3.5-Turbo on two foundational problems – segmenting conversations into coherent topics and identifying how individual turns relate to one another. With a structured but simple prompt, the model produced topic outlines that, in a manual audit, were as good as – or better than – human labels across several datasets (Fan et al. 2024: 17002). Its strength lays in detecting shifts in everyday, general-domain talk, where lexical and pragmatic cues are frequent and unambiguous. Limitations appeared in more specialized domains, and the model struggled with longer-range rhetorical connections, often linking only adjacent turns rather than capturing deeper conversational structures (Fan et al. 2024: 17003–17004). It also occasionally failed to adhere to the requested output format, requiring post-processing. Taken together, the findings show that off-the-shelf LLMs can perform general discourse segmentation with near-human reliability, but their precision declines as tasks demand more domain knowledge of complex rhetorical modeling. These results underscore both the promise and the limits of applying early-generation models to nuanced discourse annotation – while also pointing toward steady improvements as newer, more powerful models become available.

Subsequent work has explored how LLMs perform once prompts, data regimes, or training procedures are adapted to the annotation task. A recurring finding is that model performance improves substantially when the task is made cue-rich and constrained, but that gains diminish for fine-grained, context-sensitive labels.

Garg et al. (2024) illustrate the value – and the limits – of task adaptation. Testing GPT-3.5-Turbo and GPT-4 (and a fine-tuned GPT-3.5) on 6,405 interview responses labeled into eight practical categories, they compared simple zero-shot instructions, few-shot prompts, prompts that foregrounded indicative phrases, and task-specific fine-tuning. Their main takeaway is pragmatic: the fine-tuned model performed best and, in some categories, reached or exceeded common reliability cutoffs (Garg et al. 2024: 817). Yet, averaged across categories and prompting modes, no single method consistently met the field's thresholds for full automation; performance remained uneven, particularly for nuanced or low-resource

classes. The study therefore positions LLMs as promising assistants for broad, cue-rich categories but cautions against relying on them for subtle, contextually anchored codes without additional safeguards.

Ostyakova, Mikhailova, and Konovalov (2025) approach the same trade-off from a different angle: they test three comparatively affordable LLMs (GPT-3.5-Turbo, Claude-3-Haiku, Open-Mixtral) on a hierarchical speech-function scheme for 1,030 utterances. Their methodological innovations – multi-step prompts that mirror human guidelines, masking of label names to reduce hallucination, and staged decoding – improved robustness on coarse cut labels (expert  $\kappa=0.83$ ) while revealing clear limits on full label granularity. Claude-3-Haiku emerged strongest on the cut labels, GPT-3.5 performed moderately, and Mixtral lagged and occasionally failed to respect output constraints. Hybrid setups (LLM + crowd voting) improved over single-model runs but did not outperform human-only aggregation. Importantly, the authors emphasize practical frictions: model/version instability, the engineering cost of multi-step prompts, and the residual need for human adjudication on fine-grained distinctions.

Siiman et al. (2023) provide a complimentary, task-specific perspective: when categories are transparent, theory-guided, and deductive, GPT-4 can reduplicate structured human judgments reliably. Translating and coding 40 Estonian chats, they found that carefully refined instructions and embedded task context led to deductive judgements that largely matched a human’s coding and were reasonably consistent across runs (Siiman et al. 2023: 92–93). However, the model struggled with open-ended, inductive tasks (e.g. inventing a rubric), produced more variable outputs, and remained constrained by input length (Siiman et al. 2023: 93–95). Their conclusion – that LLMs are useful for transparent, theory-guided coding but less reliable for open-ended interpretation – is directly relevant to headline analysis, where labels can be made explicit and anchored in lexicogrammatical cues.

Yu (2025) shows how a pragmatic engineering approach can combine automation with targeted human oversight to yield substantial efficiency gains. After iteratively refining a prompt on 30 CSR statements, Yu applied a “MoveAnnotator-GPT” to 50 unseen texts and ran the model twice to measure stability. The initial sentence-level accuracy was high (87.14%), with the model strongest on cue-rich moves (reporting accomplishments, outlining strategy) (Yu 2025: 34–35, 42–43). The runs revealed intra-model inconsistency ( $\approx 33\%$ ) and persistent error in ambiguous or residual categories, but a pragmatic workflow – automatic coding plus human review of only the conflicted/ambiguous items – lifted corrected accuracy to 95.76% while drastically cutting manual effort (Yu 2025: 42–43, 46). Yu’s study therefore provides a concrete blueprint for hybrid workflows that preserve human judgment where it matters most.

Across these studies a consistent pattern emerges: (i) explicit, cue-rich tasks and compact, well-worded label sets boost LLM performance; (ii) fine-tuning or multi-step prompting narrows error but introduces engineering and transparency costs; (iii) different model families and versions perform idiosyncratically, creating versioning risks; and (iv) hybrid pipelines (LLM + human adjudication) balance scalability with reliability. In short, LLMs markedly expand the feasible scale of discourse annotation, but they do not (yet) eliminate the need for targeted human oversight – especially for subtle pragmatic judgments or rare categories. A final methodological caveat in this cluster of work concerns randomness and reproducibility. Several authors stress that off-the-shelf consumer interfaces may mask model routing and sampling controls; in practice, reproducibility improves when researchers lock model versions and decoding parameters via API access and set temperature to zero for deterministic decoding. Thus, while model families have advanced rapidly (the baseline studies often used GPT-3.5 or GPT-4), claims of reliability must be re-established on each new generation and task under controlled decoding and transparent prompts. Taken together, these findings justify an AI-assisted design for large-scale headline coding – provided

that (a) prompts are explicit and task-tailored, (b) model versions and decoding settings are fixed and reported, and (c) human audit layer adjudicates ambiguous or rare cases.

The insights from Eurolinguistics and recent advances in AI-assisted discourse annotation converge in important ways. Eurolinguistic studies have shown that newspapers are fertile sites for tracing both shared pragmatic tendencies and regional divergences across Europe. At the same time, evaluations of LLM demonstrate that large-scale, systematic coding of such corpora is increasingly feasible, provided prompts are explicit, model versions are fixed, and human oversight remains in place for ambiguous cases. This combination opens the door to a new kind of inquiry: the cross-national, longitudinal study of bias in political headline discourse at a scale not previously possible. The question then arises of how such a study connects to the broader work on media bias, particularly in relation to the figure of Donald Trump, whose press coverage has been uniquely polarized. It is to this field of research that the next section now turns.

### 3.3. Trump, Media Bias, and Political Reporting

Few political figures in recent decades have attracted as much sustained media attention as Donald J. Trump. From the moments of his candidacy announcement in 2015, Trump became a central figure in domestic and international news coverage, a prominence that both reflected and amplified his self-styled role as an outsider challenging established norms. The literature on Trump’s media portrayal is accordingly extensive, spanning the U.S. press and broadcasters, European quality dailies, and outlets across Asia, Africa, Latin America, and the Middle East. What emerges across these diverse settings is a consistent pattern: Trump drew disproportionate coverage compared to his rivals, and that coverage was more often negative than positive. This tendency has important implications for the present study, as it directly relates to the hypotheses on evaluative distribution (H1) and candidate-specific bias (H2), while also framing the broader question of whether such tendencies reflect national particularities or cross-national discursive alignments.

At the level of volume, Trump consistently received more attention than his opponents. In his analysis of U.S. election coverage across ten major newspapers and television broadcasts, Patterson (2016: 7) observes that Trump drew more attention than Clinton (by ~15%), a finding that holds across both print and television outlets. Importantly, Patterson shows that this heightened salience coincided with overwhelmingly negative reporting: his coverage was relentlessly negative, never turning positive and ranging from roughly two-to-one more than ten-to-one negative-to-positive (Patterson 2016: 11). International studies corroborate this asymmetry. Nord, Mancini, and Gerli (2017: 9), analyzing 933 articles from Italy, Sweden, and the United Kingdom, conclude that “Trump was considered much more newsworthy than Clinton,” a pattern evident in all three countries despite variations in tone. Similarly, Hinck (2024: 120-121), examining elite media in Chinese, Russian, Iranian, and Saudi contexts during the 2020 election, finds that Trump was more often discussed than Joe Biden (68.8% versus 61.6%), echoing the earlier 2016 imbalance. Together, these findings point to a structural feature of Trump coverage: before questions of tone are even considered, his heightened visibility placed him at the center of journalistic attention.

If visibility was the first hallmark of Trump coverage, tone was the second. Across settings, the dominant pattern was one of negativity, confirming Trump’s perception of an adversarial press and forming the backdrop for hypothesis H2, which anticipates that Clinton, Biden, and Harris receive a more positive balance of coverage than Trump. Patterson’s (2016: 11-17) study provides a U.S. benchmark: in ten leading outlets, every outlet was net negative on Trump, with figures ranging from 73% at Fox to 89%

negative at CBS. Coverage was not only critical in aggregate but consistently harsh across topic categories, from horserace reporting to leadership and experience. Clinton's coverage was also net negative, but the proportions varied more widely by outlet, and she enjoyed occasional positive spikes, particularly in the context of debates. International evidence shows a broadly similar pattern, lending support to H7 on European unity of bias and H8 on Europe versus the rest of the world. In Nord, Mancini, and Gerli's (2017: 8) comparative analysis, Trump's coverage was predominantly negative in all three countries studied – 66% in the UK, 62% in Italy, and 49% in Sweden – while Clinton's was more mixed, with Sweden leaning favorable and Italy more critical (2017: 7). In German prestige media, Cazzamatta and Hafez (2021: 131-137) report that fully 65.6% of *Der Spiegel* coverage was negative, often casting Trump as *anti-democratic*, *racist*, or *corrupt*, reinforcing a moralized register of critique. Similar tendencies appear in studies of Greek online newspapers, where reporting on Trump was overwhelmingly neutral (80.4%) but with negative framings far outweighing positive ones (Zaharopoulos 2021: 151–152), and in Malaysia and Nigeria, where Trump himself was depicted more unfavorably than the U.S. as a whole (Baghestan & Osman 2021; Baghestan & Akoje 2021). In Latin America, Colombian (Arroyave 2021) and Mexican (Lozano & Martínez-Garza 2021) likewise presented Trump in overwhelmingly critical terms, highlighting arrogance, misogyny, and hostility toward migrants. A more complex picture emerges in Russia, China, and Egypt, which function as partial exceptions to this near-universal negativity. In Russian newspapers, Bykova, Gavra, and Slutskiy (2018: 9) find that 41% of coverage between 2017 and 2018 was positive and only 21% negative, portraying Trump as a strong leader constrained by domestic elites. Yet this early positivity faded over time: Rovinskaya and Greydina (2021: 268-270) show that by 2019, Russian press and television had shifted toward predominantly negative framings, particularly in relation to NATO, Ukraine, and impeachment proceedings. Chinese coverage in *People's Daily* was more balanced overall, alternating between praise and criticism depending on bilateral relations; positive peaks (up to 85.0%) coincided with Xi's 2017 state visit, while trade-war periods produced strongly negative coverage (up to 83.3%) (Sun & Song 2021: 75–76). Egyptian outlets presented Trump more favorably still, framing the U.S. administration as a valuable partner and depicting the Obama administration in harsher terms (Allam 2021: 103). Despite these exceptions, the weight of evidence points to a consistent global tilt: Trump was more often framed negatively than positively across diverse contexts, with temporary or interest-driven deviations rarely altering the overall balance. For this thesis, these findings establish two key points. First, they justify treating candidate-specific negativity as a structural feature of Trump coverage (H2). Second, they highlight the need to examine whether Europe displays distinctive unity (H7) or whether its negativity toward Trump is qualitatively stronger than in other regions (H8).

While negativity was widespread, its intensity and orientation often reflected the ideological leanings of individual outlets and the media ecologies in which they were embedded. This is central to hypothesis H3, which predicts that right-leaning newspapers frame Trump more positively than their left-leaning counterparts, and to H11, which anticipates that in Europe such ideological divides exert less influence than in the U.S. context. In the American case, Patterson (2016: 11) shows that although every outlet was net negative on Trump, *Fox News* was the mildest in tone (73% negative), while *CBS* was the most critical (89% negative). Clinton's coverage, too, varied by outlet, with the *Los Angeles Times* offering a relatively balanced mix (53% negative, 57% positive) compared to the *Washington Post* (77% negative) and *Fox* (81% negative) (Patterson 2016: 14). Such differences, though modest, reflect the strong alignment between partisan media ecosystems and evaluative stance, with *Fox* leaning toward Republican sympathies and other outlets displaying Democratic or centrist orientations. Outside the United States, ideological divides were sometimes sharper. In Australia, Cohen, Gurney, and Castillo (2021: 58-65) show that *News Corp* papers (among them *The Australian*) tended to present Trump favorably, particularly on economic and foreign policy issues, while *Nine Entertainment's* titles (among them *The Age*)

framed him as erratic, chaotic, and untrustworthy. The *ABC* consistently framed him as destabilizing and unprecedented in his impact on democratic norms. In Poland, Płodowski (2021: 247–256) finds that right-wing outlets and the ruling Law and Justice party embraced Trump as a populist ally, in contrast to liberal media that depicted him as erratic and threatening. These cases confirm the expectation that outlet ideology shapes coverage, broadly consistent with H3. By contrast, European prestige media frequently displayed a more unified critical stance, suggesting the distinctiveness posited in H11. In *Der Spiegel*, for instance, negativity was overwhelming across the political spectrum, and critiques were framed through shared moral categories – despotism, corruption, racism – rather than partisan alignments (Cazzamatta & Hafez 2021: 131–137). Similarly, in Greek and Nigerian contexts, outlet-level differences were minimal: reporting was largely consistent across newspapers regardless of government alignment (Zaharopoulos 2021; Baghestan & Akoje 2021). These patterns suggest that in many non-U.S. settings, and in Europe in particular, Trump’s polarizing persona and policies overrode traditional left-right divides, producing an unusual degree of cross-ideological convergence. Taken together, the evidence supports the claim that ideology matters, but its weight varies by media system. In the U.S. and other Anglophone contexts, partisan alignments remained visible in tone and framing, consistent with H3. In European contexts, however, the relative uniformity of critical portrayals points toward H11: Trump elicited a broadly shared distancing that muted the effects of traditional ideological cleavages.

The temporal dynamics of Trump coverage also reveal important patterns for assessing hypothesis H4, which anticipates stability in the distribution of bias across electoral cycles. Several studies show that peaks and troughs in tone are closely tied to discrete events, but that the overall evaluative balance remains consistent. In Patterson’s (2016: 11–17) study of U.S. election coverage, negativity toward Trump dominated from August through November 2016, with only slight fluctuations. The two “best” weeks for Trump occurred immediately before Election Day, but even these periods failed to shift his coverage into net positive territory. Clinton’s reporting followed a similar trajectory: persistently negative overall, with her one positive week appearing after the first debate. These findings highlight the stability of negative framing across the campaign, even as weekly variations reflected momentary developments. Comparable patterns appear in Nord, Mancini, and Gerli’s (2017: 6–7) European study of the 2016 campaign. In Italy, Sweden, and the UK, Trump’s coverage was consistently more negative than Clinton’s, though debate periods and conventions briefly altered the ratio of positive to negative reporting. Importantly, the spikes did not reverse the general trend: Trump’s negativity remained dominant across the campaign, suggesting that event-driven shifts were short-lived against a backdrop of structural bias. Later studies extend this temporal stability beyond 2016. Hinck (2024: 117–122), examining coverage of both the 2016 and 2020 elections across Chinese, Russian, Iranian, and Saudi outlets, finds that Trump was again more frequently mentioned than his opponents and disproportionately covered in negative terms. Although the intensity varied by country – negativity intensified in Chinese outlets while softening slightly in Russian ones – the overarching pattern was consistent with 2016: Trump remained the more negatively framed candidate. The contrast with Biden, whose 2020 coverage was predominantly neutral, underscores the persistence of evaluative asymmetry across cycles. Event-specific spikes are visible in non-electoral coverage as well. Bykova, Gavra, and Slutskiy (2018: 9–12) document how Russian media shifted toward negativity in response to U.S. missile strikes in Syria or withdrawal from the Paris climate accord, while reverting to neutral or positive tones in calmer periods. Similarly, Sun and Song (2021: 75–76) trace oscillations in *People’s Daily*, with highly positive coverage around Xi’s 2017 state visit but surges of negativity during tariff escalations in 2018. Yet in both cases, these fluctuations occurred within a stable evaluative orientation; broadly supportive in Russia, more balanced in China. Taken together, these findings suggest that event-driven volatility does not undermine structural tendencies. Whether in the U.S. or abroad, the overall direction of Trump coverage – predominantly negative – remained stable across both short-term news cycles and multiple electoral campaigns. This evidence provides strong support for

H4: while the intensity of bias responds to specific events, its distribution across time exhibits remarkable consistency.

A further dimension of media bias toward Trump lies in the framing repertoires and evaluative categories that dominate his coverage. Large-scale semantic analyses confirm that moralization, rather than assessments of competence, structured much of the press discourse around the 2016 campaign. Bhatia, Goodwin, and Walasek (2018) demonstrate that partisan outlets diverged most strongly in their associations of Clinton and Trump with moral traits, concluding that morality rather than competence is the central dimension of media portrayals in that cycle. This finding resonates directly with H6, which predicts that evaluative labels disproportionately attach to Trump: if moral character becomes the primary lens through which candidates are judged, then descriptors such as *racist*, *sexist*, or *authoritarian* are more likely to cluster around him than around his Democratic opponents. Individual case studies of international coverage confirm this tendency toward moral framing. Isakhan, Nwokora, and Pan (2019) show that Arab News coverage of Trump's proposed Muslim ban overwhelmingly denounced it as racist and Islamophobic, while also casting doubt on the health of U.S. democracy. Similarly, Akdenizli, Özçetin, Gündoğdu (2021) find that Turkish newspapers, though largely neutral overall, invoked themes of sovereignty and national betrayal when covering U.S.-Turkey conflicts, frames that align with broader nationalist and moral vocabularies. Such examples illustrate how recurring repertoires – democracy, race, sovereignty – converge on moralized evaluations that extend beyond Trump's policies to his perceived character and legitimacy. Genre also plays a crucial role in intensifying these patterns. Cazzamatta & Hafez's (2021) analysis of *Der Spiegel* shows how magazines, less constrained by brevity than daily newspapers, elaborated moral vocabularies into clusters: Trump was framed as “an unreliable partner,” anti-democratic and despot,” “populist, racist, and sexist,” “unprepared and inept,” “corrupt and dishonest,” and “disdain[ful of] science” (Cazzamatta & Hafez 2021: 131). Even as the magazine condemned him, it simultaneously amplified his rhetoric, inadvertently circulating the very populist discourse it critiqued (Cazzamatta & Hafez 2021: 137, 139). In another genre, British satire and protest likewise reinforced moralized portrayals: Stokes (2021) shows how Steve Bell's cartoons and the “Trump Baby” blimp cast the president as infantile and irrational, crystallizing negative judgements through ridicule. Taken together, these studies underscore that moralization is not incidental but structural: across contexts and genres, Trump was cast less as a competent or incompetent politician than as a figure embodying vice or threat. For the present thesis, this provides a conceptual pivot. Headlines, as the most condensed form of journalistic discourse, are unlikely to sustain extended moral vocabularies, but they nonetheless crystallize evaluative stances through selective naming and value-laden labels. In this sense, headlines distill the broader moralization of Trump into compressed linguistic cues, making them a crucial site for testing H2, H3, and especially H6.

Cross-regional contrasts further sharpen the relevance of Eurolinguistics for media bias research. The evidence consistently shows that Trump's coverage was not only negative overall but also more uniformly negative across Europe than in other regions, lending support to H7–H10. Nord, Mancini, and Gerli's (2017: 7–9) comparative study of Italy, Sweden, and the UK during the 2016 election demonstrates that Trump was covered more negatively than Clinton in all three countries, with particularly strong negativity in the UK (66%) and Italy (62%). Although Sweden registered a somewhat milder pattern (49% negative, 43% neutral), the broader finding is clear: European outlets converged on a skeptical orientation toward Trump. Clinton's coverage, by contrast, was more mixed and varied across national contexts. These results prefigure H7 (European unity of bias) and H9 (cross-linguistic convergence): despite cultural and linguistic differences, European outlets aligned in their evaluative stance toward Trump. Hinck's (2024: 117–122) large-scale study of both the 2016 and 2020 elections confirms the distinctiveness of the European pattern by placing it in international perspective. Examining Arabic,

Chinese, Russian, and Iranian media, Hinck finds that Trump was again covered more often and more negatively than his Democratic opponents, but with notable variation. In 2020, Chinese and Iranian outlets produced the strongest negativity ( $m = -0.30$  and  $-0.34$ ), Saudi media displayed a milder pattern ( $m = -0.20$ ), and Russian coverage approached neutrality ( $m = -0.07$ ). Biden, by contrast, was consistently framed in more neutral or slightly positive terms. The comparison highlights two points relevant for the present study: first, that Trump's negative framing extended well beyond Europe, suggesting a global pattern; and second, that Europe's consistency across multiple subregions makes it distinctive when compared to the more varied orientations in non-European contexts. Sub-continental variation within Europe also warrants attention. Nord, Mancini, and Gerli's (2017) found that Northern Europe (Sweden) was somewhat less negative than Western Europe (UK) and Southern Europe (Italy), a contrast that complicates any claim of a monolithic European discourse. Yet even here, the evaluative balance tipped against Trump in every country surveyed, supporting H10 (regional consistency) while leaving space for degrees of variation. The persistence of negative coverage across linguistic families (Germanic in the UK and Sweden, Romance in Italy) further underpins H9 (cross-linguistic convergence). Together, these studies provide a comparative baseline for the present thesis. They suggest that while global coverage of Trump has varied by geopolitical context, European newspapers have tended to align in their negativity, transcending linguistic and regional boundaries. This evidence supports the expectation in H7–H10 that European newspapers will display shared patterns of bias that both distinguish them from non-European outlets and unify them internally despite cultural and geopolitical diversity.

A final set of observations concerns the methodological diversity that characterizes the scholarship on media bias and Trump. Researchers have pursued different units of analysis – ranging from entire articles (Patterson 2016; Nord, Mancini & Gerli 2017), to magazine covers (Cazzamatta & Hafez 2021), to television programs (Rovinskaya & Greydina 2021). Each of these units can yield valuable insights, but their heterogeneity complicates direct comparison across studies. In the present thesis, headlines were chosen as the unit of analysis. This decision is not a claim to superior validity but a practical choice: headlines are short, condensed, and easily retrievable across multiple outlets and languages. Their brevity makes them particularly well-suited to systematic coding at scale, and they crystallize evaluative stances in a form that is readily comparable across electoral cycles.

Sampling strategies in the existing literature are equally varied. Some studies restrict themselves to constructed weeks (Zaharopoulos 2021), others to the convention and debate periods (Nord, Mancini & Gerli 2017), and still others to entire multi-year spans (Bykova, Gavra & Slutskiy 2018; Sun & Song 2021). Each approach reflects deliberate methodological choices, but the lack of alignment in temporal scope means that findings cannot be straightforwardly integrated. By contrast, the present thesis selects identical event windows across the 2016, 2020, and 2024 campaigns. This design makes it possible to track changes over time with greater consistency while still remaining comparable across candidates.

Variation is also evident in how tone is measured. Patterson (2016), for example, collapsed a six-point scale into three categories, while others report only positive and negative without neutrality. These choices make sense within the scope of individual studies but complicate numerical comparability across the field. The present analysis therefore adopts a tripartite distinction – positive, negative, and neutral – that captures a fuller distribution of stance. This means its percentages cannot be directly compared with studies that exclude neutrality, but the approach aligns with the central aim of testing whether evaluative headlines outweigh neutral ones.

Finally, outlet selection has been uneven across prior studies. Some focus on television, others on children's or partisan outlets, while many – including the majority of international comparisons – concentrate

on prestige or quality newspapers (Cazzamatta & Hafez 2021; Baghestan & Akoje 2021). This thesis follows the latter convention by restricting its corpus to qualifying outlets. Such a focus inevitably excludes tabloids and online-only media, which are also influential, but it reflects the dominant practice in comparative media research and ensures continuity within existing scholarship.

Taken together, these methodological differences demonstrate the richness but also the fragmentation of prior work on Trump and media bias. They underscore the need for approaches that are both standardized and transparent across units, timeframes, and categories of tone. The present thesis seeks to address this by applying consistent headline-level coding to aligned event windows across three electoral cycles, thereby laying the groundwork for the next step: identifying what has so far remained under-tested in the literature.

### 3.4. Research Gap

Despite the remarkable growth of scholarship on media bias and the Trump presidency, important questions remain underexplored. Existing studies have shown convincingly that Trump attracted disproportionate attention, that coverage tilted strongly negative across a wide range of outlets, and that patterns of evaluation varied across national and ideological contexts. They have also demonstrated the relevance of multiple genres – from television broadcasts to magazine covers to full-length articles – and developed sophisticated coding schemes to track tone, framing, and moral evaluation. Together, this body of work has established a robust foundation for understanding how Trump was represented in the media, and it has illuminated broader dynamics of political communication in an era of populism, polarization, and global contestation.

Yet the diversity of approaches also reveals fragmentation. Much of the existing research works with heterogeneous units of analysis – articles, constructed weeks, or selected covers – making it difficult to compare findings across studies. Sampling windows are often tied to particular events, such as conventions or debates, or span extended periods without a consistent anchor, which limits the ability to assess continuity across election cycles. Measurement strategies differ as well: some studies collapse complex tone scales into simple positive/negative categories, others omit neutrality altogether, and many vary in how they operationalize bias. These choices are defensible within each individual study, but the result is a patchwork of findings that do not easily add up to a coherent, longitudinal picture.

It is here that the present thesis positions itself. By focusing on headlines, it concentrates on a textual form that is both practically retrievable and methodologically significant. Headlines condense complex issues into compressed linguistic forms, exert strong framing power, and can be collected consistently across countries, languages, and electoral cycles. Aligning identical event windows across the 2016, 2020, and 2024 campaigns provides a further layer of comparability, enabling systematic assessment of continuity and change over time. By retaining a tripartite distinction between positive, negative, and neutral evaluations, this study also avoids the loss of nuance that results when neutrality is omitted, offering a more transparent account of the distribution of bias.

At the same time, Eurolinguistic research provides a crucial motivation for the thesis. As Grzega (2025) and Giessen and Szwed (2025) emphasize, newspapers are especially fertile sites for identifying pan-European tendencies in discourse, while also revealing regional divergences. Prior studies in this tradition have demonstrated how European media frame political actors differently, how metaphors of violence and conflict recur across national boundaries, and how naming strategies encode evaluative nuances. Yet despite these advances, Eurolinguistics has not yet produced a systematic, multilingual comparison of

European headline discourse on U.S. presidential elections across multiple cycles. Questions of whether European newspaper coverage in their portrayals of Trump, whether regional or linguistic subgroups diverge, and whether European ideological divides are overridden by a shared distance from Trump remain unanswered. By bringing Eurolinguistic insights into conversation with media-bias scholarship, the present thesis aims to fill precisely this gap.

Finally, the methodological possibilities opened by artificial intelligence have yet to be fully integrated into this field. Prior work on media bias has relied heavily on manual coding, which ensures interpretive depth but limits scalability and reproducibility. Evaluations of large language models suggest that, with careful prompt design and targeted human oversight, they can reliably code large corpora while maintaining methodological transparency. To date, however, such approaches have not been applied to cross-national headline analysis of Trump coverage. This study therefore contributes not only substantively, by comparing European and non-European newspapers across three electoral cycles, but also methodologically, by testing the potential of hybrid AI-human workflows for discourse analysis.

Taken together, these considerations define the space in which the present thesis intervenes. Prior research has provided invaluable insights into media portrayals of Trump and into the discursive tendencies of European and global journalism. What has been missing is a design that combines consistent headline-level coding, aligned event windows, integration of neutrality alongside positive and negative tones, and systematic cross-European comparison informed by Eurolinguistics. By addressing these gaps, the study builds on and extends the achievements of earlier scholarship, while offering a new perspective on how political bias manifests in headline discourse across time, space, and media systems.

The preceding review provides the foundation for the dataset that follows. Eurolinguistics motivates a geographically and linguistically balanced European sample, with external comparisons to test distinctiveness. AI research justifies headline-level coding at scale, provided models and prompts are fixed and human oversight is retained. Trump/media-bias scholarship defines the empirical expectations in H1-H12. Section 4 now operationalizes these commitments by specifying a corpus of quality newspapers, aligned event windows across 2016/2020/2024, and a tripartite stance scheme (positive/negative/neutral) suited to headline analysis.

## 4. Data

The selection of countries for this study follows methodological recommendations in Eurolinguistics, which stress that a study should encompass the main geographical regions and major linguistic families of Europe. Grzega (2025: 10–11) argues that a project can only be called Eurolinguistic if it includes representatives from northern, southern, eastern, western, and central Europe, ensuring that no single area dominates the perspective. In line with this, the European corpus of this study includes the United Kingdom and Ireland in the north, Poland and Bulgaria in the east, Italy and Spain in the south, France and the Netherlands in the west, and Germany and Switzerland in the center. This geographical spread allows for the study of European headline practices across the full compass of the continent and provides a dataset that avoids the bias of focusing exclusively on Western European languages. A further guiding principle of Eurolinguistics is that representativeness should extend to linguistic families. As Grzega underlines, a Eurolinguistic study should include representatives of the main Indo-European branches – Germanic, Romance, and Slavic – as well as other major families where relevant (Grzega 2025: 10). The chosen countries reflect this principle: English, Dutch, and German represent the Germanic group, French, Italian, Spanish, and Portuguese the Romance group, and Polish and Bulgarian the Slavic group.

This balance permits an analysis of whether similarities in discourse strategies cluster along linguistic-family lines, or whether broader, cross-family European commonalities can be identified.

In addition, Eurolinguistics recognizes that “Europe” is not only a geographical concept but also a cultural and political one, defined by contrasts between centers and peripheries (Grzegza 2025: 3–5). The selection therefore deliberately includes both EU core members, such as Germany and France, and more peripheral states, such as Bulgaria and Poland, as well as the non-EU member Switzerland, whose historical trajectories and political orientations provide valuable comparative perspectives. This inclusion allows the analysis to address whether European commonalities in discourse arise more strongly within the EU core or whether they extend across the continent’s east-west and north-south divides. Finally, Eurolinguistics emphasizes that the identification of “European features” often only becomes meaningful when contrasted with patterns from outside Europe (Grzegza 2025: 8–10). The inclusion of non-European countries in this study therefore serves to test whether features observed in European newspapers are distinctive to the continent or part of wider global tendencies. By juxtaposing European newspapers with those from North America, South America, Asia, Africa, and Oceania, the study situates its findings within the broader theoretical framework of Eurolinguistics, which seeks to identify not only intra-European commonalities but also their position within the global linguistic landscape.

At the same time, certain caveats must be acknowledged regarding the scope and expectations of this selection. While the methodological rationale ensures broad geographical and linguistic representation, the analysis may reveal that patterns do not align neatly along these lines. It is possible that instead of a clear “Europe versus the rest of the world” stance, there may be significant internal diversity within Europe itself, with national traditions or political climates shaping discourse in unique ways. Likewise, similarities in reporting strategies may not necessarily follow from geographical or linguistic proximity but rather from ideological orientation, with newspapers across different countries aligning in their treatment of candidates along left-right divisions. In this sense, the framework of Eurolinguistics provides a valuable starting point for structuring the corpus, but it does not predetermine the results. Instead, the study remains open to the possibility that European commonalities emerge only in certain respects, or that headline practices are better explained by transnational political or cultural alignments than by strictly geographical or linguistic ones.

#### 4.1. Corpus and Data Selection

Building on these theoretical considerations, the next step is the practical selection of outlets. The overarching principle is to maximize representativeness while ensuring methodological consistency across regions. Three criteria guide the process. First, whenever possible, the highest circulation quality daily newspapers are chosen, with boulevard and tabloid outlets excluded to maintain comparability across serious journalistic sources. Second, to reflect the ideological dimension of discourse, each country is ideally represented by two outlets situated on opposite sides of the political spectrum, usually one center-left/left-leaning or liberal and one center-right/right-leaning or conservative. Third, the scope of the study is necessarily shaped by the availability of outlets in *Lexis Nexis*, which serves as the sole database for headline retrieval. As a result, not all ideal pairings can be realized: in several cases only one outlet meets the inclusion criteria, and in others the *Lexis Nexis* corpus imposes temporal restrictions.

Bias labels for the selected newspapers were obtained from established external sources. For European outlets, eurotopics.net was consulted; for the United States and the *South China Morning Post* Ad Fontes Media Bias Chart; for *El Universal* All Sides Media Bias; for Australia, Nigeria, *China Daily*, and *Reforma* Media Bias Fact Check; for Brazilian newspapers Wikipedia. These labels serve as a transparent

reference point for situating each outlet within the left-right spectrum, while recognizing that such classifications are necessarily simplifications and that newspapers may exhibit variation over time or across sections. Together, these outlets constitute a corpus that balances geographical, linguistic, and ideological representativeness, while remaining feasible within the constraints of *Lexis Nexis*. Table 2 and Table 3 provide an overview of all selected newspapers, their political leanings, and temporal restrictions. Figure 1 and Figure 2 show the distribution of European and non-European countries chosen for analysis.

#### 4.1.1. Europe North

In Northern Europe, the United Kingdom and Ireland were selected. For the UK, *The Guardian* (center-left) and *The Times* (conservative) were included. Both are quality dailies with among the highest circulations in their respective ideological camps, and both are available for the full time span in *Lexis Nexis*. *The Guardian* corpus includes both print and online material, while for most other newspapers such distinctions are not specified. There are no tags indicating the difference between print and online material, so all headlines are included in the analysis. Ireland was chosen as a complement to the UK to ensure Northern Europe was represented by two distinct countries. Attempts to include Scandinavia were unsuccessful: for Denmark, Finland, Iceland, Norway, and Sweden, *Lexis Nexis* primarily offered wire services, press-release distributors, or financial information providers, none of which fit the present study's criteria. Although the Irish press is generally less overtly partisan than the British, *The Irish Times* (center-left) and the *Irish Independent* (conservative) provide a usable contrast, and both are available for all cycles.

#### 4.1.2. Europe East

Eastern Europe is represented by Poland and Bulgaria. In Poland, *Gazeta Wyborcza* (liberal) is the only usable outlet in *Lexis Nexis*, available across all cycles. A conservative counterpart such as *Rzeczpospolita* would have been ideal but is not available in the database. Bulgaria is represented by *Дневник* (*Dnevnik*), which appears in *Lexis Nexis* as multiple separate section corpora (*Analizi*, *Biznes*, *Bulgaria*, *Evropa*, *World*, *Zelen*). For this study, these sections were aggregated and treated as one outlet. *Dnevnik* is classified as liberal-conservative and is available from 2017 onward, which may limit coverage for the 2016 cycle. *Kapital*, another suitable Bulgarian daily, is present in *Lexis Nexis* only between April 2017 and October 2019 and was therefore excluded.

#### 4.1.3. Europe South

Southern Europe is covered by Italy and Spain. In Italy, *Corriere della Sera* (liberal-conservative) was chosen as a consistently available outlet across all cycles. To provide a counterpart, *La Stampa* (liberal) was added, though its coverage in *Lexis Nexis* ends in June 2019, limiting it to the 2016 cycle. The high circulation *La Repubblica* would have been preferable but is not available in the database. The wire service *ANSA*, which was used for the AI test runs (see section 5.2) was excluded from the study, as the main focus is on newspaper reporting. In Spain, *El Pais* (center-left) and *El Mundo* (conservative) were selected as the country's highest circulation quality dailies, both available across all cycles.

#### 4.1.4. Europe West

For Western Europe, France and the Netherlands were chosen. In France, *Le Figaro* (liberal-conservative) is available for all cycles, while *Libération* (center-left) is included from 2018 onward, making coverage of the 2016 cycle impossible for this outlet. *Le Monde*, France's highest-circulation daily, is

absent from *Lexis Nexis*. The Netherlands is represented by *De Volkskrant* (center-left) and *De Telegraaf* (right-wing), both available for all cycles. Belgium was considered, but the Netherlands was ultimately preferred to add Dutch-language material and thus increase linguistic diversity.



Figure 1 Distribution of European countries chosen for analysis (created with mapchart.net)

| Area          | Country     | Newspaper (political leaning)                                                                                                                    |                                                    | Notes                                                                                  |
|---------------|-------------|--------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------|----------------------------------------------------------------------------------------|
| Europe North  | UK          | <i>The Guardian</i> (center-left)                                                                                                                | <i>The Times</i> (conservative)                    | <i>The Guardian</i> corpus includes both print and online content; cannot be separated |
| Europe North  | Ireland     | <i>The Irish Times</i> (center-left)                                                                                                             | <i>Irish Independent</i> (conservative)            | -                                                                                      |
| Europe East   | Poland      | <i>Gazeta Wyborcza</i> (liberal)                                                                                                                 |                                                    | -                                                                                      |
| Europe East   | Bulgaria    | <i>Дневник (Dnevnik): Dnevnik Analizi, Dnevnik Biznes, Dnevnik Bulgaria, Dnevnik Evropa, Dnevnik World, Dnevnik Zelen</i> (liberal-conservative) |                                                    | <i>Dnevnik</i> available only from June 9, 2017 onward                                 |
| Europe South  | Italy       | <i>La Stampa</i> (liberal)                                                                                                                       | <i>Corriere della Sera</i> (liberal-conservative)  | <i>La Stampa</i> available only until June 13, 2019 (2016 election only)               |
| Europe South  | Spain       | <i>El Pais</i> (center-left)                                                                                                                     | <i>El Mundo</i> (conservative)                     | -                                                                                      |
| Europe West   | France      | <i>Libération</i> (center-left)                                                                                                                  | <i>Le Figaro</i> (liberal-conservative)            | <i>Libération</i> available only from March 26, 2018 onward (excludes 2016 election)   |
| Europe West   | Netherlands | <i>De Volkskrant</i> (center-left)                                                                                                               | <i>De Telegraaf</i> (right-wing)                   | -                                                                                      |
| Europe Center | Germany     | <i>taz</i> (left-wing)                                                                                                                           | <i>Die Welt</i> (conservative)                     | -                                                                                      |
| Europe Center | Switzerland | <i>Tages-Anzeiger</i> (center-left)                                                                                                              | <i>Neue Zürcher Zeitung</i> (liberal-conservative) | <i>Tages-Anzeiger</i> no data between October 2, 2017 to November 10, 2017             |

Table 2 European newspapers chosen for analysis

| Area          | Country | Newspaper (political leaning)          |                               | Notes                                                                   |
|---------------|---------|----------------------------------------|-------------------------------|-------------------------------------------------------------------------|
| North America | USA     | <i>The New York Times</i> (skews left) | <i>USA Today</i> (middle)     | <i>The New York Times</i> corpus includes both print and online content |
| North America | Mexico  | <i>El Universal</i> (left)             | <i>Reforma</i> (center-right) | -                                                                       |

|                      |                   |                                                     |                                          |                                                                                                                                           |
|----------------------|-------------------|-----------------------------------------------------|------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------|
| <b>South America</b> | Brazil            | <i>Folha de São Paulo</i> (non-partisan, pluralist) | <i>O Globo</i> (conservative)            | <i>Folha de São Paulo</i> available only from December 23, 2019 onward (excludes 2016 election); the corpus includes only online content. |
| <b>Asia</b>          | China / Hong Kong | <i>China Daily</i> (left, state propaganda)         | <i>South China Morning Post</i> (middle) | -                                                                                                                                         |
| <b>Africa</b>        | Nigeria           | <i>Vanguard</i> (center-left)                       |                                          | -                                                                                                                                         |
| <b>Oceania</b>       | Australia         | <i>The Age</i> (center-left)                        | <i>The Australian</i> (center-right)     | -                                                                                                                                         |

*Table 3 Non-European newspapers chosen for analysis*

#### 4.1.5. Europe Center

Central Europe is represented by Germany and Switzerland. In Germany, *taz* (left-wing) and *Die Welt* (conservative) were selected. Both are available across all cycles. While *Süddeutsche Zeitung* and *Frankfurter Allgemeine Zeitung* would have been preferable due to their high circulation, neither is included in *Lexis Nexis*. Importantly, for the main study only the print corpus of *Die Welt* is used, in contrast to the test run, where *Welt (Online)* was included to boost sample size (see section 5.2.1). Switzerland is represented by *Tages-Anzeiger* (center-left) and *Neue Zürcher Zeitung* (liberal-conservative), both available across all cycles (exception: no data for *Tages-Anzeiger* for October 2, 2017 - November 10, 2017). Switzerland was chosen over Austria, as the two countries are of comparable size, but Switzerland contributes the additional perspective of a non-EU country.

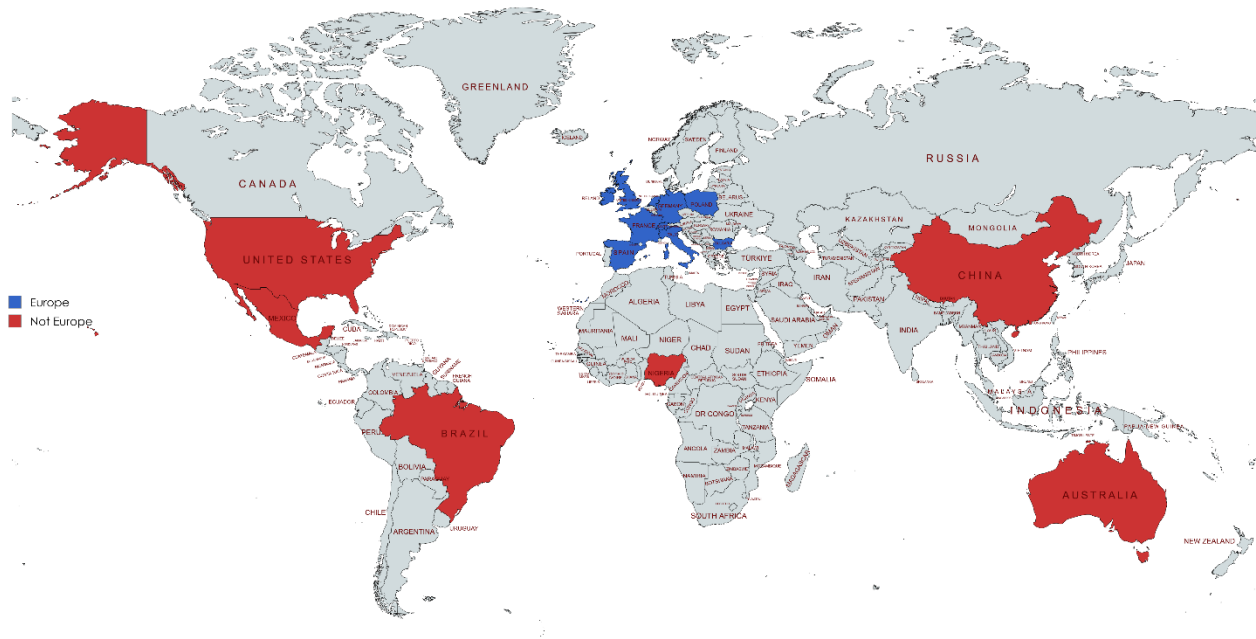


Figure 2 Distribution of European and non-European countries chosen for analysis (created with mapchart.net)

#### 4.1.6. North America

North America is represented by the United States and Mexico. For the United States, *The New York Times* (skews left) and *USA Today* (middle) were selected. While a more explicitly right-leaning national daily would have been preferable, *Lexis Nexis* does not provide access to such an outlet. The *New York Times* corpus includes both print and online articles, which are tagged in the database and will be filtered accordingly. For Mexico, *El Universal* (left) and *Reforma* (center-right) were selected, both available across all cycles. Mexico was included alongside the United States to ensure that North America was represented by two countries, and to introduce further linguistic diversity through Spanish-language material.

#### 4.1.7. South America

Brazil provides South American coverage, with *Folha de São Paulo* (non-partisan pluralist) and *O Globo* (conservative). While *Folha* is only available from December 2019 onward, excluding the 2016 election, *O Globo* is available for all cycles. The selection acknowledges that not all newspapers fit neatly into a left/right classification; *Folha*'s pluralist orientation provides a distinct perspective that broadens the comparative base.

#### 4.1.8. Asia

Asia is represented by China and Hong Kong. In China, the *People's Daily* – the most obvious candidate – was not available in *Lexis Nexis*. Instead, *China Daily* was chosen, which is published in English and functions as an official state outlet, classified as left-leaning due to its strict propaganda function. To provide balance, the *South China Morning Post* (middle) was selected as a counterpart. Based in Hong Kong, it has historically represented a freer press tradition than mainland sources. Both outlets are available for all cycles.

#### 4.1.9. Africa

Nigeria was selected to represent Africa. Here, *Vanguard* (center-left) was chosen as the most clearly classifiable daily available in *Lexis Nexis*. While *The Guardian (Nigeria)* is present in the database from 2022 onward, it was excluded in favor of *Vanguard*, which offers broader coverage across all cycles.

#### 4.1.10. Oceania

Finally, Oceania is represented by Australia, with *The Age* (center-left) and *The Australian* (center-right). Both are available across all cycles and together provide a clear left-right contrast in Australian daily journalism.

### 4.2. Time Frame

The timeframe of the study has been determined with the dual aim of comparability and feasibility. It focuses on the last three U.S. presidential races (2016, 2020, 2024), which are of particular relevance because Donald Trump, a candidate whose relationship with the press has been especially contentious (see section 2.5), participated in all three. This choice allows for a systematic comparison of his representation in the press with that of his rivals – Hillary Clinton, Joe Biden, and Kamala Harris – across different political contexts. Restricting the analysis to these three election cycles makes it possible to trace both continuity and change in reporting practices, while also ensuring that the material remains manageable in scope.

Within each election cycle, the study focuses on three anchor events: the televised debates, election day, and inauguration day. These events were selected because they are globally significant media spectacles, attracting heightened international coverage and offering natural points of comparison across different elections. They also provide steady reference points, since each occurs regularly in the cycle, and they represent moments when both candidates are highly visible in the news. To capture not only the events themselves but also their discursive buildup and aftermath, the timeframe extends to three days prior to and after each event. This ensures a corpus that reflects both anticipatory framing and retrospective evaluation, while keeping the dataset within manageable limits.

Some caveats must be acknowledged. Ideally, one might analyze full campaigns or even entire presidencies, but the scope of this project does not permit such breadth and depth. Focusing on selected dates is

therefore a methodological trade-off, privileging depth and comparability over comprehensiveness. At the same time, this limitation offers a potential avenue for future research, which could expand the analysis to cover longer periods or different types of political events. Similarly, while the reliance on specific dates and short reporting windows necessarily narrows the perspective, it also highlights moments when newspapers are most likely to reveal their ideological stances through headlines. Thus, the timeframe reflects a balance between feasibility and theoretical relevance, allowing for the analysis of discourse at moments of peak global attention while remaining sensitive to the practical constraints of data collection.

## 5. Methodology

The methodology of the study builds on the design of the corpus (see section 4.1) by specifying how the collected headlines were systematically analyzed. The analysis proceeded in several stages, beginning with the development and application of a coding scheme for variables such as bias, names as given, value judgement, and naming strategy. These procedures were first tested in pilot runs with both human and AI coders (see section 5.2), then implemented on the full dataset with AI coding (see section 0). Each stage of the process is described in detail to ensure transparency and reproducibility, from the operationalization of coding categories to the assessment of intercoder reliability.

### 5.1. Coding Scheme

#### 5.1.1. Bias

Bias in headlines was defined as an evaluative judgement directed at the candidate. The coding distinguished three main categories – neutral, positive, and negative. The guiding principle was that bias must directly concern the candidate, not merely policies, institutions, or external actors mentioned in the headline. The biased stance may be conveyed through verbs, adjectives, adverbs, or nouns with strong evaluative force.

##### 5.1.1.1. *Neutral Bias*

Headlines were coded as neutral, when they reported actions, statements, or events without evaluative language directed at the candidate. Neutral headlines could describe presidential duties, quotations, or factual developments. For example,

(52) Nördlich von Pjöngjang Nordkorea scheitert bei Raketentest – Trump reagiert<sup>32</sup> (*Welt (Online)*: April 29, 2017)

was coded as neutral, since it merely notes Trump’s reaction. Similarly,

(53) Donald Trump warns of a ‘major, major conflict’ with North Korea (*The Guardian*: April 28, 2017)

was treated as neutral: the negativity lies in the conflict itself, not in Trump’s person. An Italian example is

---

<sup>32</sup> North of Pyongyang North Korea fails missile test – Trump responds

(54) Trump taglia tasse in Usa, Wall Street vola poi ripiega<sup>33</sup> (*ANSA*: April 26, 2017),

where the positive and negative effects on the stock market cancel each other out, yielding a neutral evaluation.

#### 5.1.1.2. *Positive Bias*

Headlines were coded as positive when they attributed beneficial qualities or actions to the candidate, such as benevolence, competence, or constructive leadership. These framings often highlight helpfulness, praise, or opportunities created by the candidate. For instance,

(55) Abu Mazen, nuove opportunità di pace grazie a Trump<sup>34</sup> (*ANSA*: May 3, 2017)

portrays Trump as facilitating peace, a clearly positive role. Another case is

(56) US-Präsident Trump sagt Merkel telefonisch Hilfe zu<sup>35</sup> (*Welt (Online)*: April 25, 2017),

which may be read as positive because it presents Trump as offering assistance.

#### 5.1.1.3. *Negative Bias*

Headlines were coded as negative when they framed the candidate in unfavorable terms, whether through pejorative word choice, depiction of harmful actions, or hostile reactions from others. For example,

(57) US-Präsident Trump prahlt mit seinen Taten und ätzt gegen die Presse<sup>36</sup> (*Welt (Online)*: April 30, 2017)

is clearly negative due to the verbs *prahlen* ('boast') and *ätzen* ('lash out'). In English,

(58) Worried world urges Trump not to pull out of Paris climate agreement (*The Guardian*: May 7, 2017)

conveys negativity through the collective worry and urging directed at Trump's intentions. Another example,

(59) Trump insiste, ho chiesto a Gentiloni saldare conti Nato<sup>37</sup> (*ANSA*: April 24, 2017)

was coded as negative because the verb *insiste* ("insists") portrays Trump as forceful or demanding.

#### 5.1.1.4. *Ambiguous and Mixed Cases*

Some headlines presented ambiguity, where reasonable arguments could support either a neutral or a positive/negative interpretation. For example,

---

<sup>33</sup> Trump cuts taxes in the US, Wall Street soars then falls back

<sup>34</sup> Mahmoud Abbas, new opportunities for peace thanks to Trump

<sup>35</sup> US President Trump promises Merkel assistance in a telephone call

<sup>36</sup> US President Trump brags about his achievements and attacks the press

<sup>37</sup> Trump insists, I asked Gentiloni to settle NATO accounts

(60) Federal grant that Trump wanted to abolish is a lifeline for poor families (*The Guardian*: May 6, 2017)

could be coded as neutral (policy preference without judgement) or negative (portraying Trump as wanting to harm vulnerable groups). Similarly,

(61) Steuerreform Jetzt erfährt die Welt, was Trump wirklich kann<sup>38</sup> (*Welt (Online)*: April 25, 2017)

could be read as neutral (merely reporting his chance to demonstrate ability) or positive (suggesting that he is proving his competence). Another form of ambiguity arises when headlines combine both positive and negative framings. By default, such cases were treated as neutral, unless one evaluative direction was clearly dominant. For example,

(62) Trump übers Präsidentenamt: “Ich hatte schon Dinge, die härter waren”<sup>39</sup> (*Welt (Online)*: April 30, 2017)

might be neutral as a factual statement, but depending on perspective could be read positively (experience, competence) or negatively (arrogance, lack of seriousness).

Taken together, this typology ensures a transparent and systematic approach to identifying evaluative stances in headlines across languages and outlets. By distinguishing between positive, negative, and neutral cases, the coding scheme captures both overt judgements and more subtle or contested framings. The explicit treatment of ambiguous headlines acknowledges the interpretive flexibility inherent in political discourse and provides a safeguard against arbitrary decisions. This structured framework thus balances clarity with sensitivity to nuance, making it suitable for cross-linguistic comparison and for testing inter-coder reliability between human and AI coders.

#### 5.1.2. Name as Given

The second dimension of the coding scheme concerned the names as given, i.e., the exact form of reference used for the candidate in the headline. This variable was recorded verbatim to preserve the linguistic choices of the newspapers and to allow precise cross-checking with AI outputs. Every occurrence of the candidate’s name in a headline was recorded, even when it appeared more than once, and the entire string was retained in its original form. Thus, designations such as

(63) Harris

(64) Hillary Clinton

(65) Mr. Trump

(66) President Biden

or longer expressions like

(67) last-minute democratic candidate Kamala Harris

were all coded in full, without normalization or reduction. This decision was motivated by two considerations. First, exact transcription of naming practices makes it possible to analyze subtle cross-linguistic

---

<sup>38</sup> Tax reform Now the world is finding out what Trump is really capable of

<sup>39</sup> Trump on the presidency: “I’ve had things that were harder.”

and cross-outlet variation in how candidates are referred to, whether through bare surnames, titles, or extended role descriptions (see Section 2.3). Second, it provides a necessary basis for the variable value judgment (see Section 5.1.3), which evaluates whether the form of naming itself contains evaluative content. Treating the full string as the unit of analysis ensures that qualifiers embedded in the reference (e.g. *embattled Trump*, *Sleepy Joe Biden*) are preserved for later coding.

It must be noted, however, that the use of family names as search terms in *Lexis Nexis* introduces a structural limitation: references that do not contain the candidate's surname are unlikely to appear in the dataset. As a result, nicknames or role-based designations that omit the family name – such as *the Don*, *the Republican candidate*, or *Cackling Kamala* – may be underrepresented or absent altogether. This caveat does not undermine the consistency of the analysis but should be borne in mind when interpreting the distribution of naming strategies and value judgements across the corpus.

### 5.1.3. Value Judgement

A third dimension of the coding scheme focused on value judgement. This variable sought to capture whether the name as given itself contained evaluative content, such as pejorative or laudative descriptors attached to the candidate's name. The rationale for including this category is grounded in Kuypers (2014: 5) observation that labeling constitutes a powerful source of “hidden bias” in mainstream news. As Kuypers argues, extreme modifiers are disproportionately applied to conservatives and Republicans, with terms such as *right-wing*, *ultra-conservative* or similar formulations carrying implicit negative evaluations. Conversely, laudative naming, though less common, may also occur, for instance in depictions of a politician as a *peacemaker* or *savior*. By systematically coding value judgement, this study aims to test empirically whether such asymmetries appear in the selected international corpus, and whether they are attached differentially to Kamala Harris, Joe Biden, Hillary Clinton, or Donald Trump.

In practice, the coding of value judgement was restricted strictly to the name as given, rather than to the headline as a whole. This ensured a clear distinction between the evaluative force of a candidate's naming and the broader stance of the headline, which was captured under the separate category of bias. For example, the headline

(68) Casa Bianca, ‘molto buona’ la telefonata Trump-Putin<sup>40</sup> (*ANSA*: May 2, 2017)  
referred to Trump simply as *Trump*, which was coded as carrying no evaluative content. Similarly, the headline

(69) Malcolm Turnbull's meeting with Donald Trump delayed by healthcare bill (*The Guardian*: May 4, 2017)

contained the straightforward reference *Donald Trump*, again coded as *none*. Even when headlines contained strong evaluative language overall, such as

(70) USA Zwischen knallhart und konfus – Trumps Außenpolitik<sup>41</sup> (*Welt (Online)*: April 26, 2017)

the value judgement of the name as given remained *none*, since the evaluation applied not to the naming but to the policy described.

---

<sup>40</sup> White House: Trump-Putin phone call ‘very good’

<sup>41</sup> USA Between tough and confused – Trump's foreign policy

The test dataset analyzed during the AI trial runs contained no pejorative or laudative valued judgements in the name as given. All references were coded as *none*. Nevertheless, retaining this variable in the scheme is methodologically important. The absence of variation in the pilot material does not preclude the possibility that evaluative naming will occur in the larger datasets, especially during heated campaign moments where nicknames, extreme ideological labels, or laudatory descriptors may emerge. For instance, a possible headline such as

(71) Sleepy Joe Biden returns to campaign trail

would be coded as *pejorative*, since the nickname *Sleepy Joe* conveys an explicitly negative characterization. Conversely, a hypothetical Italian headline like

(72) Il pacificatore Trump guida i colloqui di pace<sup>42</sup>

would be coded as *laudative*, given that the evaluative label directly presents Trump in a positive light.

To safeguard against misclassification, all cases coded as containing a pejorative or laudative value judgement will be systematically double-checked by a human coder in the main study. This step ensures that any evaluative naming, should it appear, is reliably identified and distinguished from broader headline bias.

#### 5.1.4. Naming Strategies

The final dimension of the coding scheme concerned naming strategies, which describe how candidates are referred to in headlines. The categories follow van Leeuwen's (2008) typology, distinguishing between nomination (e.g., by proper names or titles), categorization (by roles or social attributes), and appraisal (by evaluative labels). This framework was selected because naming choices are never neutral: as Fowler (1991) and Taylor (2019) argue, even seemingly factual references carry ideological distinctions and contribute to shaping reader's perceptions.

In practice, the range of strategies observable in the dataset was constrained by the search procedure. Since only family names (*Trump*, *Clinton*, *Biden*, *Harris*) were used as search terms, the corpus predominantly yielded instances of formal nomination (surname only) and titled nomination (titles such as *Vice President Kamala Harris* or *Secretary Clinton*). Other strategies outlined by van Leeuwen, such as informal naming (*Joe*, *Donald*) or relation identifications (*Obama's vice president*), were unlikely to appear in this dataset and were indeed absent in the pilot analysis. Nevertheless, these categories were retained in the coding scheme to preserve systematic comparability and to allow for the possibility that such forms may occur in the expanded corpus, especially during more heated phases of election coverage.

Examples from the test dataset illustrate the main categories encountered:

- Formal nomination (surname only, ± honorific):

(73) Glass-Steagall-Gesetz Trump prüft die Aufspaltung großer Wall-Street-Banken<sup>43</sup> (*Welt (Online)*: May 2, 2017)

---

<sup>42</sup> Peacemaker Trump leads peace talks

<sup>43</sup> Glass-Steagall Act Trump considers breaking up large Wall Street banks

where the surname *Trump* appears in isolation;

- Semiformal nomination (given name + surname):  
(74) The lesson from Donald Trumps first 100 days: resistance is not futile (*The Guardian*: April 28, 2017)

using the full name *Donald Trump*:

- Titled nomination (honorification):  
(75) US-Präsident Trump sagt Merkel telefonisch Hilfe zu<sup>44</sup> (*Welt (Online)*: April 25, 2017)

which foregrounds Trump's official role as U.S. President.

While less frequent in the test material, other plausible forms demonstrate the broader scope of the coding scheme:

- Titled nomination (affiliations): e.g. *Bill Clintons Frau Hillary*<sup>45</sup>
- Categorization – functionalization: e.g. *procuratrice generale della California Harris*<sup>46</sup>
- Categorization – classification: e.g. *der New Yorker Milliardär*<sup>47</sup> (Trump)
- Categorization – relational identification: e.g. *Biden's running mate* (Harris)
- Categorization – physical identification: e.g. *die Frau mit dem Hosenanzug*<sup>48</sup> (Clinton)
- Appraisal: e.g. *der Held der MAGA-Vereinigung*<sup>49</sup> (Trump)

By systematically distinguishing naming strategies, the coding captures how candidates are constructed not only through evaluative language but also through the very terms of reference that position them socially and politically.

## 5.2. AI Test Runs

### 5.2.1. Data Collection and Sampling

The test run was designed to evaluate and compare human coding with coding performed by an artificial intelligence (AI) system to determine the most appropriate methodology for the main analysis. The rationale behind the test run was to establish whether AI-based coding could replicate, complement, or substitute human coding in systematic discourse analysis, and to test the stability of results across prompts and coding conditions. To achieve this, a carefully structured dataset of headlines in German, English, and Italian was created and analyzed. All data was collected via *Lexis Nexis*, a large database of international news outlets. The initial plan was to select headlines from three major newspapers, namely *Die Welt* (Germany), *The Guardian* (UK), and *Corriere della Sera* (Italy). These newspapers were chosen to ensure linguistic diversity and to allow testing of the AI's ability to process data in three distinct

---

<sup>44</sup> US President Trump promises Merkel assistance in a phone call

<sup>45</sup> Bill Clinton's wife Hillary

<sup>46</sup> Attorney General of California Harris

<sup>47</sup> The New York billionaire

<sup>48</sup> The woman with the pantsuit

<sup>49</sup> The hero of the MAGA-camp

languages. The goal for the test run was to obtain at least 50 headlines per language to provide a sufficient basis for comparative coding and reliability testing.

An initial search was conducted for the period of June 13-19, 2015, centered on June 16, 2015, the date on which Donald Trump formally announced his candidacy for the presidency. The search term was “Trump.” This time frame was chosen because the event was expected to trigger international reporting. However, the search yielded too few results for the chosen outlets. Even after extending the time frame to seven days before and after the announcement (June 9-23, 2015), the number of headlines was insufficient. Given the goal of ensuring a minimum of 50 headlines per language, the initial time frame was abandoned. A new period was then selected, namely Trump’s first 100 days in office. The symbolic milestone of 100 days occurred on April 30, 2017. To capture reporting around this event, the time frame was set to April 23 to May 7, 2017 (seven days before and after the milestone). This adjustment was expected to yield a larger number of relevant headlines across the target newspapers. In the German case, it was necessary to switch from *Die Welt* (print edition) to *Welt (Online)*, because the print edition alone produced too few hits. The search produced the following raw results: *Welt (Online)*: 192 results, *The Guardian* 275 results, *Corriere della Sera* 116 results. The results were downloaded in Excel format. An overview of the successive filtering steps for each outlet is presented in Table 4, with the detailed procedures described in the main text.

| Outlet                     | Raw results  | After duplicate removal | After filtering for “Trump” in headline | After removing non-Trump mentions & extra text |
|----------------------------|--------------|-------------------------|-----------------------------------------|------------------------------------------------|
| <i>Welt (Online)</i>       | 192          | 162                     | 50                                      | 45                                             |
| <i>The Guardian</i>        | 275          | 274                     | 99                                      | 70                                             |
| <i>Corriere della Sera</i> | 116          | 116                     | 15                                      | - (too few, replaced)                          |
| <i>ANSA</i>                | 594          | 479                     | 283 (random sample of 70)               | 52                                             |
| <b>Total</b>               | <b>1,177</b> | <b>1,031</b>            | <b>447</b>                              | <b>167</b>                                     |

Table 4 Corpus search results and filtering

Since *Lexis Nexis* returns results for the entire article, i.e., not necessarily the headline, duplicates and irrelevant entries needed to be removed. After the first round of duplicate removal, the datasets contained: *Welt (Online)*: 162 (30 duplicates removed), *The Guardian* 274 (1 duplicate removed), *Corriere della Sera* 116 (no duplicates). The high number of duplicates in *Welt (Online)* is most likely attributable to the corpus including both the online and the print versions of the same article. The next step was to filter for headlines only, since *Lexis Nexis* also indexes hits where the search term appears in the body text but not in the headline. Only headlines mentioning “Trump” were retained. After this filtering step, the datasets contained: *Welt (Online)*: 50 results, *The Guardian* 99 results, *Corriere della Sera* 15 results. Because the Italian dataset was too small, an additional source was added: the Italian news agency *ANSA*. The *Lexis Nexis* search for *ANSA* yielded 594 entries, of which 115 were duplicates. After duplicate removal, 479 entries remained. Filtering for “Trump” in the headline yielded 283. To reduce the dataset to a manageable size, a random sample of 70 headlines was drawn using the Excel formula =RAND(). Each row was assigned a random number, sorted in ascending order, and the first 70 entries were selected. Next, additional cleaning was performed. Mentions that did not refer to Donald J. Trump but to other individuals or entities (e.g., *Ivanka Trump*, *Melania Trump*, *Trump Tower*, *Trump rally*, ...) were removed. This resulted in the exclusion of nine entries from *Welt (Online)*, eight from *The Guardian*, and one from *ANSA*. Furthermore, both *The Guardian* and *ANSA* files included not only the headline but also the first lines of the article. These were manually removed to ensure that only headlines were analyzed.

After this filtering step, the final datasets comprised *Welt (Online)*: 45 headlines, *The Guardian*: 70 headlines, and *ANSA*: 52 headlines.

The headlines were retained in their original languages. They were not translated, normalized, or corrected for spelling or grammar. This decision was taken to avoid introducing interpretive bias and to ensure that the coding reflected the linguistic choices of the original texts. The final datasets were organized into three separate Excel files (German, English, Italian). Each file included the headline, publication, date, and outlet. These metadata were retained for organizational purposes but not used in the analysis. The final sample size across the three languages was 167 headlines.

### 5.2.2. Human Coding

All human coding was conducted by the author. This ensured consistency in coding across the three languages. For Italian, occasional consultation of an online dictionary was necessary. When headlines referred to individuals or events unfamiliar to the author, clarifications were sought using ChatGPT, although the coding decisions were ultimately made independently.

The categories applied in the human coding process are summarized in Table 5, while their theoretical grounding and application are discussed in the main text.

| Variable               | Coding categories                      | Example(s)                                              |
|------------------------|----------------------------------------|---------------------------------------------------------|
| <b>Bias</b>            | positive / negative / neutral          | Trump praised for... → positive                         |
| <b>Name as given</b>   | Exact string from headline             | Donald Trump, President Trump                           |
| <b>Value judgement</b> | none / pejorative / laudative          | Embattled Trump → pejorative                            |
| <b>Naming strategy</b> | Based on van Leeuwen (2008) categories | <i>Trump</i> → formal; <i>Donald Trump</i> → semiformal |

Table 5 Coding scheme overview

The coding scheme was implemented in four columns for each headline. The first column captured bias, defined as the overall evaluative stance toward Donald J. Trump and coded as *positive*, *negative*, or *neutral*. The second column recorded the name as given in the headline, meaning the exact wording used to refer to Trump, such as *Trump*, *Donald Trump*, *President Trump*, or modified forms. The third column concerned value judgement, which assessed whether the chosen form of naming carried an evaluative content and was coded as *none*, *pejorative*, or *laudative*. Finally, the fourth column identified the naming strategy, drawing on van Leeuwen's (2008) framework (see section 2.3). In this dataset, the most common strategies included *Formal* (surname only, with or without honorific), *Semiformal* (given name plus surname), and *titulated nomination* (titles or ranks). While other strategies listed by van Leeuwen were theoretically possible, they appeared only rarely in the present material. This scheme was developed in advance, based on prior work by the author on headline analysis in English and German (Holler 2023). It had been refined through earlier experimentation, ensuring that the categories were theoretically grounded and practically usable. Coding was recorded directly into the Excel files. Each headline was coded once by the author. The AI coding results were added later to the same Excel files but in separate sheets, ensuring that there was no risk of AI influencing the human coding.

### 5.2.3. AI Coding

The AI analysis was conducted using the GPT-5-chat-latest model in the OpenAI Playground. This model was chosen because it allowed the adjustment of key parameters, namely temperature, top\_p, and max

tokens. For all runs, the parameters were fixed at *temperature* = 0, *top\_p* = 1, and *max tokens* = 16,384. The settings were chosen to ensure reproducibility, since temperature values above zero would introduce randomness into the outputs. Three separate runs were carried out. They are outlined in Table 6.

| Run | Prompt used | Purpose                              | Notes                                  |
|-----|-------------|--------------------------------------|----------------------------------------|
| 1   | Prompt 1    | Main comparison with human coding    | Full dataset (per language)            |
| 2   | Prompt 1    | Replication for stability check      | Same setup as Run 1                    |
| 3   | Prompt 2    | Sensitivity check (prompt variation) | Focus-actor schema; stricter filtering |

Table 6 AI run design

Run 1 served as the main analysis: headlines were coded with Prompt 1, and the resulting dataset was used for direct comparison with the human coding. Run 2 replicated this procedure with Prompt 1 on the same dataset to test the stability of the model’s results across repeated runs. Run 3 employed Prompt 2, which differed slightly in structure and placed stronger emphasis on coding Donald J. Trump as the sole “focus actor,” thereby allowing an assessment of the system’s sensitivity to prompt design. For each run, the headlines were submitted in batches by language. Since the datasets were relatively small (45-70 headlines each), all headlines for a given language could be processed simultaneously without exceeding model limits. Both prompts required the model to output results in JSON format. These JSON outputs were then converted to CSV format and copy-pasted into the Excel sheets alongside the human-coded results. Prompt 1 instructed the model to code each headline into the same four categories used for human coding (*bias*, *name as given*, *value judgement*, *naming strategy*) plus an exclusion flag and short notes. Prompt 2 used a slightly different schema: it treated Donald J. Trump as the only focus actor and represented a more restrictive coding regime.

The following two prompts were used for AI coding:

### 5.2.3.1. Prompt 1

#### Developer message

You are a meticulous content coder. You will analyze newspaper headlines in their original language (German, English, Italian). Infer labels only from the headline text. If a headline does not refer to Donald J. Trump, label it `exclude=true`.

#### Prompt message

##### TASK

Code each headline according to these fields:

- *bias*: one of {positive, negative, neutral}. Judge the overall evaluative stance toward Donald J. Trump conveyed by the headline wording.
- *name\_as\_given*: copy EXACTLY the main form used to refer to Donald J. Trump in the headline (e.g., “Trump”, “Donald Trump”, “President Trump”). If multiple forms appear, choose the **\*\*primary\*\*** one. Mention the other forms in notes.
- *value\_judgement*: one of {none, pejorative, laudative}. This captures explicit evaluative modifiers in the name itself (e.g., “embattled Trump” → pejorative, “peacemaker Trump” □ laudative).
- *naming\_strategy*: Use only one of:
  - “Nomination: Formal (surname only, ± honorific)” (e.g., “Trump,” “Mr. Trump”)
  - “Nomination: Semiformal (given name + surname)” (e.g., “Donald Trump,” Donald J. Trump”)
  - “Nomination: Informal (given name only)” (e.g., “Donald,” “Don”)
  - “Nomination: Non-proper-name nomination (unique role in context)” (e.g., “the President,” “the former President,” “the Republican presidential candidate”)

- “Nomination: Titled nomination – honorification (titles/ranks)” (e.g., “President Trump,” “Former President Trump”)
- “Nomination: Titled nomination – affiliations (kinship/personal/work ties)” (e.g., “Donald Trump, Ivanka Trump’s father,” “Melania Trump’s husband, Donald”)
- “Categorization: Functionalization (what they do/roles & activity nouns)” (e.g., “Trump the businessman,” “Trump the incumbent”)
- “Categorization: Identification – classification (social categories)” (e.g., “Trump the wealthy New Yorker,” “the billionaire Trump”)
- “Categorization: Identification – relational (defined by relations)” (e.g., “Barron Trump’s father,” “Biden’s 2020 opponent”)
- “Categorization: Identification – physical (salient physical traits)” (e.g., “the orange man,” “the blond-haired man”)
- “Appraisal: Evaluative labels found in discourse” (e.g., “the polarizer,” “the villain to his critics”)
- exclude: {true,false}. Set true if the headline’s “Trump” refers to someone/something other than **Donald J. Trump**, or if he is **not mentioned**.
- notes: short justification (max 25 words).

#### RULES

- Do not translate. Judge in the headline’s language.
- If sentiment is unclear, use neutral.
- If the form of reference is unclear, pick the most prominent mention; if none, set exclude=true.
- Output **one JSON object per headline**, exactly with keys: id, bias, name\_as\_given, value\_judgement, naming\_strategy, exclude, notes.

#### OUTPUT

Return only the JSON lines, one per headline, e.g.:

```
{“id”: 1, “bias”: “neutral”, “name_as_given”: “Trump”, “value_judgement”: “none”, “naming_strategy”: “Formal (surname only, ± honorific)”, “exclude”: false, “notes”: “matter-of-fact wording”}
{“id”: 2, “bias”: “...”, “name_as_given”: “...”, “value_judgement”: “...”, “naming_strategy”: “...”, “exclude”: false, “notes”: “...”}
...
```

### 5.2.3.2.Prompt 2

#### Developer message

You are an impartial, multilingual discourse coder for news headlines. Your job is to classify each headline strictly from its text without adding outside knowledge. Work in any language the headline uses.  
Return JSON only (no markdown, no prose). Follow the schema and rules exactly.  
Focus actor filter: Only report results for Donald J. Trump (the focus actor). Ignore other actors entirely in all fields except brief mentions inside notes when needed for clarity.

#### Prompt message

For each headline, output an object with exactly these fields:

1. “bias\_by\_entity” — an object containing only one entry for the focus actor if and only if the headline refers to them. Key must be the canonical name “Donald J. Trump” and value one of “positive” | “negative” | “neutral” describing the stance toward the focus actor in the headline. If the focus actor is not mentioned at all, output an empty object {}.
2. “names\_as\_given” — an array of the exact surface strings as they appear in the headline for the focus actor only, including qualifiers/adjectives/titles/affiliations (e.g., “President Trump”, “the billionaire Trump”). Do not include other actors. Do not translate or normalize.
3. “value\_judgment” — one of “pejorative” | “laudative” | “both” | “none”, considering only labels that target the focus actor. If evaluative labels apply to someone else or a different group, treat as “none” for this field (explain in notes if helpful).
4. “naming\_strategy” — exactly one label chosen based on the focus actor’s most central mention (see Choosing rules below). Use only one of:
  - “Nomination: Formal (surname only, ± honorific)”
  - “Nomination: Semiformal (given name + surname)”
  - “Nomination: Informal (given name only)”
  - “Nomination: Non-proper-name nomination (unique role in context)”

- “Nomination: Titled nomination – honorification (titles/ranks)”
- “Nomination: Titled nomination – affiliations (kinship/personal/work ties)”
- “Categorization: Functionalization (what they do/roles & activity nouns)”
- “Categorization: Identification – classification (social categories)”
- “Categorization: Identification – relational (defined by relations)”
- “Categorization: Identification – physical (salient physical traits)”
- “Appraisal: Evaluative labels found in discourse”

5. “notes” — a free text string for clarifications (may be empty). Use this for: (a) noting that the focus actor is not present and output is intentionally empty; (b) explaining that evaluative labels target someone else; (c) listing additional naming strategies that appeared but were not selected; (d) language specific disambiguation.

If a single headline is provided, return a single JSON object. If multiple headlines are provided, return a JSON array of objects in the same order.

Coding rules

- Use only the headline text. Do not infer facts or context beyond the words shown.
- Multilingual: Detect the language automatically. Do not translate; preserve original strings in `names_as_given`.
- Focus actor only: Ignore other actors/parties entirely in `bias_by_entity` and `names_as_given`. Do not include them as keys or strings. You may mention them briefly in notes if they determine stance toward the focus actor.
- Actor detection for Trump: Treat as the focus actor only when the headline uses an unambiguous form such as “Trump”, “Donald Trump”, “President Trump”, or a role clearly tied to Trump (e.g., “the former U.S. president” when the headline context makes Trump the only plausible referent). Do not treat generic “Donald” as Trump unless the same headline also includes “Trump”.
- Value judgment determination: Consider only evaluative labels that apply to the focus actor (e.g., nicknames like “Sleepy Joe” do not apply to Trump; ignore them). Party descriptors alone are not evaluative.
- Choosing a single naming strategy (focus actor based): Select the strategy that best represents the primary mention of Donald J. Trump. Heuristics:

1. Use the mention that is syntactically/topically central (subject or leading NP) referring to Trump.

2. If multiple mentions refer to Trump, choose the earliest.

- Quoted speech & framing: If the headline reports someone else’s evaluation of Trump, code stance toward Trump accordingly and set `value_judgment` as pejorative/laudative/both as applicable.
- Headlines without Trump: If Trump is not mentioned at all, return an object with `bias_by_entity`: {} and `names_as_given`: [], set `value_judgment` to “none”, choose `naming_strategy` as “Nomination: Non-proper-name nomination (unique role in context)” only if a role phrase clearly stands in for Trump; otherwise set `naming_strategy` to “Categorization: Functionalization (what they do/roles & activity nouns)”. In both cases, explain in notes that the focus actor is not present.

- Multiple headlines input: Treat each non empty line as a separate headline.

- No extra fields beyond the five specified keys.

Output format

- Strictly JSON. No backticks, no markdown, no explanations. Ensure valid UTF 8 and valid JSON.

#### 5.2.4. Comparison of Human and AI Coding

The coded datasets were then imported into SPSS for systematic comparison. In Excel, each coder’s results had been stored in separate sheets. In SPSS, the data was reorganized into adjoining columns with descriptive headings, making it possible to compare coders side by side. Before conducting statistical tests, it was necessary to ensure consistent terminology. Krippendorff’s Alpha (Hayes & Krippendorff 2007) requires identical category labels across coders. While “bias,” “name as given,” and “value judgement” were already consistent, the “naming strategy” field required minor adjustments. For example, AI output sometimes used “Nomination: Formal (surname only, ± honorific),” while the human coding used “Formal (surname only, ± honorific).” Such wording discrepancies were standardized. Crucially, no substantive coding decisions were altered; only superficial wording was harmonized. Krippendorff’s Alpha (nominal) was calculated for each variable (bias, name as given, value judgement, naming strategy) and for each dataset (German, English, and Italian). The procedure was conducted for each of the three comparisons: human vs. AI (main run, Prompt 1), AI vs. AI (replication run, Prompt 1), human vs. AI (sensitivity run, Prompt 2). To provide robust estimates, confidence intervals were calculated using 1,000

bootstrap samples. To account for unequal sample sizes across languages, all datasets were eventually pooled to generate an overall global estimate of agreement. The results are presented in Table 8.

| Alpha Value            | Reliability                                                |
|------------------------|------------------------------------------------------------|
| $\alpha = 1$           | Indicates perfect agreement among raters                   |
| $\alpha \geq 0.80$     | Indicates a reliable rating                                |
| $\alpha [0.67 - 0.79]$ | Indicates moderate agreement                               |
| $\alpha < 0.67$        | Indicates poor agreement                                   |
| $\alpha = 0$           | Indicates no agreement; similar to a random rating pattern |
| $\alpha < 0$           | Indicates systematic disagreement                          |

Table 7 How to interpret Krippendorff's Alpha (Marzi, Balzano & Marchiori 2024: 3)

The reliability analysis revealed variation across variables, languages, and runs. For the category *bias*, agreement between human and AI coding in the main run (Prompt 1) was moderate to high overall ( $\alpha = 0.78$ , 95% CI [0.68, 0.85]), with particularly strong results for German ( $\alpha = 0.83$ ) and English ( $\alpha = 0.81$ ), but weaker reliability for Italian ( $\alpha = 0.67$ ). When comparing the two AI runs using the same prompt, reliability rose to  $\alpha = 0.92$  overall, suggesting that the model was internally consistent. In contrast, agreement between human coding and AI coding with the second prompt was slightly lower ( $\alpha = 0.72$  overall), reflecting some sensitivity to prompt design.

For *name as given*, reliability was consistently very high. Human and AI coding in the main run agreed almost perfectly ( $\alpha = 0.98$  overall), with both AI runs producing identical results ( $\alpha = 1.0$ ). Agreement remained strong under the second prompt ( $\alpha = 0.95$ ). This indicates that both human and AI coders were highly consistent in recording the exact form of reference used for Donald Trump across all languages.

The variable *value judgement* proved more difficult. In the German and English datasets, nearly all headlines were coded as “none.” As a result, Krippendorff's Alpha returned a value of 0 with a confidence interval spanning -1 to 1, indicating that the lack of category variation made a reliability estimate impossible. In the Italian dataset, agreement was complete but again without variation, which also prevented calculation of Alpha. The observed lack of agreement between human and AI coding under Prompt 2 for value judgement in English ( $\alpha = -0.01$ , 95% CI [-0.65, 0.46]) and Italian ( $\alpha = -0.03$ , 95% CI [-1, 0.74]) is unlikely to reflect a fundamental inability of AI to perform this task, but rather highlights the sensitivity of results to prompt design. This finding underscores the importance of testing alternative prompts, evaluating them against human coding, and refining instructions before drawing substantive conclusions in an AI-assisted content analysis.

For *naming strategy*, agreement was once again high. In the main run, the overall Alpha was 0.95, while AI vs. AI reliability was perfect ( $\alpha = 1.0$ ). Agreement remained equally high between human and AI under the second prompt ( $\alpha = 0.95$  overall). For the Italian dataset specifically, all headlines referred to Donald Trump simply as “Trump,” and the naming strategy was therefore consistently coded as “Formal (surname only,  $\pm$  honorific).” Since both human and AI coders selected the same category in every instance, Alpha could not be computed due to lack of variations. This uniformity, however, indicates complete agreement.

Taken together, these findings suggest that AI coding can achieve near-perfect agreement with human coders for structural variables such as *name as given* and *naming strategy*, while agreement for more interpretive variables such as *bias* and especially *value judgement* is more modest and strongly influenced by prompt design. Nonetheless, the overall Krippendorff's Alpha of 0.78 for *bias* indicates a sufficiently

high level of reliability to justify extending the analysis to larger headline datasets with confidence. At the same time, the weaker and often indeterminate results for *value judgment* caution against over-reliance on this variable in subsequent analysis. To counter this limitation, all value judgements that are not coded as “none” will be systematically double-checked by a human coder in the next stages of analysis.

#### *5.2.4.1. Nature of Disagreements between Human and AI Coding*

The discrepancies between human and AI coding in the main test run were relatively few in number (4 cases in German, 7 in English, 10 in Italian out of 167 headlines in total), but they reveal important insights into how bias can be interpreted differently. Three main patterns emerged. First, there was a distinction between a narrow and a broad definition of bias. In several cases, the AI coded headlines as negative because they concerned policies or associations connected to Donald Trump, whereas the human coder judged them to be neutral since the criticism was not directed at Trump personally, as in these examples:

(76) Mar-a-Lago US -Außenministerium macht Werbung für Trumps Privatklub<sup>50</sup> (*Welt (Online)*: April 25, 2017)

(77) Trump all'AP, possibile usciremo da accordo con Iran<sup>51</sup> (*ANSA*: April 24, 2017)

Second, the discrepancies reflected differences in how ambiguous positive framings were treated. The human coding sometimes leaned toward a positive interpretation of headlines that could be read as praise or benevolence, such as

---

<sup>50</sup> Mar-a-Lago US State Department promotes Trump's private club

<sup>51</sup> Trump to AP: We may withdraw from agreement with Iran

| Variable               | Language | Human vs. AI (main run)                | AI vs. AI (same prompt)                | Human vs. AI (second prompt)           |
|------------------------|----------|----------------------------------------|----------------------------------------|----------------------------------------|
| <b>Bias</b>            | German   | $\alpha = 0.83$ [0.66, 0.96]           | $\alpha = 1$ [1, 1]                    | $\alpha = 0.73$ [0.54, 0.88]           |
|                        | English  | $\alpha = 0.81$ [0.67, 0.95]           | $\alpha = 0.87$ [0.76, 0.97]           | $\alpha = 0.78$ [0.61, 0.92]           |
|                        | Italian  | $\alpha = 0.67$ [0.47, 0.83]           | $\alpha = 0.91$ [0.78, 1]              | $\alpha = 0.63$ [0.44, 0.81]           |
|                        | Overall  | $\alpha = 0.78$ [0.68, 0.85]           | $\alpha = 0.92$ [0.85, 0.97]           | $\alpha = 0.72$ [0.62, 0.81]           |
| <b>Name as is</b>      | German   | $\alpha = 0.94$ [0.82, 1]              | $\alpha = 1$ [1, 1]                    | $\alpha = 0.89$ [0.72, 1]              |
|                        | English  | $\alpha = 1$ [1, 1]                    | $\alpha = 1$ [1, 1]                    | $\alpha = 1$ [1, 1]                    |
|                        | Italian  | n/a (complete agreement, no variation) | n/a (complete agreement, no variation) | n/a (complete agreement, no variation) |
|                        | Overall  | $\alpha = 0.98$ [0.90, 1]              | $\alpha = 1$ [1, 1]                    | $\alpha = 0.95$ [0.88, 1]              |
| <b>Value judgement</b> | German   | $\alpha = 0$ [-1, 1]                   | $\alpha = 1$ [1, 1]                    | $\alpha = 0$ [-1, 1]                   |
|                        | English  | $\alpha = 0$ [-1, 1]                   | $\alpha = 1$ [1, 1]                    | $\alpha = -0.01$ [-0.65, 0.46]         |
|                        | Italian  | n/a (complete agreement, no variation) | n/a (complete agreement, no variation) | $\alpha = -0.03$ [-1, 0.74]            |
|                        | Overall  | $\alpha = 0$ [-1, 1]                   | $\alpha = 1$ [1, 1]                    | $\alpha = -0.06$ [-0.56, 0.39]         |
| <b>Naming strategy</b> | German   | $\alpha = 0.89$ [0.72, 1]              | $\alpha = 1$ [1, 1]                    | $\alpha = 0.89$ [0.71, 1]              |
|                        | English  | $\alpha = 1$ [1, 1]                    | $\alpha = 1$ [1, 1]                    | $\alpha = 1$ [1, 1]                    |
|                        | Italian  | n/a (complete agreement, no variation) | n/a (complete agreement, no variation) | n/a (complete agreement, no variation) |
|                        | Overall  | $\alpha = 0.95$ [0.88, 1]              | $\alpha = 1$ [1, 1]                    | $\alpha = 0.95$ [0.88, 1]              |

Table 8 Intercoder reliability (Krippendorff's Alpha) for human and AI coding across variables, languages, and runs (confidence intervals based on 1000 bootstrap samples)

(78) US-Präsident Trump sagt Merkel telefonisch Hilfe zu<sup>52</sup> (*Welt (Online)*: April 25, 2017)

(79) Steuerreform Jetzt erfährt die Welt, was Trump wirklich kann<sup>53</sup> (*Welt (Online)*: April 25, 2017)

while the AI defaulted to a neutral reading in the absence of overtly positive evaluative markers. Third, some disagreements stemmed from inherently ambiguous or multi-layered wording, where both coding choices could be defended as valid. Headlines that presented competing perspectives or contained phrasing such as

(80) One nation, two Trumps: America as divided as ever after first 100 days (*The Guardian*: April 27, 2017)

(81) Reform von Obamacare Trumps Sieg ist eigentlich eine List der Geschichte<sup>54</sup> (*Welt (Online)*: May 5, 2017)

were interpreted differently depending on which nuance was foregrounded.

These patterns show that the divergences were not due to misinterpretations or errors by the AI, but rather to different but plausible coding logics. They underscore the fact that bias is not always clear-cut and that intercoder disagreement, whether between humans or between human and AI, often reflects the interpretive flexibility inherent in political discourse. For the purpose of this study, the relatively small number of disagreements and the overall reliability score (Krippendorff's Alpha = 0.78) justify proceeding with larger-scale AI coding. At the same time, the identified areas of ambiguity will inform the subsequent analysis, where cases of potential evaluative uncertainty can be flagged and, where necessary, double-checked by a human coder.

#### 5.2.4.2. Summary of Methodological Design

The test run proceeded in a clearly defined sequence. Headlines were sampled systematically from *Lexis Nexis*, filtered, cleaned, and organized into three language-based datasets. Human coding was carried out with a theoretically grounded scheme developed in advance and refined through prior research. AI coding was conducted with GPT-5 in three runs, using two prompts to test reproducibility and sensitivity. The comparison of human and AI results was performed in SPSS using Krippendorff's Alpha with bootstrapped confidence intervals. This methodological design ensured transparency, reproducibility, and robustness. The detailed documentation of every step – from corpus selection and sampling to coding procedures and statistical analysis – provides a solid foundation for evaluating the potential role of AI coding in systematic discourse analysis of multilingual newspaper headlines. Building on these results, the subsequent part of the thesis extends the analysis to larger datasets of headlines drawn from different timeframes, newspapers, and languages. Having established through the test run that AI coding can achieve a satisfactory degree of reliability – particularly for structural variables such as bias, naming, and naming strategies – the AI will be applied to these expanded corpora to trace patterns of representation across a broader range of contexts. This step allows the methodological insights gained in the pilot phase to be scaled up for substantive analysis, while maintaining safeguards such as human verification of value judgement where necessary.

---

<sup>52</sup> US President Trump promises Merkel assistance in a phone call

<sup>53</sup> Tax reform Now the world is learning what Trump is truly capable of

<sup>54</sup> Reform of Obamacare Trump's victory is actually a trick of history

### 5.3. Main Data Analysis

The main analysis extended the validated procedures of the test run (see section 5.2) to the complete international corpus of headlines. Since the pilot phase demonstrated sufficient reliability for bias (Krippendorff's  $\alpha = 0.78$  overall) and near-perfect agreement for structural variables such as naming and naming strategies, the AI model was applied to the full dataset without further human comparison. The coding scheme, categories, and AI model remained unchanged to ensure comparability across datasets. All coding was performed with the GPT-5-chat-latest model using Prompt 1 (see section 5.2.3.1), adapted for each candidate by substituting the focus name (Trump, Clinton, Biden, Harris). Apart from this substitution, the prompt and technical parameters (temperature = 0, top\_p = 1, max tokens = 16,384) remained identical to those used in the pilot phase. As in the pilot, all headlines were analyzed in their original languages, with spelling and grammar preserved exactly as published. Metadata on outlet and publication date were recorded, and tags for print/online content were retained where available. This continuity guarantees methodological consistency between the test run and the main analysis.

#### 5.3.1. Filtering Procedures

The filtering of raw datasets followed the same general steps as in the test run, with minor adjustments to increase efficiency given the much larger volume of material. Since some newspapers included not only the headline but also part of the lead text in the headline field, a character limit was imposed. Using the Excel formula =LEFT(A2; 100), all entries were truncated to 100 characters. This value was chosen on the basis of the pilot dataset, where headlines averaged around 64 characters in length; the slightly higher limit ensured that longer headlines would still be captured while eliminating irrelevant lead text. After truncation, duplicates were removed, and instances of the search terms (#Trump#, #Clinton#, #Biden#, #Harris#) were identified through text-to-columns processing. Any lines not containing the relevant candidate were excluded.

A further step concerned the distinction between print and online content. Where tagging was available, only print headlines were retained. This was possible for *The New York Times*, which provides explicit tagging, but not for *The Guardian*, whose dataset merges print and online material in *Lexis Nexis*. As a result, the *New York Times* corpus contains print-only headlines, while the *Guardian* corpus includes both print and online content.

#### 5.3.2. Dataset Size and Composition

The main data collection produced a large set of headlines across the three election cycles. In total, 55,634 entries were initially downloaded from *Lexis Nexis* before filtering. After applying the filtering procedures described in section 5.2.1, the dataset was reduced to 11,096 candidate-focused headlines. This step ensured that only relevant headlines were retained, while duplicates and extraneous lead text were removed. All results are presented in Table 9.

|                | Raw dataset   | Filtered dataset |
|----------------|---------------|------------------|
| <b>2016</b>    | 18,253        | 3,695            |
| <b>Clinton</b> | 6,750         | 734              |
| <b>Trump</b>   | 11,503        | 2,961            |
| <b>2020</b>    | 20,082        | 3,335            |
| <b>Biden</b>   | 8,083         | 1,269            |
| <b>Trump</b>   | 11,999        | 2,066            |
| <b>2024</b>    | 17,299        | 4,066            |
| <b>Biden</b>   | 1,222         | 265              |
| <b>Harris</b>  | 4,432         | 543              |
| <b>Trump</b>   | 11,645        | 3,258            |
| <b>Total</b>   | <b>55,634</b> | <b>11,096</b>    |

Table 9 Number of headlines retrieved before and after filtering, by election cycle and candidate

For the 2016 cycle, 18,253 raw entries were retrieved, but only 3,695 qualified headlines remained after filtering (2,961 referring to Trump, 734 to Clinton). In 2020, 20,082 raw entries were collected, of which 3,335 headlines were retained (2,066 referring to Trump, 1,269 to Biden). The 2024 cycle produced 17,299 raw entries, which after filtering yielded 4,066 usable headlines (3,258 referring to Trump, 265 to Biden, and 543 to Harris). Out of the remaining 11,096 headlines, the AI analysis identified an additional 157 headlines that did not refer to the candidates in question but, for example, their spouses, resulting in a total of 10,939 headlines that referred to Clinton, Biden, Harris, and Trump. The substantial reduction between the raw and final datasets reflects the same issues encountered during the pilot phase – namely, the presence of duplicates, irrelevant items, and extended lead text that was not part of the actual headline. While this reduction was pronounced, it ensured that the final corpus consisted solely of headlines directly relevant to the candidates under study, thereby providing a consistent and comparable dataset across election cycles and candidates.

### 5.3.3. Reliability and Limitations

The reliability of the coding procedures was supported by the pilot phase, where intercoder agreement between human and AI coding reached a Krippendorff's  $\alpha$  of 0.78 for bias and near-perfect agreement for structural variables such as naming and naming strategies. This level of reliability provides confidence that the AI model could be applied consistently to the full dataset without further human comparison. Throughout the main coding process, individual lines were periodically cross-checked by a human coder, and the AI outputs were consistently accurate.

At the same time, several limitations must be acknowledged. First, corpus composition was constrained by database availability. While *The New York Times* corpus included online content that could be removed during the filtering process, the *Guardian* corpus necessarily combines print and online content, as *Lexis Nexis* does not differentiate between the two. In addition, some high-circulation newspapers, such as *Le Monde* or *Süddeutsche Zeitung*, could not be included because they were not accessible through the database. Second, the distribution of candidate coverage is highly uneven. Donald Trump received nearly three times as many headlines as his Democratic competitors combined, reflecting international patterns of news selection. This imbalance is itself an instance of coverage bias (D'Alessio & Allen 2000), but it also means that results for Clinton, Biden, and Harris rest on smaller datasets and must be interpreted with caution.

Finally, while AI-based coding ensured consistent treatment across thousands of headlines, it remains an automated process. Subtle ambiguities in meaning may always be open to interpretation. However, human spot-checking throughout the analysis confirmed that the AI's coding decisions were sound, and the consistent application of the same coding scheme across all datasets mitigates the risk of systematic distortion.

Taken together, the corpus construction, transparent filtering, and pre-registered AI-assisted coding pipeline provide a consistent basis for hypothesis testing. The pilot established acceptable intercoder reliability for stance ( $\alpha = .78$  overall) and near-perfect agreement on structural variables (names and naming strategy), while the main run applied a fixed prompt and deterministic coding across all languages, outlets, and cycles. With these safeguards in place – and with all headlines analyzed in their original language and reported using identical categories – the study now turns to the empirical results. Section 6 reports each hypothesis in turn, pairing descriptive distributions (tables/figures) with the relevant statistical tests and a brief verdict (Hypothesis supported/partially supported/not supported).

## 6. Results

This section presents the outcomes of twelve preregistered hypotheses that trace headline bias across candidates, outlets, time, and European subfields of comparison. For each hypothesis, the subsection restates H and H0, reports descriptive counts and proportions, provides the inferential test (with effect sizes where applicable), and closes with a concise verdict. The sequence moves from overall stance distribution (H1) to candidate-specific asymmetries (H2) and ideological alignment (H3), then to temporal dynamics (H4) and the linguistic micro-mechanisms of naming (H5-H6). The Eurolinguistic analyses follow (H7-H12), comparing Europe with non-European regions, language families, and sub-regions, and contrasting European with North American ideological effects and naming conventions. Throughout, results are reported without interpretation beyond what the data warrant; broader implications are reserved for the discussion.

### 6.1. H1: Distribution of Bias

H1: Headlines are more often evaluative than neutral. Specifically, the combined share of positive and negative headlines outweighs the neutral category. This reflects Trump's own framing of the media as adversarial: if he perceived relentless negativity, then a baseline assumption is that evaluative stance (whether hostile or laudatory) dominates headline reporting.

H0: The evaluative share is not greater than the neutral share.

Out of 10,939 headlines, 6,018 (55%) were coded as neutral and 4,921 (45%) as evaluative (Table 10, Figure 3). A chi-square goodness-of-fit test confirmed that this distribution deviated significantly from an even split,  $\chi^2 (N = 10,939) = 110.01, p < .001$ . Neutral headlines were overrepresented (+548.5), while evaluative headlines were underrepresented (-548.5) relative to expected values.

These results run counter to H1, which predicted that evaluative stances would dominate headline reporting. Instead, neutral headlines prevailed across the corpus, suggesting that international newspapers tended to frame U.S. presidential candidates in a more factual or non-evaluative manner. Accordingly, H0 is supported, as evaluative stances were less frequent than hypothesized.

|                   | Frequency     |
|-------------------|---------------|
| <b>Neutral</b>    | 6,018         |
| <b>Evaluative</b> | 4,921         |
| <b>Total</b>      | <b>10,939</b> |

Table 10 Overall stance - counts

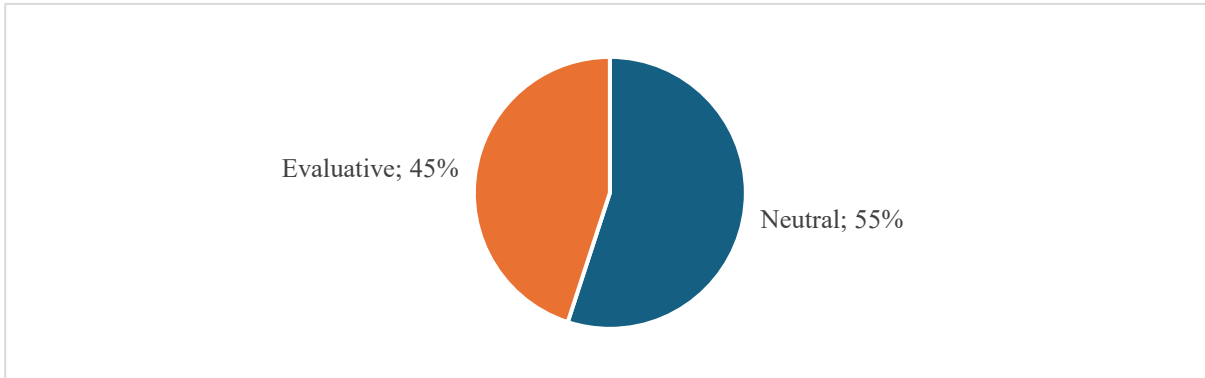


Figure 3 Overall stance – percent

## 6.2. H2: Candidate-Specific Bias

H2: Across the international corpus, Hillary Clinton, Joe Biden, and Kamala Harris receive a more positive balance of coverage than Donald J. Trump, who is disproportionately framed in negative terms. This hypothesis is directly aligned with Trump’s claim that mainstream journalism is structurally biased against him.

H0: The distribution of headline bias is independent of candidate.

Bias distribution varied substantially across candidates (Table 11, Figure 4). Donald Trump received 55.2% neutral, 6.9% positive, and 37.9% negative headlines (N = 8,160). By contrast, Hillary Clinton’s headlines were 43.9% neutral, 33.2% positive, and 22.9% negative (N = 707). Joe Biden’s distribution was 58.5% neutral, 23.1% positive, and 18.4% negative (N = 1,531). Kamala Harris showed a similar pattern with 57.1% neutral, 25.5% positive, and 17.4% negative headlines (N = 541). A chi-square test of independence confirmed that these differences were highly significant,  $\chi^2(6, N = 10,939) = 929.23$ ,  $p < .001$ .

|                | Neutral      | Positive     | Negative     | Total         |
|----------------|--------------|--------------|--------------|---------------|
| <b>Trump</b>   | 4,503        | 566          | 3,091        | <b>8,160</b>  |
| <b>Clinton</b> | 310          | 235          | 162          | <b>707</b>    |
| <b>Biden</b>   | 896          | 354          | 281          | <b>1,531</b>  |
| <b>Harris</b>  | 309          | 138          | 94           | <b>541</b>    |
| <b>Total</b>   | <b>6,018</b> | <b>1,293</b> | <b>3,628</b> | <b>10,939</b> |

Table 11 Bias by candidate -counts

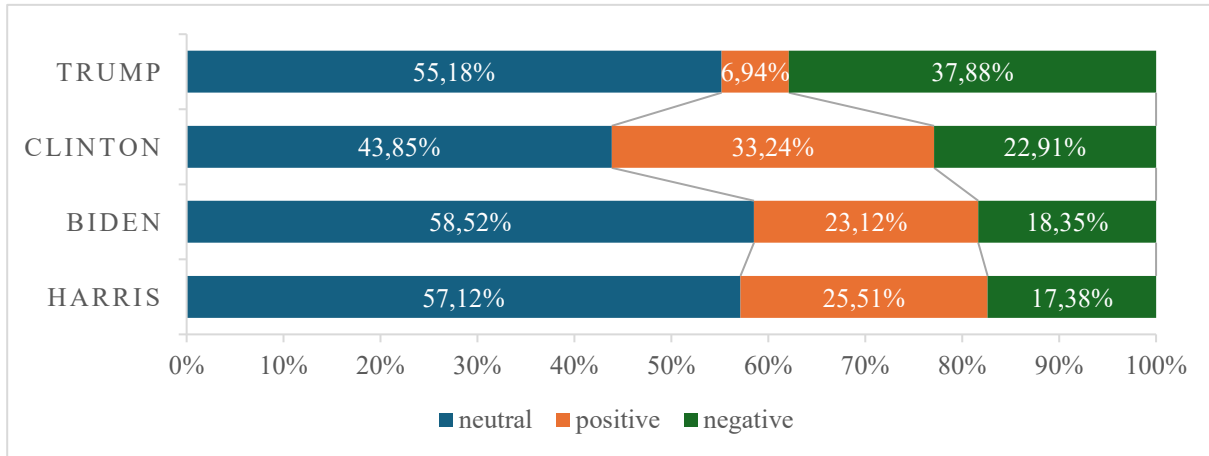


Figure 4 Bias by candidate – percent

The results support H2. Trump was disproportionately framed in negative terms compared to his Democratic counterparts, who consistently received a higher share of positive coverage and a lower share of negative coverage. The data therefore confirm that headline bias is not independent of candidate identity. Instead, Trump was subject to systematically more negative portrayals, while Clinton, Biden, and Harris benefitted from a more positive balance of coverage.

### 6.3. H3: Bias and Ideological Stance of Newspapers

H3: The orientation of bias in headlines correlates with the ideological leaning of outlets. Right-leaning newspapers are more likely to frame Trump positively and Democratic candidates negatively, whereas left-leaning outlets display the reverse tendency. As Higdon, Marmol, and Huff (2022: 265) note, legacy media tend to reproduce their audiences' ideological expectations, presenting partisan actors either as "heroes or archvillains."

H0: The distribution of headline bias is independent of outlet ideology, for each candidate.

Headline bias varied systematically with outlet ideology for both Trump and Democratic candidates (Table 12, Figure 5). For Trump, left-leaning outlets produced only 5.4% positive headlines, compared with 39.6% negative ones ( $N = 4,850$ ). In contrast, right-leaning outlets framed Trump somewhat more favorably, with 9.2% positive and 36.6% negative ( $N = 2,374$ ). Liberal and middle outlets fell between these poles but still leaned negative, with 7.9% positive versus 31.3% negative in liberal outlets ( $N = 555$ ) and 10.5% positive versus 32.8% negative in middle outlets ( $N = 381$ ). Democratic candidates showed the opposite pattern. Left-leaning outlets presented them most positively (27.6% positive vs. 17.5% negative,  $N = 1,616$ ), while right-leaning outlets were less favorable (26.5% positive vs. 21.6% negative,  $N = 827$ ). Liberal and middle outlets again occupied intermediate positions, with 21.6% positive coverage versus 17.0% negative for middle outlets ( $N = 153$ ) and liberal outlets diverging from this pattern by assigning Democrats more negative (26.8%) than positive (15.9%) coverage ( $N = 183$ ). All chi-square tests confirmed significant associations between outlet ideology and bias distributions, including for the full dataset,  $\chi^2(2, N = 10,939) = 863.55, p < .001$ .

|                             | Neutral      | Positive     | Negative     | Total         |
|-----------------------------|--------------|--------------|--------------|---------------|
| <b>Democratic candidate</b> | 1,515        | 727          | 537          | <b>2,779</b>  |
| left                        | 887          | 446          | 283          | <b>1,616</b>  |
| liberal                     | 105          | 29           | 49           | <b>183</b>    |
| middle                      | 94           | 33           | 26           | <b>153</b>    |
| right                       | 429          | 219          | 179          | <b>827</b>    |
| <b>Trump</b>                | 4,503        | 566          | 3,091        | <b>8,160</b>  |
| left                        | 2,663        | 264          | 1,923        | <b>4,850</b>  |
| liberal                     | 337          | 44           | 174          | <b>555</b>    |
| middle                      | 216          | 40           | 125          | <b>381</b>    |
| right                       | 1,287        | 218          | 869          | <b>2,374</b>  |
| <b>Total</b>                | <b>6,018</b> | <b>1,293</b> | <b>3,628</b> | <b>10,939</b> |

Table 12 Bias by outlet ideology and candidate - counts

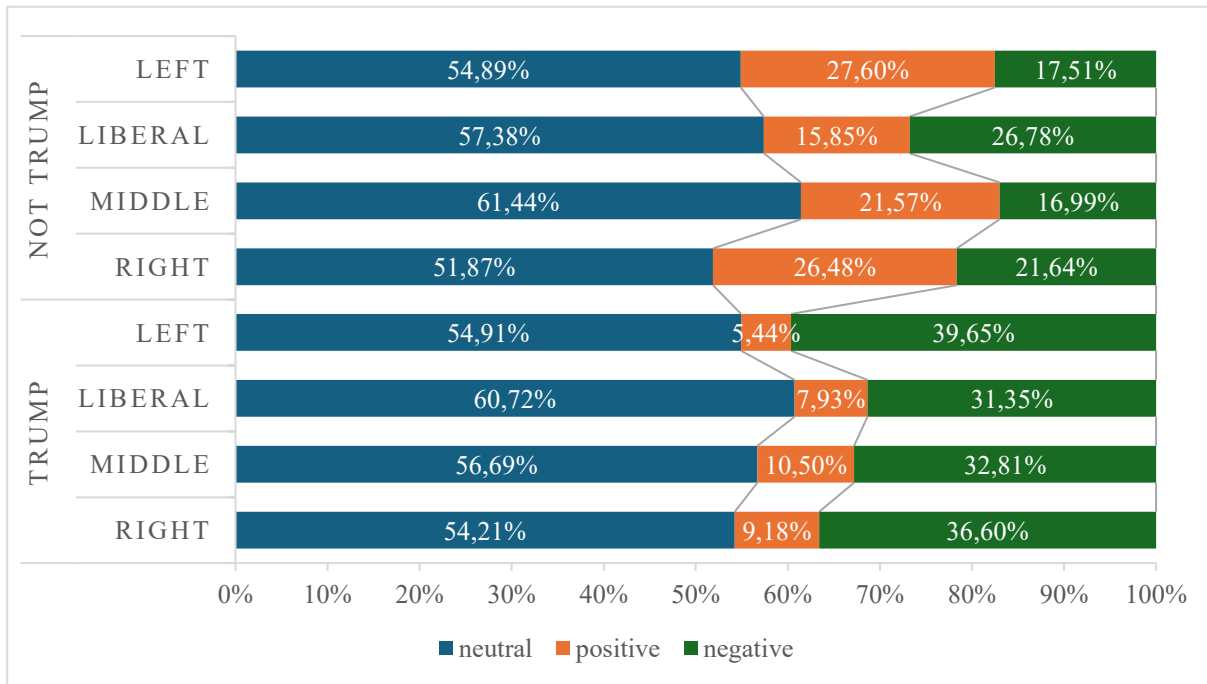


Figure 5 Bias by outlet ideology and candidate – percent

These results provide partial support for H3. Bias was significantly associated with outlet ideology, and the directional tendencies were consistent with the hypothesis: right-leaning outlets were relatively more favorable toward Trump, while left-leaning outlets were relatively more favorable toward Democratic candidates. However, these effects were relative rather than absolute. All ideological groups were still net negative toward Trump, and neutral reporting remained the dominant category across candidates. Moreover, liberal outlets did not align neatly with the expected pattern, as they framed Democratic candidates more negatively than positively. Overall, the findings indicate that ideology shaped headline bias, but within the broader tendency toward neutrality or negativity across the corpus.

#### 6.4. H4: Stability Across Electoral Cycles

H4: The distribution of bias remains consistent across the 2016, 2020, and 2024 campaigns. If Trump’s characterization of legacy media as persistently “bad” were accurate, negative bias would not be episodic but sustained across cycles.

H0: The distribution of headline bias differs by year.

Headline bias distribution showed noticeable shifts across the three election cycles (Table 13, Figure 6). In 2016, negative headlines were relatively frequent (37.6%), with neutral coverage at 51.9% and positive coverage at 10.5% (N = 3,631). By 2020, the proportion of negative coverage had declined to 32.2%, while neutral reporting rose modestly to 55.1% and positive to 12.7% (N = 3,286). In 2024, this trend toward neutrality continued, with 57.8% of headlines classified as neutral, 29.9% as negative, and 12.3% as positive (N = 4,022). A chi-square test of independence confirmed that these changes were statistically significant,  $\chi^2(4, N = 10,939) = 55.07, p < .001$ .

|              | Neutral      | Positive     | Negative     | Total         |
|--------------|--------------|--------------|--------------|---------------|
| <b>2016</b>  | 1,884        | 382          | 1,365        | <b>3,631</b>  |
| <b>2020</b>  | 1,810        | 417          | 1,059        | <b>3,286</b>  |
| <b>2024</b>  | 2,324        | 494          | 1,204        | <b>4,022</b>  |
| <b>Total</b> | <b>6,018</b> | <b>1,293</b> | <b>3,628</b> | <b>10,939</b> |

Table 13 Bias by campaign year - counts

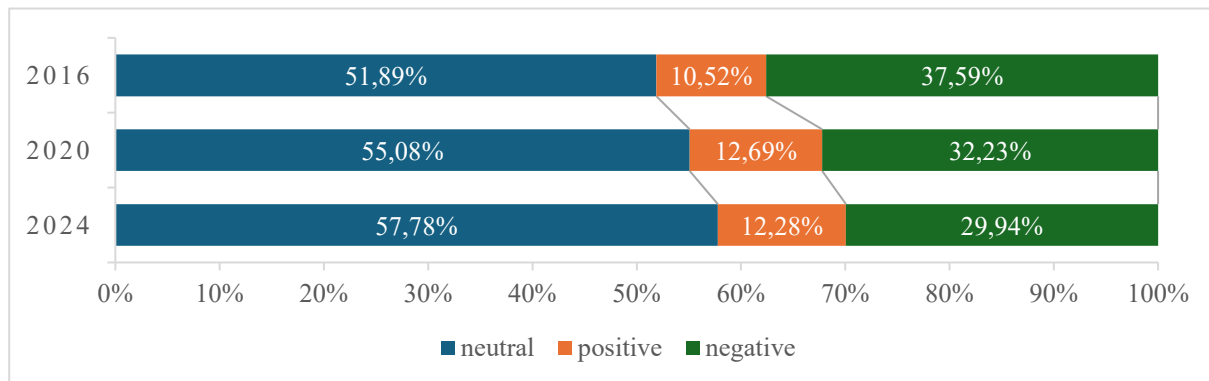


Figure 6 Bias by campaign year - percent

These findings do not support H4. Rather than remaining stable, the distribution of headline bias shifted across cycles, with negative coverage gradually decreasing and neutral coverage increasing. While Trump’s criticism of the media presupposes sustained negativity, the data show that headline framing became less negative over time. The results therefore align with H0, indicating that the distribution of bias is contingent on campaign cycle rather than constant across them.

#### 6.5. H5: Naming Practices

H5: Legacy media adhere primarily to neutral naming conventions (surname only or first name + surname) rather than overtly evaluative forms (nicknames, epithets, or derogatory designations). This builds

on van Leeuwen's (2008) typology of naming strategies and Würth's (2016) findings that newspapers tend to rely on conventional nomination practices.

H0: Neutral naming is not more frequent than evaluative naming.

Naming conventions in headlines were overwhelmingly neutral (Table 14, Figure 7). Across the corpus, 10,735 references (98.1%) adhered to conventional forms such as surname only or given name plus surname, while only 204 cases (1.9%) used marked forms such as nicknames, epithets, or derogatory designations. A chi-square goodness-of-fit test confirmed that this distribution was highly significant,  $\chi^2(1, N = 10,939) = 10,138.22, p < .001$ .

|                | Frequency     |
|----------------|---------------|
| <b>Neutral</b> | 10,735        |
| <b>Marked</b>  | 204           |
| <b>Total</b>   | <b>10,939</b> |

Table 14 Naming practices - counts

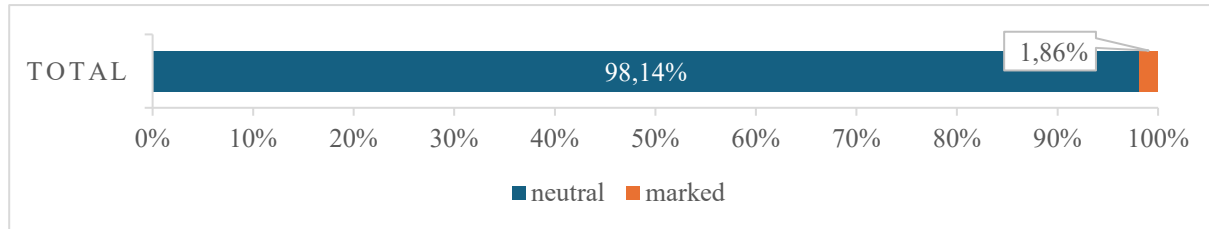


Figure 7 Naming practices - percent

These findings support H5. Legacy media overwhelmingly relied on neutral naming conventions rather than evaluative or marked forms. The near absence of nicknames or epithets suggests strong adherence to conventional journalistic practice in headline formulation, consistent with earlier findings on nomination strategies in international press coverage.

## 6.6. H6: Value-Judgement Labels

H6: When evaluative labels are attached to candidates' names, they occur more frequently with Donald Trump than with his Democratic counterparts. This hypothesis echoes Kuypers' (2014: 5) observation that value-laden descriptors are disproportionately applied to Republicans.

H0: The frequency of evaluative labels is independent of candidate; Trump does not receive such labels more often than Democrats.

The distribution of value-judgment labels differed significantly by candidate (Table 15, Figure 8). For Democratic candidates, 96.0% of references contained no evaluative label, 1.6% were laudative, and 2.5% pejorative ( $N = 2,779$ ). For Trump, evaluative naming was more common: 93.7% none, 1.1% laudative, and 5.3% pejorative ( $N = 8,160$ ). Thus, Trump's share of pejorative labels was more than double that of Democratic candidates. A chi-square test of independence confirmed this difference was statistically significant,  $\chi^2(2, N = 10,939) = 39.56, p < .001$ .

|                  | None          | Laudative  | Pejorative | Total         |
|------------------|---------------|------------|------------|---------------|
| <b>Not Trump</b> | 2,667         | 43         | 69         | <b>2,779</b>  |
| <b>Trump</b>     | 7,643         | 89         | 428        | <b>8,160</b>  |
| <b>Total</b>     | <b>10,310</b> | <b>132</b> | <b>497</b> | <b>10,939</b> |

Table 15 Value-judgement labels by candidate - counts

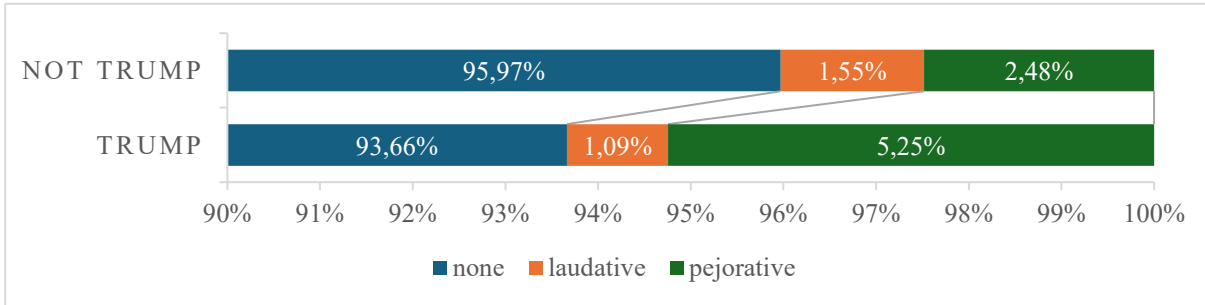


Figure 8 Value-judgement labels by candidate - percent

These results support H6. Evaluative labels occurred disproportionately in reference to Trump, primarily in the form of pejorative characterizations. Democratic candidates were not immune to evaluative naming, but such instances were comparatively rare and more evenly distributed between laudative and pejorative forms. The findings therefore align with Kuypers' (2014) observation that Republicans are more frequently subjected to value-laden descriptors in mainstream news coverage.

### 6.7. H7: European Unity of Bias

H7: European newspapers display a shared discursive pattern of negativity toward Donald Trump, suggesting a common transnational orientation. This draws on Eurolinguistic research (Grzega 2025; Giesen & Szwed 2025) which emphasizes the identification of pan-European discursive features.

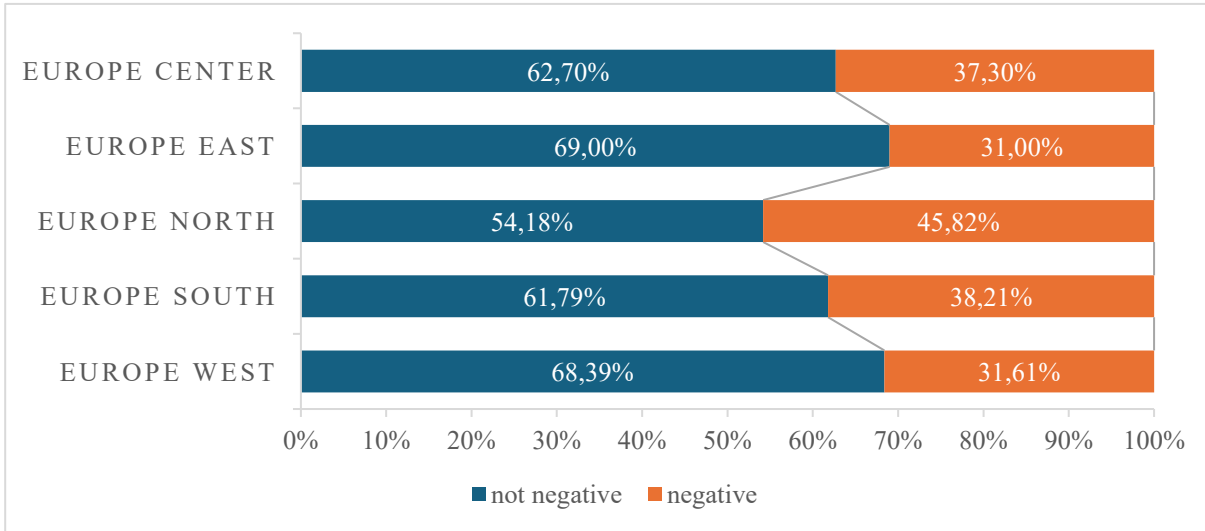
H0: Among European outlets, there is no pan-European negative orientation toward Trump; negative share is not higher than non-negative, or negativity varies widely across countries.

The distribution of bias toward Trump in European outlets shows both commonalities and internal variation (Table 16, Figure 9). Across all European regions combined, 39.9% of headlines were negative and 60.1% were not negative (N = 4,880). Negativity was present in every region, but its intensity varied: it was strongest in Northern Europe (45.8% negative) and weakest in Eastern and Western Europe (31.0% and 31.6% respectively). A chi-square test of independence confirmed that these regional differences were statistically significant,  $\chi^2$  (4, N = 4,880) = 76.28,  $p < .001$ , though the effect size was weak (Cramer's V = .125).

|                      | Not negative | Negative     | Total        |
|----------------------|--------------|--------------|--------------|
| <b>Europe Center</b> | 400          | 238          | <b>638</b>   |
| <b>Europe East</b>   | 492          | 221          | <b>713</b>   |
| <b>Europe North</b>  | 1,264        | 1,069        | <b>2,333</b> |
| <b>Europe South</b>  | 393          | 243          | <b>636</b>   |
| <b>Europe West</b>   | 383          | 177          | <b>560</b>   |
| <b>Grand Total</b>   | <b>2,932</b> | <b>1,948</b> | <b>4,880</b> |

*Table 16 Bias in Trump headlines by European sub-region - counts*

These findings provide partial support for H7. While negativity toward Trump was evident across all European sub-regions, suggesting a broadly shared discursive orientation, it did not dominate coverage overall, and significant regional variation undermines the idea of a fully uniform “pan-European” stance. Instead, the results indicate a shared critical tendency with differences in intensity, particularly between Northern Europe and other regions.



*Figure 9 Bias in Trump headlines by European sub-region - percent*

## 6.8. H8: Europe Versus the Rest of the World

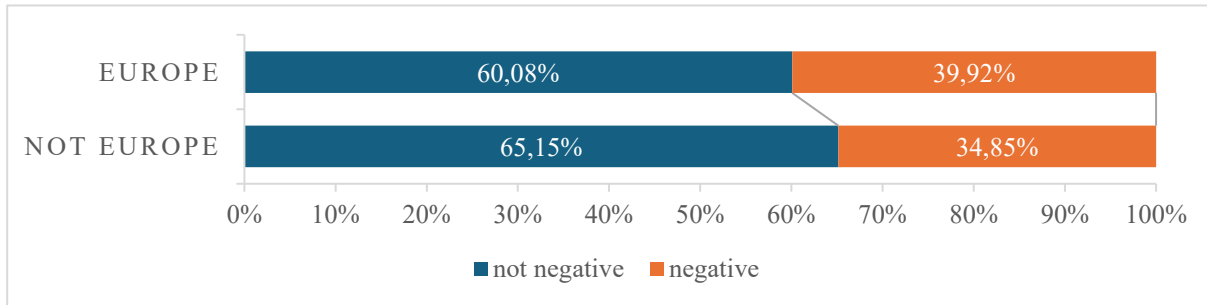
H8: European outlets frame Trump more negatively than newspapers from non-European regions, reflecting structural and cultural differences in political communication.

H0: Trump-related bias does not differ between European and non-European outlets.

A comparison of European and non-European outlets shows that European coverage of Trump was more negative (Table 17, Figure 10). In Europe, 39.9% of headlines were negative and 60.1% not negative (N = 4,880), while in non-European outlets the proportions were 34.9% negative and 65.1% not negative (N = 3,280). The difference was statistically significant,  $\chi^2 (1, N = 8,160) = 21.43, p < .001$ , although the association was weaker (Cramer's V = .051). Risk estimates indicated that Trump was about 15% more likely to receive negative framing in Europe than elsewhere (OR = 1.15, 95% CI [1.08, 1.21]).

|                   | Not negative | Negative     | Total        |
|-------------------|--------------|--------------|--------------|
| <b>Europe</b>     | 2,932        | 1,948        | <b>4,880</b> |
| <b>Not Europe</b> | 2,137        | 1,143        | <b>3,280</b> |
| <b>Total</b>      | <b>5,069</b> | <b>3,091</b> | <b>8,160</b> |

Table 17 Stance towards Donald Trump by region - counts



These findings support H8. European newspapers framed Trump more negatively than non-European outlets, suggesting a distinctive regional orientation in political communication. However, the weak effect size indicates that this divergence, while consistent, was relatively modest in scale.

## 6.9. H9: Cross-linguistic Convergence

H9: Across Europe's major linguistic families (Germanic, Romance, Slavic, etc.), newspapers exhibit comparable patterns of bias in reporting on U.S. presidential candidates. This reflects a shared European

Figure 10 Stance towards Donald Trump by region - percent

discursive orientation rather than language-family-specific traditions.

H0: The distribution of bias differs across language families.

Bias distributions across Europe's three major language families followed broadly comparable patterns (Table 18, Figure 11). Germanic outlets produced 37.1% negative headlines, Romance 32.3%, and Slavic 28.8%, with neutral coverage remaining the majority in all three groups. Positive headlines were rare across the board (11.1-12.0%). A chi-square test of independence detected statistically significant variation across language families,  $\chi^2 (4, N = 6,630) = 34.36, p < .001$ , but the effect size was weak (Cramer's  $V = .051$ ).

|                 | Neutral      | Positive   | Negative     | Total        |
|-----------------|--------------|------------|--------------|--------------|
| <b>Germanic</b> | 2,269        | 536        | 1,653        | <b>4,458</b> |
| <b>Romance</b>  | 685          | 139        | 393          | <b>1,217</b> |
| <b>Slavic</b>   | 574          | 106        | 275          | <b>955</b>   |
| <b>Total</b>    | <b>3,528</b> | <b>781</b> | <b>2,321</b> | <b>6,630</b> |

Table 18 Bias by European language family - counts

These results offer partial support for H9. While minor differences were observed – Germanic outlets were somewhat more negative, Slavic somewhat less – the overall distributional pattern was consistent across families: negative coverage outweighed positive, with neutrality dominant. The findings suggest a shared European discursive orientation, with only limited divergence linked to language structure types.

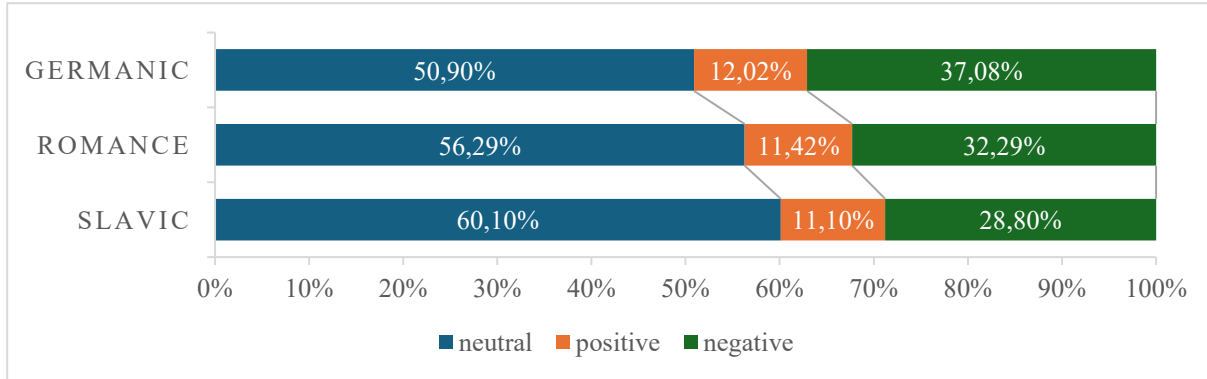


Figure 11 Bias by European language family - percent

### 6.10. H10: Regional Consistency

H10: Despite cultural and geopolitical diversity, European newspapers converge on a broadly similar evaluative stance toward Trump and his opponents. Bias does not systematically differ across Northern, Southern, Western, Eastern, or Central Europe, indicating pan-European alignment in headline discourse.

H0: The distribution of bias differs across European sub-regions.

Headline bias distribution showed both convergence and divergence across European sub-regions (Table 19, Figure 12). For Democratic candidates, coverage was broadly stable: Biden and Clinton displayed little systematic variation, with chi-square tests of independence confirming no significant differences across sub-regions,  $\chi^2 (8, N = 977) = 7.70, p = .463$ ;  $\chi^2 (8, N = 414) = 7.87, p = .447$ . Harris showed a marginally significant result,  $\chi^2 (8, N = 359) = 15.75, p = .046$ , though the effect size was weak (Cramer's  $V = .148$ ). By contrast, Trump's coverage varied more substantially: in Northern Europe, 45.8% of headlines were negative compared with 31.0% in Eastern Europe and 31.6% in Western Europe. This variation was statistically significant,  $\chi^2 (8, N = 4,880) = 83.27, p < .001$ , although the effect size was small (Cramer's  $V = .092$ ).

Taken together, these results offer only partial support for H10. While Democratic candidates were covered in a broadly consistent manner across Europe, Trump's portrayal diverged between sub-regions, with particularly negative framing in Northern Europe. Nevertheless, the strength of these associations was weak (Cramer's  $V \leq .15$ ), and the dominant evaluative pattern remained similar across the continent: neutral headlines prevailed, positive coverage was limited, and negative reporting outweighed positive. The results therefore indicate a mixed picture: European coverage does exhibit signs of pan-national convergence, especially for Democrats, but regional divergence remains evident in Trump's case, partially aligning with H0.

|                      | <b>Neutral</b> | <b>Positive</b> | <b>Negative</b> | <b>Total</b> |
|----------------------|----------------|-----------------|-----------------|--------------|
| <b>Biden</b>         | 571            | 215             | 191             | <b>977</b>   |
| <b>Europe Center</b> | 42             | 21              | 18              | <b>81</b>    |
| <b>Europe East</b>   | 114            | 28              | 35              | <b>177</b>   |
| <b>Europe North</b>  | 283            | 114             | 97              | <b>494</b>   |
| <b>Europe South</b>  | 62             | 28              | 23              | <b>113</b>   |
| <b>Europe West</b>   | 70             | 24              | 18              | <b>112</b>   |
| <b>Clinton</b>       | 172            | 135             | 107             | <b>414</b>   |
| <b>Europe Center</b> | 25             | 22              | 19              | <b>66</b>    |
| <b>Europe East</b>   | 1              |                 |                 | <b>1</b>     |
| <b>Europe North</b>  | 89             | 82              | 64              | <b>235</b>   |
| <b>Europe South</b>  | 34             | 17              | 12              | <b>63</b>    |
| <b>Europe West</b>   | 23             | 14              | 12              | <b>49</b>    |
| <b>Harris</b>        | 191            | 93              | 75              | <b>359</b>   |
| <b>Europe Center</b> | 9              | 12              | 5               | <b>26</b>    |
| <b>Europe East</b>   | 27             | 18              | 19              | <b>64</b>    |
| <b>Europe North</b>  | 102            | 42              | 37              | <b>181</b>   |
| <b>Europe South</b>  | 20             | 12              | 8               | <b>40</b>    |
| <b>Europe West</b>   | 33             | 9               | 6               | <b>48</b>    |
| <b>Trump</b>         | 2,594          | 338             | 1,948           | <b>4,880</b> |
| <b>Europe Center</b> | 353            | 47              | 238             | <b>638</b>   |
| <b>Europe East</b>   | 432            | 60              | 221             | <b>713</b>   |
| <b>Europe North</b>  | 1,137          | 127             | 1,069           | <b>2,333</b> |
| <b>Europe South</b>  | 345            | 48              | 243             | <b>636</b>   |
| <b>Europe West</b>   | 327            | 56              | 177             | <b>560</b>   |
| <b>Total</b>         | <b>3,528</b>   | <b>781</b>      | <b>2,321</b>    | <b>6,630</b> |

*Table 19 Bias by European sub-region and candidate - counts*

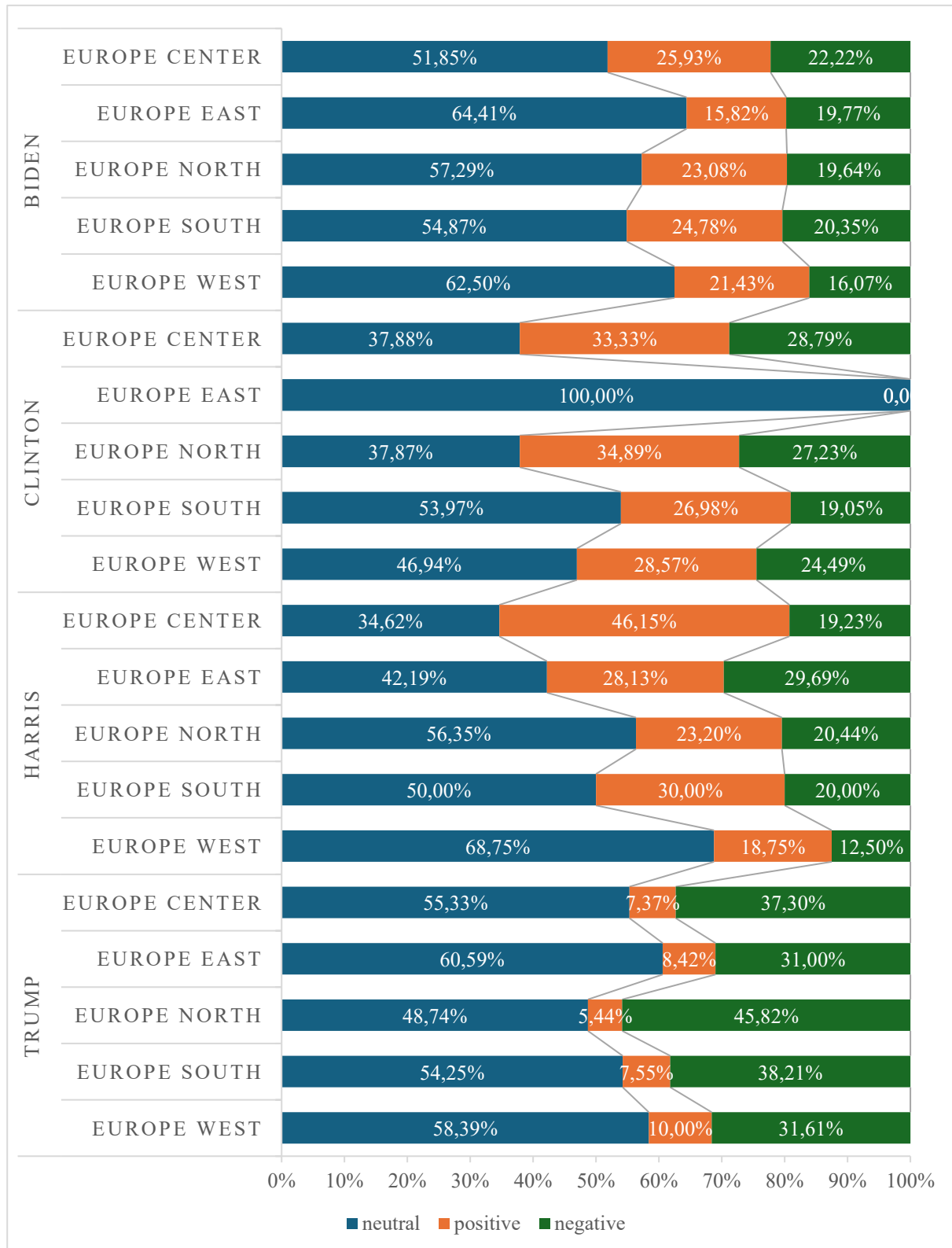


Figure 12 Bias by European sub-region and candidate - percent

### 6.11. H11. European Ideological Distinctiveness

H11: Within Europe, the ideological stance of newspapers (left-leaning vs. right-leaning) has less impact on their portrayal of Donald Trump than in the North American media context. This suggests a shared European distance from Trump that overrides traditional partisan divides.

H0: In Europe, the association between ideology and bias toward Trump is at least as strong as (or stronger than) in North America; there is no attenuation in Europe.

Headline bias toward Donald Trump varied significantly by outlet ideology in both Europe and North America (Table 20, Figure 13). In Europe, left-leaning outlets produced 44.7% negative coverage, 50.4% neutral, and 4.9% positive (N = 2,539). Right-leaning European outlets were less negative, with 35.8% negative, 54.7% neutral, and 9.5% positive (N = 1,786). In North America, the ideological gap was narrower: left-leaning outlets recorded 41.6% negative, 62.0% neutral, and 5.1% positive (N = 1,275), while right-leaning outlets showed 32.9% negative, 54.3% neutral, and 4.1% positive (N = 315). A chi-square tests of independence confirmed significant associations in both regions (Europe:  $\chi^2$  (2, N = 4,325) = 55.40,  $p < .001$ ; North America:  $\chi^2$  (2, N = 1,590) = 8.40,  $p < .015$ ).

|                      | Neutral      | Positive   | Negative     | Total        |
|----------------------|--------------|------------|--------------|--------------|
| <b>Europe</b>        |              |            |              |              |
| left                 | 2,257        | 294        | 1,774        | <b>4,325</b> |
| right                | 1,280        | 125        | 1,134        | <b>2,539</b> |
| <b>North America</b> |              |            |              |              |
| left                 | 977          | 169        | 640          | <b>1,786</b> |
| right                | 961          | 78         | 551          | <b>1,590</b> |
| left                 | 790          | 65         | 420          | <b>1,275</b> |
| right                | 171          | 13         | 131          | <b>315</b>   |
| <b>Total</b>         | <b>3,218</b> | <b>372</b> | <b>2,325</b> | <b>5,915</b> |

Table 20 Bias towards Donald Trump by outlet ideology in Europe vs North America – counts

Comparisons of effect size further underscore the difference. In Europe, ideology was more predictive of Trump-related bias (Cramer's  $V = .113$ ) than in North America (Cramer's  $V = .073$ ). This indicates that partisan divides in Trump coverage were not attenuated in the European press, as hypothesized, but instead more pronounced in North America. These findings therefore do not support H11. Rather than showing a shared European distance that overrides partisan divides, the results align with H0, demonstrating that ideological stance had a stronger impact on Trump coverage in Europe than across the Atlantic.

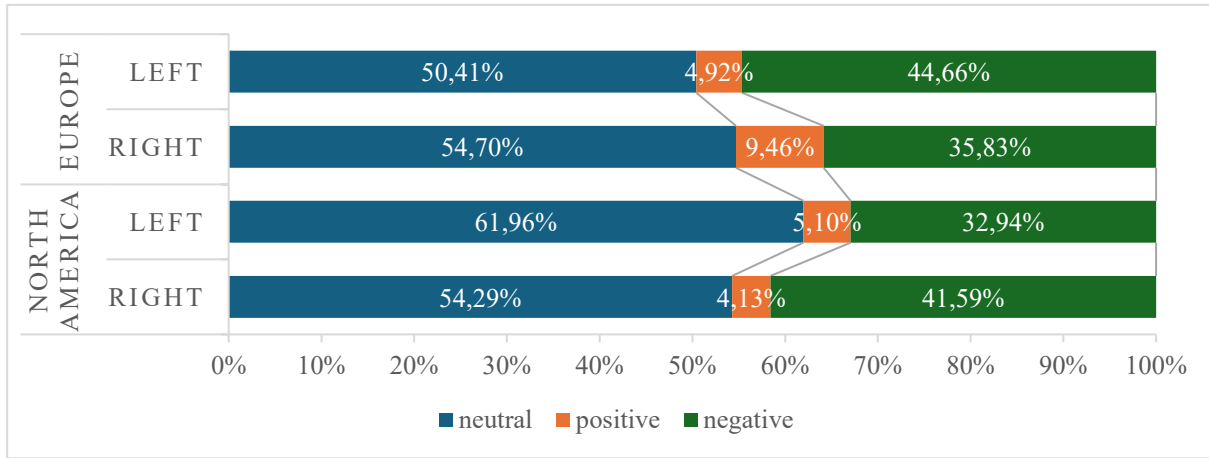


Figure 13 Bias towards Donald Trump by outlet ideology in Europe vs North America - percent

## 6.12. H12. Unified Naming Conventions

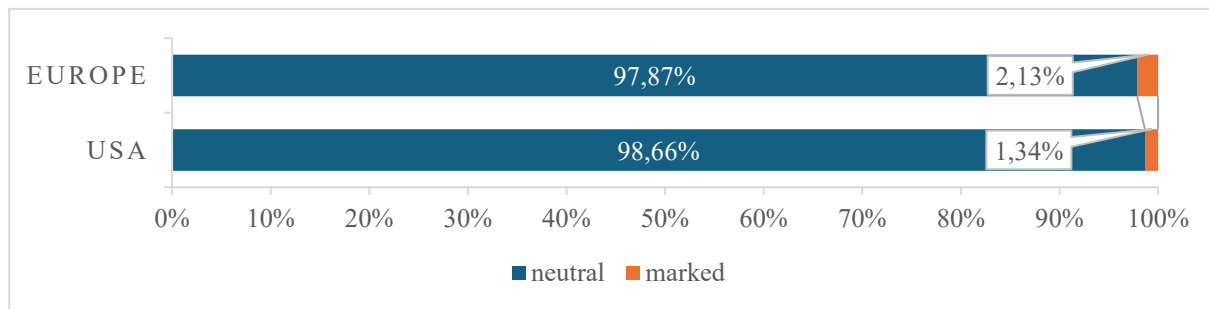
H12: European newspapers converge on neutral naming practices (surname, or first + surname), with minimal use of evaluative nicknames or derogatory labels. This uniformity reflects a European journalistic convention distinct from U.S. patterns.

H0: Within Europe, neutral naming is not more frequent than evaluative naming.

Naming practices were overwhelmingly neutral in both Europe and the USA (Table 21, Figure 14). In Europe, 97.9% of references were neutral, while only 2.1% were marked (N = 6,630). In the USA, neutral naming was even more dominant at 98.7%, with just 1.3% marked forms (N = 1,495). A chi-square test of independence indicated a statistically significant but marginal difference between the regions,  $\chi^2 (1, N = 8,125) = 3.91, p = .048$ . The odds ratio (.624, 95% CI [.389, 1.000]) suggests that European outlets were somewhat more likely than U.S. outlets to employ named marking, though the effect size was very small.

|        | Neutral | Marked | Total |
|--------|---------|--------|-------|
| Europe | 6,489   | 141    | 6,630 |
| USA    | 1,475   | 20     | 1,495 |
| Total  | 7,964   | 161    | 8,125 |

Table 21 Naming practices in Europe vs USA - counts



*Figure 14 Naming practices in Europe vs USA - percent*

These results provide only partial support for H12. While the data confirm that European outlets converge strongly on neutral naming practices, the same pattern is observed in U.S. coverage. The small, statistically significant difference indicates slightly greater tolerance for marked naming in Europe, but this effect is negligible in practical terms. Overall, the findings suggest that neutral naming is a transatlantic journalistic convention rather than a uniquely European one, and thus the evidence aligns more closely with H0 than with H12.

### 6.13. Summary of Findings

Across 10,939 headlines, neutrality predominated, and evaluative headlines did not outweigh neutral ones (H1 not supported). Candidate-specific patterns were pronounced: Trump received substantially more negative and less positive coverage than Clinton, Biden, or Harris (H2 supported). Outlet ideology correlated with bias in the expected directions – right-leaning outlets were relatively more favorable to Trump and less favorable to Democrats, with the reverse for left-leaning outlets – yet all ideological groups remained net negative toward Trump and largely neutral overall (H3 partially supported). Bias distributions shifted across cycles, with negativity declining and neutrality increasing from 2016 to 2024 (H4 not supported). Naming practices were overwhelmingly neutral (surname or fist + surname), with marked/evaluative forms rare (H5 supported), but when value-labels did appear they attached disproportionately to Trump and were typically pejorative (H6 supported). Within Europe, negativity toward Trump was evident in all sub-regions but varied in intensity – strongest in the North – indicating shared tendencies with regional modulation (H7 partially supported). Europe framed Trump more negatively than non-European outlets, although the difference was modest (H8 supported). Language-family patterns were broadly convergent, with small differences (H9 partially supported). Sub-regional consistency held more clearly for Democrats than for Trump (H10 partially supported). Contrary to expectations, ideological effects on Trump coverage were at least as strong – indeed stronger – in Europe than in North America (H11 not supported). Finally, neutral naming was the norm on both sides of the Atlantic, with only a negligible difference in marked forms (H12 not supported/only weakly supported). Overall, the results reveal a nuanced landscape: no monolithic anti-Trump stance, but persistent asymmetries by candidate and ideology, generally small effect sizes, and strong cross-system adherence to neutral naming.

## 7. Discussion

### 7.1. Introduction: Framing the Discussion

The purpose of this study was to investigate how major newspapers across Europe and beyond have represented Donald Trump and his Democratic rivals during the last three U.S. presidential election cycles (2016, 2020, 2024). Through a systematic discourse-analytic approach, supported by large-scale AI-assisted coding, the study examined more than 10,900 headlines to explore patterns of bias, naming

practices, evaluative strategies, and cross-national tendencies. At the heart of the analysis lay a series of questions central to the field of media discourse research: To what extent does political bias permeate headline reporting? How are candidates differentially portrayed across ideological, geographical, and linguistic contexts? And what do these patterns reveal about broader tendencies in European and global media discourse? The discussion that follows aims to situate these findings within the wider scholarly literature and to unpack their theoretical implications. It critically engages with longstanding debates on media bias – from D’Alessio and Allen’s (2000) conclusion that presidential campaign coverage shows little net bias to Kuypers’ (2014, 2020) contention that the press is systematically liberal – and examines how the data from this thesis confirm, complicate, or challenge those narratives. It also explores the significance of naming and labeling strategies, the structural dynamics of media coverage, and the elusive notion of a “shared European discourse.” Finally, the discussion reflects on the methodological contributions and limitations of AI-assisted discourse analysis and considers what this means for future research.

A key principle underpinning the discussion is Bendel Larcher’s (2023) reminder that discourse analysts can only make claims about the specific discourses they have examined, not about all possible manifestations of those discourses. The present findings therefore speak only to the headlines, outlets, and timeframes included in the corpus. They cannot, and do not seek to, generalize to the entirety of media discourse surrounding Trump or other candidates. Yet within those limits, the results offer robust evidence that enriches – and in some cases complicates – the existing understanding of political media bias.

## 7.2. Media Bias and the Myth of Partisan Imbalance

### 7.2.1. Revisiting D’Alessio & Allen: Evidence of Bias or Not?

D’Alessio and Allen’s (2000) meta-analysis of U.S. presidential campaign coverage between 1948 and 1996 concluded that, over the long term and across media types, there was no significant net bias. They found that isolated instances of slant were offset by countervailing examples, and that newspapers in particular showed negligible overall partisan imbalance. On the surface, some of the present findings resonate with this assessment: neutrality indeed dominates the coverage of all candidates. Across the corpus, approximately 54.5% of headlines referring to Democratic candidates and 55.2% of those referring to Trump were coded as neutral (see Table 22 and Figure 15).

This predominance of neutrality complicates common assumptions about media partisanship. If most headlines are neither explicitly positive nor negative, then media discourse may be less ideologically driven than public debate often assumes. However, neutrality alone does not mean that bias is absent. When we examine the distribution of biased reporting, clear asymmetries emerge. Coverage of Democratic candidates skews modestly positive – with 26.2% of headlines coded as positive and 19.3% as negative – whereas Trump receives overwhelmingly negative treatment, with 37.9% negative coverage and only 6.9% positive.

|                  | Neutral      | Positive     | Negative     | Total         |
|------------------|--------------|--------------|--------------|---------------|
| <b>Not Trump</b> | 1,515        | 727          | 537          | <b>2,779</b>  |
| <b>Trump</b>     | 4,503        | 566          | 3,091        | <b>8,160</b>  |
| <b>Total</b>     | <b>6,018</b> | <b>1,293</b> | <b>3,628</b> | <b>10,939</b> |

*Table 22 Stance towards Democratic candidates and Trump – counts*

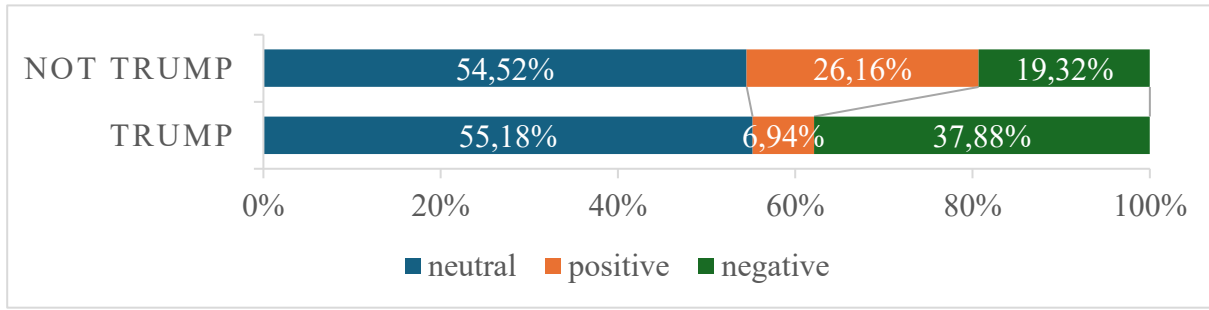


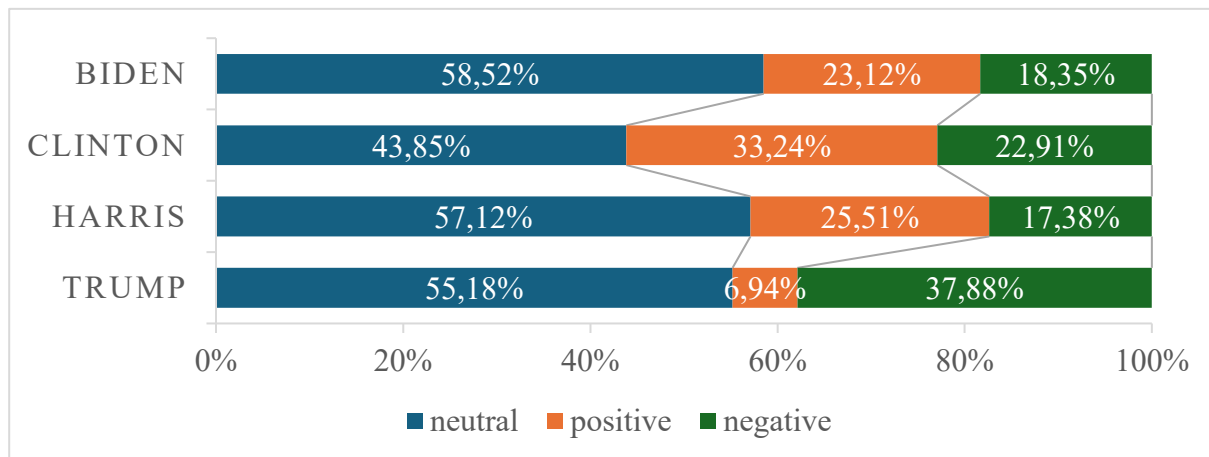
Figure 15 Stance towards Democratic candidates and Trump - percent

These disparities suggest that, while neutrality remains the most common evaluative stance, the balance of bias is far from even. The difference in positive coverage between Trump and the Democrats is 19.2 percentage points, and the gap in negative coverage is 18.6 points. This complicates D'Alessio and Allen's conclusions: even if neutrality predominates, the non-neutral portion of reporting is heavily asymmetrical. This nuance underscores Bendel Larcher's point about the scope of discourse analysis: while D'Alessio and Allen may have found negligible bias in the corpus they studied, the present findings reveal a strong imbalance within the headlines examined here.

|                | Neutral      | Positive     | Negative     | Total         |
|----------------|--------------|--------------|--------------|---------------|
| <b>Biden</b>   | 896          | 354          | 281          | <b>1,531</b>  |
| <b>Clinton</b> | 310          | 235          | 162          | <b>707</b>    |
| <b>Harris</b>  | 309          | 138          | 94           | <b>541</b>    |
| <b>Trump</b>   | 4,503        | 566          | 3,091        | <b>8,160</b>  |
| <b>Total</b>   | <b>6,018</b> | <b>1,293</b> | <b>3,628</b> | <b>10,939</b> |

Table 23 Stance towards the individual candidates – counts

Furthermore, candidate-specific differences add another layer of complexity (see Table 23 and Figure 16). Joe Biden and Kamala Harris both receive a similar distribution of coverage (Biden: 58.5% neutral, 23.1% positive, 18.4% negative; Harris: 57.1% neutral, 25.5% positive, 17.4% negative), but Hillary Clinton stands out as a notable exception. Her coverage is less neutral (43.9%) and more positive (33.2%) than that of other Democrats. This suggests that evaluative dynamics vary not only between parties but also among individual candidates, challenging the idea of a uniform bias pattern.



*Figure 16 Stance toward the individual candidates - percent*

It is also essential to contextualize these findings within the broader methodological framework. As Bendel Larcher emphasizes, discourse analysis can only make claims about the specific discourses it investigates. The present results reflect the corpus analyzed – a set of quality newspapers across specific election cycles – and cannot speak to all outlets or contexts. It remains possible that a wider dataset, including tabloid media, television broadcasts, or different timeframes, would yield different results. Nevertheless, the data here strongly suggest that bias exists and that its distribution is neither trivial nor random.

#### 7.2.2. Evaluating Kuypers’ “Conspiracy of Shared Values” Claim

Kuypers (2014, 2020) offers a starkly different interpretation of media bias, arguing that mainstream outlets are structurally liberal due to a “conspiracy of shared values” among journalists. This thesis – that newsrooms overwhelmingly consist of left-leaning, highly educated elites who produce a liberal-leaning news agenda – is influential but contentious. The results of the present study offer little support for this claim.

When coverage is disaggregated by ideological orientation, the differences are remarkably small (see Table 24 and Figure 17). Left-leaning outlets produced 54.9% neutral coverage, compared to 53.6% in right-leaning outlets. Positive coverage accounted for 11.0% on the left and 13.7% on the right; negative coverage was 34.1% on the left and 32.7% on the right. These differences – all within a range of three percentage points – hardly indicate systemic liberal bias.

This pattern holds even in the more granular analysis of Trump coverage (see Table 12 and Figure 5 in Section 6.3). Left-leaning outlets are indeed slightly more negative toward Trump (by about 3.05 percentage points) and slightly less positive (by 3.7 points) than right-leaning ones. Yet the differences are minimal, and both ideological camps remain overwhelmingly negative in their portrayal of Trump. Even right-leaning newspapers, which Kuypers might predict to be sympathetic, devote roughly a third of their headlines to negative coverage of him. Similarly, coverage of Democratic candidates shows no significant ideological divergence: positive reporting is almost identical (27.6% left vs. 26.5% right), and negative reporting differs by only about four points (17.5% vs. 21.6%).

|                | Neutral      | Positive     | Negative     | Total         |
|----------------|--------------|--------------|--------------|---------------|
| <b>Left</b>    | 3,550        | 710          | 2,206        | <b>6,466</b>  |
| <b>Liberal</b> | 442          | 73           | 223          | <b>738</b>    |
| <b>Middle</b>  | 310          | 73           | 151          | <b>534</b>    |
| <b>Right</b>   | 1,716        | 437          | 1,048        | <b>3,201</b>  |
| <b>Total</b>   | <b>6,018</b> | <b>1,293</b> | <b>3,628</b> | <b>10,939</b> |

Table 24 Bias by outlet political alignment - counts

These findings challenge the “shared values” thesis in two key ways. First, they demonstrate that bias is not monopolized by one side of the political spectrum; rather, it is a structural feature of political reporting present across ideological contexts. Second, they suggest that evaluative tendencies are driven by factors beyond journalists’ personal politics – possibly including news values, audience expectations, or the inherent newsworthiness of Trump’s behavior – although such causal claims cannot be verified here.

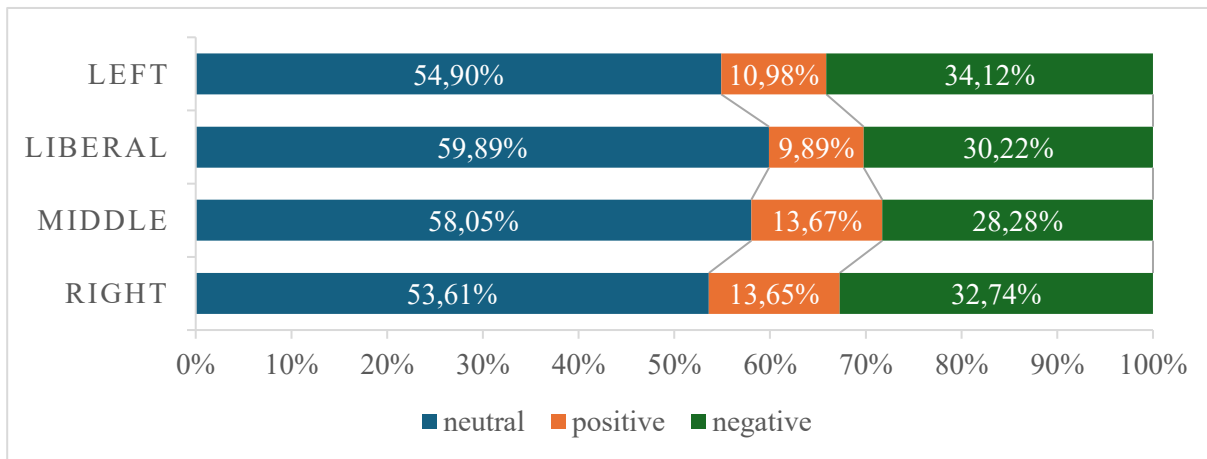


Figure 17 Bias by outlet political alignment - percent

The gap between perception and reality may help explain why Kuypers’ thesis persists despite empirical evidence to the contrary. The *hostile media effect* – the tendency of partisans to perceive neutral coverage as biased against their side – likely plays a role. Trump and his supporters frequently denounce the press as hostile, but this perception is not borne out by the data. With over half of coverage neutral and roughly equal levels of negativity from both left and right outlets, the notion of a uniformly liberal press is untenable with regard to headlines. Instead, the evidence points to a more complex and nuanced media landscape than Kuypers’ model allows.

### 7.2.3. Is Trump Treated “Unfairly”?

From the moment he launched his 2016 campaign, Donald Trump positioned himself in opposition to the media, repeatedly accusing mainstream outlets of bias and unfair treatment. The results of this study allow for a more measured assessment of that claim. On the one hand, it is true that Trump’s coverage is more negative than that of his opponents. With nearly 38% of all headlines critical of him, the press does not portray him in a uniformly positive light. However, this does not mean that the press is “overwhelmingly hostile.”

If hostility were truly overwhelming, one would expect negative reporting to comprise a majority of coverage. Instead, nearly two-thirds of all Trump-related headlines are either neutral or positive. Neutral coverage alone accounts for 55.2%, and positive coverage – while comparatively low – still appears in roughly 7% of cases. In this context, Trump’s claim of systemic hostility is difficult to sustain. Rather than an antagonistic press, the data depict a media environment in which neutrality dominates and negativity, while significant, does not define the whole.

The *hostile media effect* offers a plausible explanation for this disconnect. Individuals who feel strongly about an issue – especially political elites – often interpret even neutral reporting as biased against them. Trump’s combative rhetoric about the press likely amplifies this effect, fostering a perception of hostility that outstrips empirical reality. Moreover, such rhetoric shapes public discourse, fueling polarizing narratives about “fake news” and “enemy of the people” media.

This dynamic highlights a broader point: public debates about media bias are often inflamed more by political rhetoric than by actual content. The data from this study provide an important corrective to those debates. They show that while bias exists – and while Trump receives more negative coverage than his opponents – the notion of an overwhelmingly hostile press is not supported. Instead, the dominant feature of media coverage remains neutrality, a finding that challenges both Trump’s claims and broader populist narratives about media corruption.

### 7.3. Patterns of Bias and Naming Practices

A second major axis of interpretation concerns *how* bias manifests – not just in overall tone, but in the nuanced ways that candidates are named, framed, and morally evaluated. If the previous section demonstrated that the press is neither monolithically hostile nor systematically partisan, the micro-level analysis of linguistic strategies reveals the mechanisms by which evaluative tendencies are expressed. This section examines three interrelated dimensions: candidate-specific differences in evaluative tone, the predominance of neutral naming conventions and the selective use of value-laden labels, and the role of moralization as a framing strategy. Together, these findings illuminate how journalistic discourse constructs political meaning even when overt bias is rare.

#### 7.3.1. Candidate-Specific Differences and the Nature of Evaluative Bias

The candidate-level results reveal important asymmetries in how different actors were framed (see Table 23 and Figure 16 in Section 7.2.1). Although neutrality predominates across the corpus, the distribution of evaluative headlines is far from even. Donald Trump is uniquely characterized by a high volume of negative coverage (37.9%) and an exceptionally low share of positive headlines (6.9%), while Democratic candidates benefit from a more favorable balance. Joe Biden and Kamala Harris, for instance, receive roughly one quarter positive and less than one fifth negative coverage, while Hillary Clinton stands out as an anomaly, with less than half of her coverage neutral and a striking 33.2% positive – a figure that exceeds Trump’s positive share by more than fourfold. These disparities are not trivial. They indicate that while the press is broadly neutral, negative and positive bias do not fall evenly across candidates, but instead cluster in predictable ways that disadvantage Trump and, in Clinton’s case, confer a more favorable evaluative climate.

The qualitative evidence illustrates these tendencies vividly. Trump is more frequently cast as aberrant, dangerous, or unfit for office, as in headlines such as:

- (82) Trump, the Worst of America (*The New York Times*: October 17, 2016)
- (83) Trump vulgaire et sexiste: on ne civilise pas par décret<sup>55</sup> (*Le Figaro*: October 21, 2016)
- (84) Shameless Trump taking on the Fed (*The Irish Times*: September 27, 2016)

These examples encapsulate the evaluative tone that dominates much of Trump's coverage: they portray him in starkly negative terms, emphasizing moral failure (*worst of America*), character flaws (*vulgar and sexist*), or temperament (*shameless*). By contrast, Democratic candidates are more likely to be framed in positive or competence-oriented language:

- (85) Sharp lawyer knew all the right people Kamala Harris has navigated her steady path towards high office (*The Times*: November 6, 2024)
- (86) Biden, calm and firm, the right man to mend US (*The Age*, January 21, 2021)
- (87) Werner Vogels: 'Clinton goed geïnformeerd en bedachtzaam'<sup>56</sup> (*De Telegraaf*: November 5, 2016)

Such portrayals reinforce a competence-oriented narrative, presenting Democrats as capable, resilient, or statesmanlike – qualities rarely attributed to Trump in headline discourse. Hillary Clinton's portrayal is even more revealing. Her relatively high share of positive headlines reflects a more polarized framing, with strong support alongside moments of intense criticism:

- (88) Debate Takeaways: Hillary Clinton Digs In and Prevails (*The New York Times*: September 28, 2016)
- (89) Die Frauen lassen Hillary Clinton im Stich Wutbürger und wenig Gebildete für Trump<sup>57</sup> (*Neue Zürcher Zeitung*, November 11, 2016)

The contrast between praise for debate performance and criticism of her appeal to voters illustrates how candidate-specific narratives develop around key campaign moments. These subtleties complicate simplistic models of media bias and show that evaluative asymmetry is both real and multifaceted.

### 7.3.2. Naming Strategies and Value-Laden Labels

A striking feature of the corpus is the overwhelming prevalence of neutral naming conventions. Across more than 10,900 headlines, 98.1% adhered to conventional forms – surname alone or first name plus surname – with only 1.9% employing marked or evaluative alternatives. This finding aligns with previous scholarship (van Leeuwen 2008; Würth 2016), which has consistently shown that mainstream journalism favors neutral nomination practices in news discourse. Even in highly polarized political environments, journalists largely resist the temptation to embed judgment into the names of political actors.

Yet this statistical dominance of neutrality should not obscure the analytical significance of the exceptions. When marked naming does occur, it frequently serves as a powerful shorthand for evaluation. A closer qualitative examination of the labeled examples offers further insight into how these discursive asymmetries manifest. Trump was described using terms that targeted not only his political actions but also his character and behavior:

---

<sup>55</sup> Trump is vulgar and sexist: you can't civilize someone by decree

<sup>56</sup> Werner Vogels: 'Clinton is well-informed and thoughtful'

<sup>57</sup> Women abandon Hillary Clinton Angry citizens and the less educated vote for Trump

- (90) Groper in Chief (*The New York Times*: October 9, 2016)
- (91) Caricature de Reagan<sup>58</sup> (*Le Figaro*: January 19, 2017)
- (92) Trump, el salvador<sup>59</sup> (*Reforma*: January 21, 2025)

illustrate the range from strongly pejorative to laudative framing. Other examples such as

- (93) Shameless Trump (*The Irish Times*: September 27, 2016)
- (94) Revolting Trump (*The Australian*: October 10, 2016)

reflect explicitly moral judgments, whereas descriptors like

- (95) Vitaler Trump<sup>60</sup> (*Die Welt*: October 14, 2020)
- (96) TRUMP THE COMEBACK KING (*South China Morning Post*: November 7, 2024)

frame him more positively. Laudative labels were considerably rarer but still present, underscoring the diversity of evaluative possibilities even within a predominantly negative field. Similar patterns appeared, albeit less frequently, for Democratic candidates. Biden was referred to as

- (97) Résilient et compatissant<sup>61</sup> (*Le Figaro*: November 4, 2020)

while Clinton was described as

- (98) Frail Failure (*The New York Times*: October 12, 2016)

and Harris as

- (99) Deeply insecure (*The Irish Times*: November 5, 2024)

The asymmetry is thus not absolute – pejorative and laudative labels occur across the partisan divide – but the concentration of negative epithets in reference to Trump remains a robust feature of the data.

Despite the presence of such evaluative labels, their low overall frequency places them firmly in the realm of exceptional commentary rather than standard journalistic practice. Neutral naming conventions – for example, *President Biden*, *VP Harris*, or *Donald Trump* – remained the norm across outlets and languages. Moreover, no systematic use of ideological labels (such as *liberal* or *ultra-right*, or *conservative*) was detected, nor were there identifiable gendered patterns in naming. This absence is itself noteworthy: it suggests that journalists in this dataset rarely foreground ideological positioning or gender identity in headline-level discourse, preferring instead to frame evaluation through subtler lexical choices or through the selection of contextual modifiers.

### 7.3.3. Moralization as a Media Strategy

---

<sup>58</sup> A caricature of Reagan

<sup>59</sup> Trump, the savior

<sup>60</sup> Lively Trump

<sup>61</sup> Resilient and compassionate

The selective use of evaluative labels points to a broader pattern in the press's treatment of Trump and his opponents: the dominance of *moralization* over assessments of competence or policy. This finding resonates with Bhatia, Goodwin, and Walasek's (2018) observation that media discourse in the 2016 campaign was structured less by judgements of effectiveness than by judgements of character. The same dynamic is evident in the present dataset. Even when labels are rare, those that appear frequently target moral legitimacy rather than technical ability.

Consider, for example, the starkly moralizing tone of headlines such as:

- (100) Prepare for America to be set back 50 years – Jaded Clinton allowed mad-eyed opponent to set agenda (*The Irish Times*: November 10, 2016)
- (101) Disruptive, intolerant and populist icon: how Trump is seen as America's Mao (*South China Morning Post*: January 20, 2017)
- (102) Están jugando con nosotros – Donald Trump, un impresentable<sup>62</sup> (*El Mundo*: October 1, 2020)

Each of these headlines goes beyond policy disagreement or electoral contest. They frame Trump as regressive, intolerant, or fundamentally unfit – judgments that speak to perceived vice rather than to the success or failure of political strategy. The comparison to Mao, in particular, situates Trump within a moral universe defined by authoritarianism and extremism, while *impresentable* reduces his legitimacy to a question of character.

This emphasis on morality is not limited to pejorative framings. Positive labels, too, often invoke virtues such as resilience or salvation rather than competence:

- (103) Das Ende von Biden in den USA, Trump als Friedensbringer<sup>63</sup> (*taz*: January 20, 2025)
- (104) Beurzen zien Trump plots als heiland<sup>64</sup> (*De Volkskrant*: November 7, 2016)

The characterization of Trump as a *Friedensbringer* ('bringer of peace') or *heiland* ('savior') highlights how evaluative language, whether positive or negative, frequently invokes moral registers. Such framings reinforce the finding that political discourse around Trump is deeply moralized – not merely an assessment of policies or governance but a referendum on character, virtue, and legitimacy itself.

The implications of this pattern are significant. Even though marked naming and overt labeling are statistically rare, their moralizing content suggests that headlines participate in a broader discursive economy in which political actors are judged against ethical rather than merely political standards. This may help explain why public perceptions of Trump's coverage diverge so sharply from the data: even a minority of moralizing headlines can dominate public memory and shape narratives if they align with existing beliefs or are amplified by partisan actors.

Moreover, the rarity of ideological labeling underscores an important shift in how media bias is expressed. Rather than explicitly situating candidates within a left-right spectrum, the press often relies on implicit moral judgments – a strategy that resonates across cultural contexts and languages. This finding aligns with studies such as Cazzamatta and Hafez (2021), who documented how European magazines

---

<sup>62</sup> They are playing with us – Donald Trump, an unrepresentable individual

<sup>63</sup> The end of Biden in the US, Trump as peacemaker

<sup>64</sup> Stock markets suddenly see Trump as a savior

repeatedly cast Trump as *anti-democratic*, *corrupt*, and *racist*, and with Isakhan, Nwokora, and Pan (2019), who observed similar patterns in Arab coverage of the Muslim ban. The present analysis extends these observations to the headline level, demonstrating that even in their most condensed form, journalistic texts are saturated with moralized cues.

In all, patterns of bias in naming and framing are more complex than simple measures of positivity and negativity suggest. Although neutral naming overwhelmingly dominates, candidate-specific asymmetries persist, with Trump subject to disproportionately negative coverage and more frequent – though still rare – pejorative labels. These labels, when they occur, are less about ideology than morality, casting political actors as virtuous or corrupt, saviors or threats. The result is a discourse in which evaluative bias operates less through consistent slant and more through exceptional yet symbolically powerful acts of naming and moral judgment.

## 7.5. Temporal and Structural Dynamics of Coverage

### 7.5.1. Stability and Change Across Electoral Cycles

A central question in longitudinal analyses of media bias concerns the temporal stability of evaluative patterns across election cycles. Much of the existing literature has suggested that negativity toward Donald Trump was both persistent and structurally embedded throughout his political career. Patterson's (2016) study of the 2016 U.S. presidential election, for example, found that negativity in major U.S. outlets remained remarkably stable from the early primaries through to election day. Even short-term fluctuations – such as those following debates or major campaign events – did not alter the fundamental evaluative balance. Similar conclusions were reached by Nord, Mancini, and Gerli (2017) in their cross-European comparison, which showed that Trump's coverage remained more negative than Clinton's throughout the 2016 cycle, even if peaks and troughs reflected event-specific dynamics.

The present study partially confirms this stability thesis but also complicates it by revealing significant shifts in the distribution of tone between 2016, 2020, and 2024. As shown in Table 13 and Figure 6 in Section 6.4, the proportion of neutral coverage increased steadily across the three cycles, while negativity – though still the most frequent evaluative category aside from neutrality – declined in relative terms. During the 2016 election, evaluative coverage was marked by a pronounced degree of polarization rather than unidirectional negativity (see Table 25 and Figure 18). Trump's media portrayal was dominated by neutral reporting (53.8%), but the share of negative coverage (41.1%) was substantially higher than that of positive reporting (5.0%), underscoring the skepticism and critical scrutiny directed at his candidacy. Coverage of his Democratic opponent, Hillary Clinton, exhibited a different distribution: neutrality was lower (43.9%), while positive coverage (33.2%) and negative coverage (22.9%) were more evenly balanced, reflecting both supportive and critical narratives surrounding her campaign. These dynamics are consistent with observations of a highly charged and polarized media environment in 2016, in which Trump's unprecedented political style attracted substantial critical attention, while Clinton's portrayal was more balanced.

|                  | Neutral      | Positive     | Negative     | Total         |
|------------------|--------------|--------------|--------------|---------------|
| <b>Not Trump</b> | 1,515        | 727          | 537          | <b>2,779</b>  |
| <b>2016</b>      | 310          | 235          | 162          | <b>707</b>    |
| <b>2020</b>      | 776          | 337          | 154          | <b>1,267</b>  |
| <b>2024</b>      | 429          | 155          | 221          | <b>805</b>    |
| <b>Trump</b>     | 4,503        | 566          | 3,091        | <b>8,160</b>  |
| <b>2016</b>      | 1,574        | 147          | 1,203        | <b>2,924</b>  |
| <b>2020</b>      | 1,034        | 80           | 905          | <b>2,019</b>  |
| <b>2024</b>      | 1,895        | 339          | 983          | <b>3,217</b>  |
| <b>Total</b>     | <b>6,018</b> | <b>1,293</b> | <b>3,628</b> | <b>10,939</b> |

*Table 25 Bias by candidate by campaign year – counts*

By 2020, neutrality remained the dominant register for both candidates, yet the tonal balance shifted in notable ways. Trump continued to receive more negative (44.8%) than positive coverage (4.0%), with neutrality (51.2%) remaining roughly stable. Media reporting on Joe Biden, by contrast, was considerably more restrained: neutral coverage increased markedly to 61.3%, while the share of positive coverage declined to 26.6% and negative reporting dropped sharply to 12.2%. This shift suggests that Biden benefited from a comparatively measured media tone, whereas Trump’s portrayal continued to be shaped by controversy and critical evaluation – albeit within a less sensationalized framework than in 2016.

The 2024 campaign marked a further recalibration in journalistic approaches. Neutral reporting on Trump increased to its highest level across the three election cycles (58.9%), while negative coverage declined (30.6%) and positive coverage rose (10.5%). Coverage of Kamala Harris displayed a parallel trend: neutral reporting dominated (53.3%), positive coverage fell (19.3%), and negative reporting increased to 27.5%. These results indicate a gradual normalization of Trump’s media portrayal. Although still framed more negatively than his Democratic opponents, Trump was no longer treated as the unprecedented disruptor he was in 2016. Instead, reporting increasingly focused on his candidacy as one among several competing political phenomena, framed in a more dispassionate, reportorial tone.

It is important, however, to stress the limits of causal inference in this context. The available data cannot determine whether rising neutrality is the result of strategic newsroom decisions, changes in Trump’s political behavior, or broader transformations in the media ecosystem. It can only demonstrate that evaluative tone shifted over time in ways that complicate the notion of complete temporal stability posited by earlier scholarship. Nonetheless, this finding is important. It suggests that media bias is not a static property of journalistic discourse but a dynamic phenomenon that evolves in response to political context, audience expectations, and the institutional self-perceptions of the press.

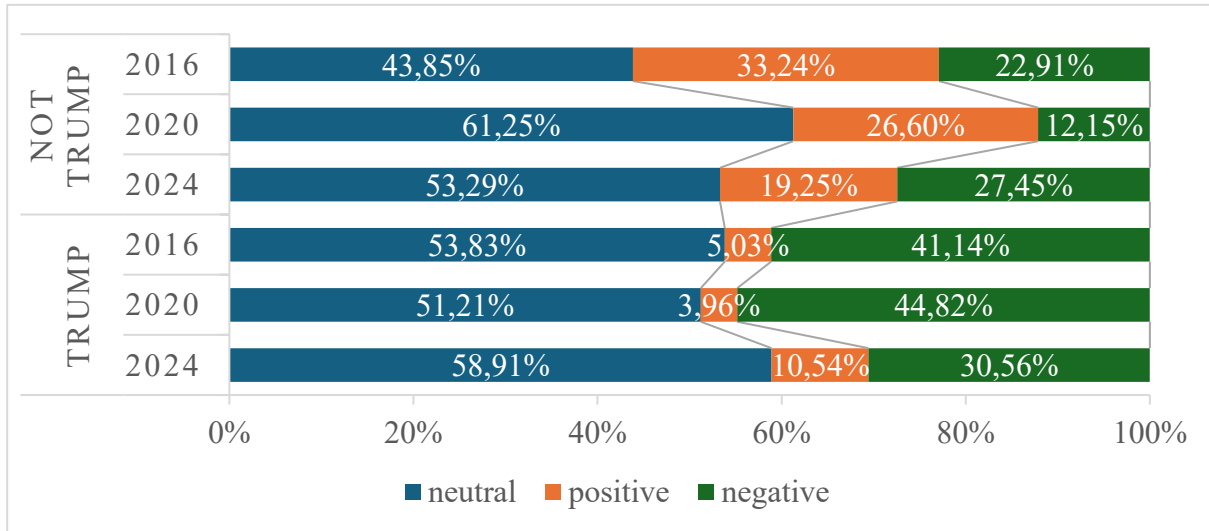


Figure 18 Bias by candidate by campaign year - percent

### 7.5.2. Gatekeeping and Coverage Bias

If changes in evaluative tone reveal shifting discursive norms, patterns of volume and visibility expose the structural dynamics of media attention. Across all three electoral cycles, Donald Trump dominated the news agenda by a substantial margin. As shown in Table 25 in the previous section, Trump accounted for nearly three times as many headlines as all Democratic candidates combined – 8,160 versus 2,779 – a difference confirmed as statistically significant by a chi-square goodness-of-fit test,  $\chi^2(1, N = 10,939) = 2646.97, p < .001$ . This overwhelming visibility reflects two interrelated phenomena identified in media bias research: *gatekeeping* and *coverage bias* (D'Alessio & Allen 2000).

*Gatekeeping* refers to the preferential selection of certain stories over others – in this case, the disproportionate focus on Trump-related events and statements. *Coverage bias*, by contrast, describes the unequal distribution of journalistic attention once a topic has been selected. Both dynamics are evident here: not only did newspapers consistently prioritize Trump-related news, but they also devoted far more space and frequency to covering him than to any of his Democratic challengers. This imbalance mirrors the findings of Patterson (2016) and Nord, Mancini, and Gerli (2017), who likewise documented Trump's outsized media presence during the 2016 campaign, but it extends these observations across three electoral cycles and multiple linguistic contexts.

The consequences of this structural imbalance are the following. First, disproportionate visibility amplifies Trump's agenda-setting power. Even when coverage is neutral or negative, the sheer volume of reporting ensures that Trump remains at the center of public attention. This visibility effect can shape perceptions of political importance: as agenda-setting theory predicts, frequent mention of a candidate signals relevance, potentially elevating their salience in public consciousness regardless of evaluative tone.

Second, the dominance of Trump-related headlines has methodological implications for interpreting bias metrics. Because Trump headlines so heavily outnumber those about his opponents, they inevitably shape the aggregate distribution of evaluative categories. For example, if Trump's coverage is more negative, the overall negativity of the dataset will be skewed upward, even if coverage of other candidates is more positive. Conversely, if neutrality rises in Trump coverage (as it does in 2024), it will influence the total

proportion of neutral headlines across the entire corpus. Scholars must therefore interpret aggregate bias figures with caution, contextualizing them within the structural asymmetry of coverage volume.

Finally, this imbalance underscores a deeper tension in democratic media systems. On one hand, newsworthiness – defined by novelty, conflict, and audience interest – naturally gravitates toward a figure like Trump, whose statements and actions reliably generate public engagement. On the other hand, such asymmetry risks distorting the democratic information environment by crowding out coverage of alternative candidates and policy positions. This is particularly consequential for figures like Kamala Harris, who received comparatively limited headline attention, potentially affecting public visibility and perceived viability.

In sum, the structural dynamics of Trump coverage reveal both continuity and change. The evaluative tone of reporting shifted over time, with rising neutrality and declining negativity complicating prior claims of temporal stability. Yet the dominance of Trump in media attention remained consistent, producing significant gatekeeping and coverage bias that shaped the contours of political discourse. Together, these findings underscore a central conclusion of this study: media bias is not only a matter of tone but also of structure and scale – and both dimensions must be analyzed to understand the press’s role in shaping contemporary electoral politics.

## **7.6. European Perspectives: A Mosaic Rather Than a Monolith**

The question of whether European press discourse exhibits shared tendencies in its portrayal of Donald Trump and other U.S. presidential candidates lies at the heart of hypotheses H7–H10. These hypotheses anticipated that European media, despite their linguistic, cultural, and political diversity, would display a broadly consistent evaluative orientation – one that might distinguish them from other regions and suggest the presence of a shared European discourse. The results of the present study provide partial support for this assumption. Across the ten European countries represented in the dataset, an overarching pattern emerges: neutral coverage is often dominant, negative reporting outweighs positive reporting, and positive framing remains a marginal phenomenon – at least when looking at the aggregate numbers (see Table 19 and Figure 12 in Section 6.10). At the same time, closer inspection reveals notable cross-national variation in magnitude, distribution, and evaluative tone. These divergences complicate any straightforward claim of a unified “European voice” and suggest instead that European press discourse, while structurally similar in broad outline, is characterized by substantial internal heterogeneity.

One of the most robust findings is the predominance of neutral reporting. In coverage of Trump, nine out of ten European countries classified more than half of all headlines as neutral, with six surpassing the 56% threshold (see Table 26 and Figure 19/Figure 20). Even in the United Kingdom – the only outlier – neutrality still formed the largest share of reporting, albeit slightly below the majority mark at 47.4%.

|                    | Neutral      | Positive   | Negative     | Total        |  |                    | Neutral | Positive | Negative | Total        |
|--------------------|--------------|------------|--------------|--------------|--|--------------------|---------|----------|----------|--------------|
| <b>Biden</b>       | 571          | 215        | 191          | 977          |  | <b>Harris</b>      | 191     | 93       | 75       | <b>359</b>   |
| <b>Bulgaria</b>    | 28           | 12         | 4            | 44           |  | <b>Bulgaria</b>    | 9       | 5        | 1        | <b>15</b>    |
| <b>France</b>      | 40           | 9          | 7            | 56           |  | <b>France</b>      | 19      | 3        | 2        | <b>24</b>    |
| <b>Germany</b>     | 26           | 8          | 4            | 38           |  | <b>Germany</b>     | 8       | 8        | 4        | <b>20</b>    |
| <b>Ireland</b>     | 64           | 48         | 21           | 133          |  | <b>Ireland</b>     | 25      | 12       | 12       | <b>49</b>    |
| <b>Italy</b>       | 19           | 2          | 6            | 27           |  | <b>Italy</b>       | 9       |          | 2        | <b>11</b>    |
| <b>Netherlands</b> | 30           | 15         | 11           | 56           |  | <b>Netherlands</b> | 14      | 6        | 4        | <b>24</b>    |
| <b>Poland</b>      | 86           | 16         | 31           | 133          |  | <b>Poland</b>      | 18      | 13       | 18       | <b>49</b>    |
| <b>Spain</b>       | 43           | 26         | 17           | 86           |  | <b>Spain</b>       | 11      | 12       | 6        | <b>29</b>    |
| <b>Switzerland</b> | 16           | 13         | 14           | 43           |  | <b>Switzerland</b> | 1       | 4        | 1        | <b>6</b>     |
| <b>UK</b>          | 219          | 66         | 76           | 361          |  | <b>UK</b>          | 77      | 30       | 25       | <b>132</b>   |
|                    |              |            |              |              |  |                    |         |          |          |              |
| <b>Clinton</b>     | 172          | 135        | 107          | 414          |  | <b>Trump</b>       | 2,594   | 338      | 1,948    | <b>4,880</b> |
| <b>Bulgaria</b>    | 0            | 0          | 0            | 0            |  | <b>Bulgaria</b>    | 95      | 16       | 47       | <b>158</b>   |
| <b>France</b>      | 14           | 4          | 5            | 23           |  | <b>France</b>      | 151     | 18       | 93       | <b>262</b>   |
| <b>Germany</b>     | 15           | 13         | 10           | 38           |  | <b>Germany</b>     | 195     | 23       | 127      | <b>345</b>   |
| <b>Ireland</b>     | 22           | 26         | 15           | 63           |  | <b>Ireland</b>     | 289     | 42       | 213      | <b>544</b>   |
| <b>Italy</b>       | 14           | 5          | 5            | 24           |  | <b>Italy</b>       | 113     | 8        | 53       | <b>174</b>   |
| <b>Netherlands</b> | 9            | 10         | 7            | 26           |  | <b>Netherlands</b> | 176     | 38       | 84       | <b>298</b>   |
| <b>Poland</b>      | 1            | 0          | 0            | 1            |  | <b>Poland</b>      | 337     | 44       | 174      | <b>555</b>   |
| <b>Spain</b>       | 20           | 12         | 7            | 39           |  | <b>Spain</b>       | 232     | 40       | 190      | <b>462</b>   |
| <b>Switzerland</b> | 10           | 9          | 9            | 28           |  | <b>Switzerland</b> | 158     | 24       | 111      | <b>293</b>   |
| <b>UK</b>          | 67           | 56         | 49           | 172          |  | <b>UK</b>          | 848     | 85       | 856      | <b>1,789</b> |
|                    |              |            |              |              |  |                    |         |          |          |              |
| <b>Total (all)</b> | <b>3,528</b> | <b>781</b> | <b>2,321</b> | <b>6,630</b> |  |                    |         |          |          |              |

Table 26 Bias by country and candidate - counts

This finding mirrors results from prior scholarship emphasizing the predominance of descriptive, factual reporting in mainstream political coverage (Patterson 2016; Nord, Mancini & Gerli 2017) and supports Holler's (2023) observation that neutrality typically outweighs evaluative tone in cross-national press discourse. It indicates that, despite increasing political polarization and rhetorical intensification, European newspapers remain committed to professional norms of descriptive reporting, particularly in headline framing.

This shared emphasis on neutrality is accompanied by a persistent asymmetry between negative and positive coverage. Across all countries, negative reporting on Trump ranged from 28.2% to 47.9%, whereas positive coverage never exceeded 12.8%. A different imbalance appears in the treatment of Democratic candidates. For Hillary Clinton, neutral coverage fell below the 50% threshold in several contexts – Germany (39.8%), Ireland (34.9%), the Netherlands (34.6%), Switzerland (35.7%) and the UK (39.0%) – with negative and positive coverage varying widely between countries. Coverage of Joe Biden was more neutral overall, surpassing 60% in France, Poland, and Italy, with more positive than negative coverage in all countries but Italy, Poland, Switzerland, and the UK. Kamala Harris's portrayal showed another kind of pattern, albeit one that is based on a much smaller dataset: neutral coverage dominated in Bulgaria (60.0%), France (79.2%), Ireland (51.0%), Italy (81.8%), the Netherlands, and the UK (both 58.3%), while in Germany neutral and positive coverage were equally distributed at 40.0% each and in Poland neutral and negative coverage even at 36.7% each. In both Spain (41.4%) and Switzerland (66.7%) positive coverage was dominant.

These differences point to the enduring influence of national journalistic traditions, political cultures, and media-system logics, even within a shared European media landscape: the result is a mosaic rather than a monolith.

## **7.7. Comparative Perspective: Relating to Existing Scholarship**

Situating the present findings within the broader body of scholarship on Trump, media bias, and political reporting reveals a complex picture of continuity and change. While many patterns identified by previous research are corroborated – particularly regarding Trump's disproportionate visibility and persistent negative framing – the results of this thesis nuance, refine, and in some cases complicate those conclusions. Rising neutrality, variation across national contexts, and the uneven intensity of evaluative framing all point to a more differentiated reality than earlier studies suggest. The following sections systematically revisit the major contributions discussed in Section 3.3 and examine their relationship to the present results.

### **7.7.1. Patterson (2016): Visibility and Negativity Revisited**

Patterson's (2016) study of U.S. media coverage during the 2016 election cycle established two core empirical patterns that have shaped subsequent debate: Trump's unparalleled visibility and the overwhelmingly negative tone of his coverage. According to Patterson's analysis, Trump received approximately 15% more attention than Hillary Clinton, and his coverage was relentlessly critical, ranging from roughly two-to-one to more than ten-to-one in negative-to-positive ratios across major outlets. This heightened salience and sustained negativity, Patterson argued, were structural features of the media environment rather than mere reflections of episodic controversies.

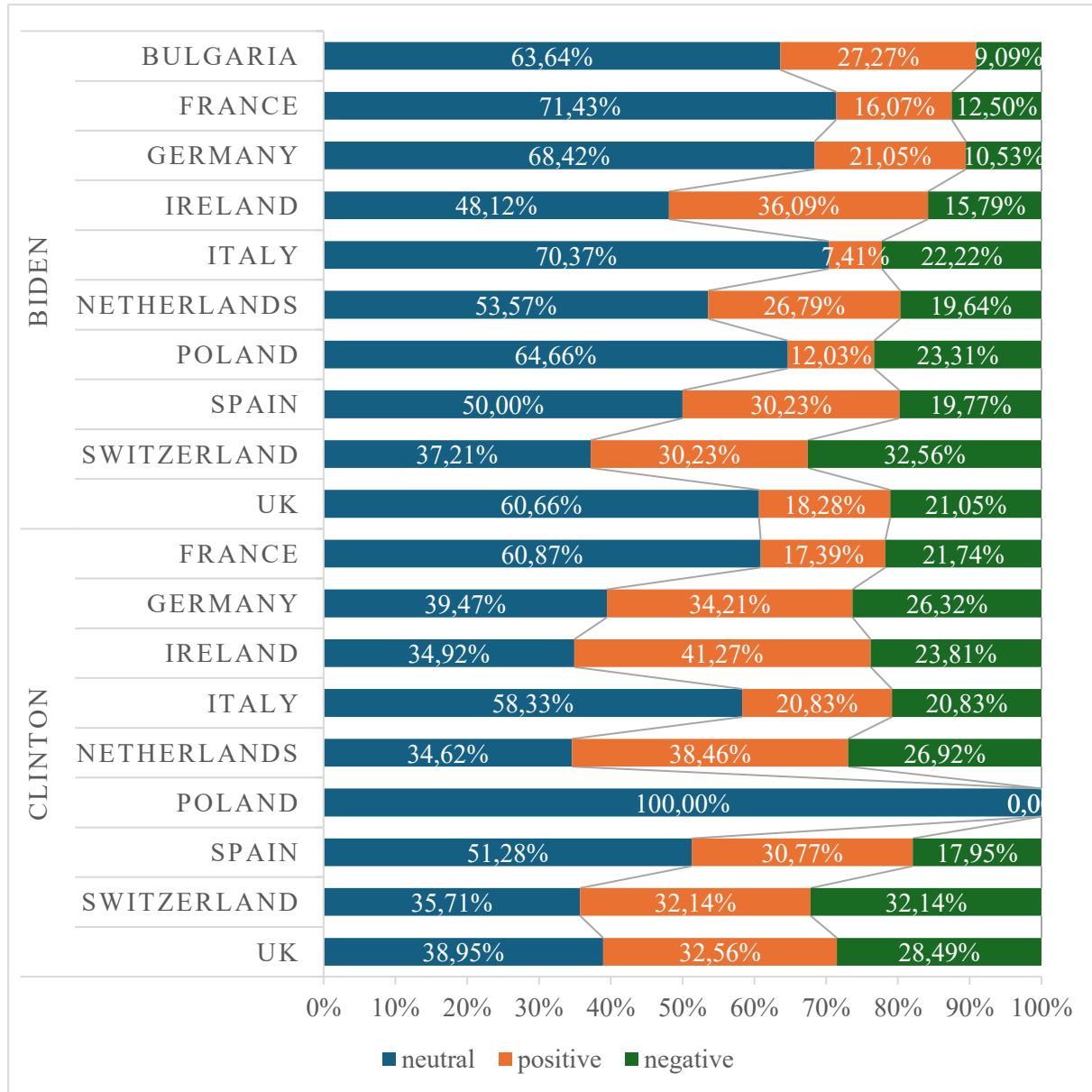


Figure 19 Bias by country and candidate - percentage

The present study confirms both of Patterson's key observations but also introduces important qualifications. First, Trump's dominance in media attention is even more pronounced in the cross-national context: across the three electoral cycles and 10,939 coded headlines, Trump was the subject of nearly three times as many headlines as all Democratic candidates combined (8,160 versus 2,779; see Table 25 in Section 7.5.1). This significant imbalance, confirmed statistically by a chi-square goodness-of-fit test ( $\chi^2$  (1, N = 10,939) = 2646.97,  $p < .001$ ), clearly illustrates both *gatekeeping* and *coverage bias* (D'Alessio & Allen 2000).

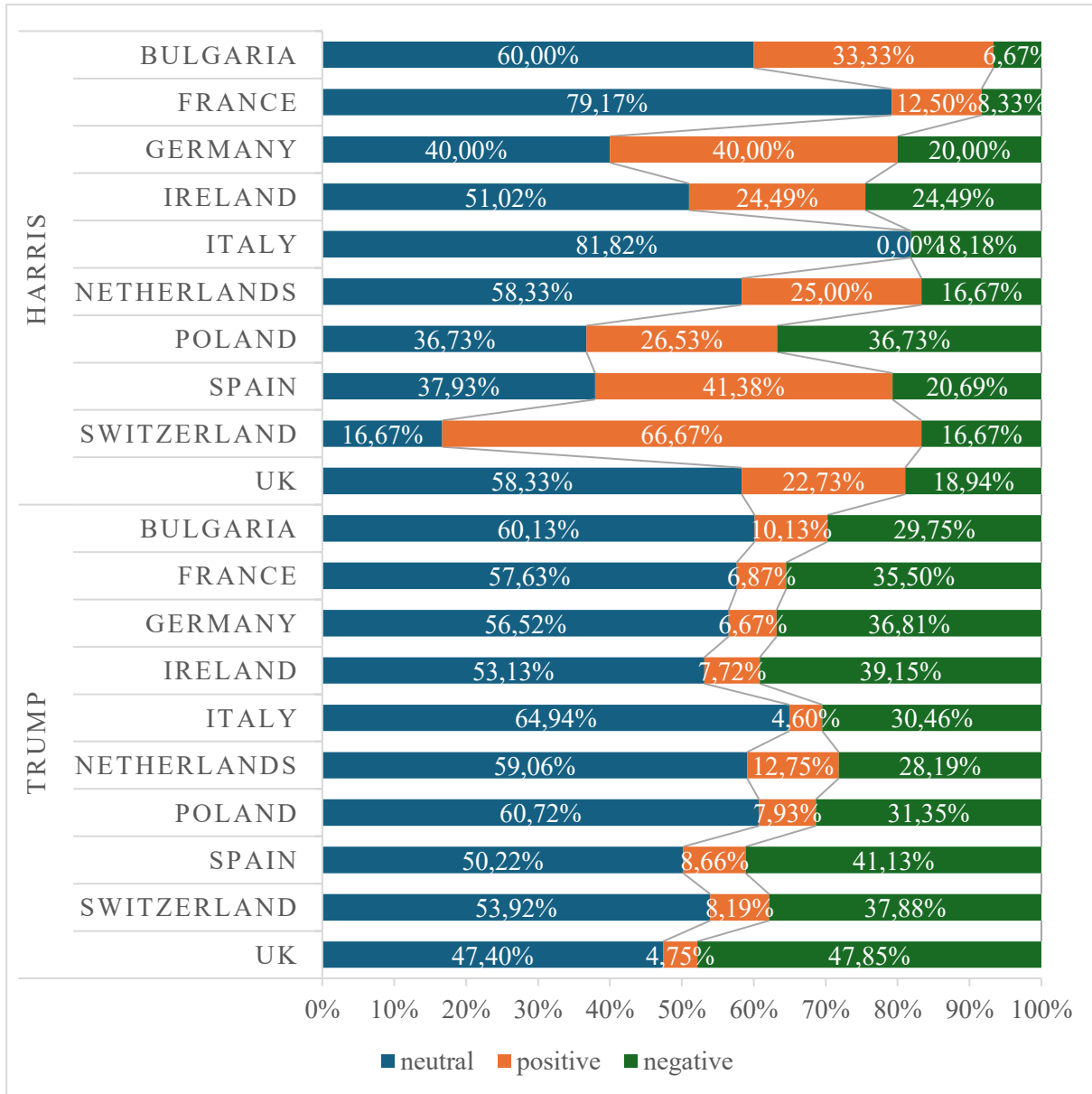


Figure 20 Bias by country and candidate - percentage

At the same time, however, the evaluative balance differs from Patterson's results in important ways. Whereas he found near-universal negativity, the present findings show that neutrality dominates across all cycles. Approximately 55.2% of Trump's coverage is neutral overall, with negative coverage accounting for 37.9% and positive coverage only 6.9% (see Table 11 and Figure 4 in Section 6.2). These results confirm that Trump continues to be scrutinized more harshly than his opponents – whose coverage is both more neutral and more positive – but they complicate the notion of a uniformly adversarial press. The rise in neutrality suggests that while negativity remains a defining feature, journalistic discourse may be shifting toward a more balanced evaluative distribution, possibly reflecting professional norms of objectivity or a recalibration of media strategy in response to bias accusations.

### 7.7.2. Nord, Mancini & Gerli (2017): Negative Dominance and the Neutrality Turn

The comparative study by Nord, Mancini, and Gerli (2017) similarly identified negative dominance as a transnational feature of Trump coverage. Analyzing reporting in Italy, Sweden, and the United Kingdom, they found that Trump's coverage was more negative than Clinton's in all three countries – 66% in the UK, 62% in Italy, and 49% in Sweden – while Clinton's tone varied more widely and included more positive elements. This study established that the pattern observed by Patterson was not confined to the U.S. media system but extended across European contexts, thereby supporting the notion of a cross-national consensus on Trump.

The findings of the present study align with Nord, Mancini, and Gerli's conclusions in terms of direction but diverge significantly in terms of magnitude. Negative framing does indeed outweigh positive framing across the European press, but the proportion of neutral coverage is far higher than Nord and colleagues reported. In nine out of ten European countries examined, neutral coverage of Trump exceeded 50%, with six surpassing 56% (see Table 26 and Figure 20 in Section 7.6). Even in the United Kingdom, where reporting was most polarized, neutrality accounted for 47.4% of headlines – a substantial share that challenges the narrative of negativity as the dominant mode. Negative reporting, by contrast, ranged from 28.2% to 47.9%, and positive coverage remained marginal, generally below 13%.

These results suggest a “neutrality turn” in European political journalism. The persistence of a negative skew is undeniable, but the scale of negativity appears to have moderated since 2016, replaced by a more balanced distribution that prioritizes descriptive reporting over overt evaluation. This shift might reflect the institutionalization of Trump as a political figure over time – an initial wave of moralized condemnation giving way to more routinized coverage – but definitive causal explanations lie beyond the scope of the present data.

### 7.7.3. Cazzamatta & Hafez (2021): Moralized Critique and Selective Labeling

Cazzamatta and Hafez (2021) emphasize a second dimension of Trump coverage that goes beyond tone: the moralization of discourse. Their study of *Der Spiegel*'s coverage during Trump's presidency documents a persistent framing of Trump as *anti-democratic*, *racist*, *corrupt*, *sexist*, and *disdain[ful of] science*. Importantly, they argue that this moralized register reflects not only ideological opposition but also a broader shift in media discourse from policy-based critique to character-based condemnation.

The present findings strongly support this dimension of their argument. Even though overtly evaluative labeling is rare – appearing in only 1.9% of naming instances – when such labels do occur, they overwhelmingly adopt a moral or characterological frame. Headlines such as

(105) The bad boy is no president But it doesn't matter that Clinton the honor student crushed Trump (*USA Today*: September 28, 2016)

(106) Surely revolting Trump can't bounce back from this (*The Australian*: October 10, 2016)  
exemplify how the press reduces complex political realities to simplified narratives of vice and threat.

At the same time, the moralizing frame is not exclusively pejorative. Some headlines cast Trump in redemptive or heroic terms, such as

(107) FIGHTER TRUMP WILL KEEP THROWING PUNCHES (*The Australian*: September 28, 2020)

(108) Trump, el nuevo ‘hombre teflón’<sup>65</sup> (*El Pais*: November 10, 2016)

These instances, though rare, show that moralization can also serve to elevate rather than denigrate, illustrating the dual valence of value-laden discourse. Nevertheless, the preponderance of negative moral labels supports Cazzamatta and Hafez’s contention that media framing often transcends policy evaluation and enters the realm of ethical judgement.

#### 7.7.4. Hinck (2024): Global Negativity and Regional Variation

Hinck’s (2024) comparative study of election coverage across Chinese, Russian, Iranian, and Saudi media extends the literature beyond Western contexts, demonstrating that negative framing of Trump is a global phenomenon but that its intensity varies substantially by region. Chinese and Iranian outlets, for example, exhibited the most negative coverage ( $m = -0.30$  and  $-0.34$ , respectively), while Russian media approached neutrality ( $m = -0.07$ ). European outlets, by contrast, displayed a more consistent negativity, suggesting a distinctive regional discourse.

The results of the present study broadly corroborate Hinck’s conclusions while adding further nuance. Negativity toward Trump does indeed span regions, but the dominance of neutrality in the dataset complicates the picture of a monolithically hostile global press. Moreover, while European coverage is consistently more negative than positive, significant intra-European variation exists, undermining the notion of a single “European voice.” Differences between countries – and even between outlets within the same country – suggest that national media systems, political cultures, and journalistic traditions mediate the intensity and framing of negativity. This evidence partially supports Hinck’s hypotheses on global patterns but also indicates that variation is not confined to cross-regional comparisons: it exists within Europe as well.

#### 7.7.5. Bykova et al. (2018) and Rovinskaya & Greydina (2021): Temporal Shifts and the Softening of Tone

Finally, Bykova, Gavra, and Slutskiy (2018) and Rovinskaya and Greydina (2021) highlight an important temporal dimension: shifts in evaluative tone over time. Russian media, for example, initially portrayed Trump positively, framing him as a strong leader constrained by domestic elites, but later coverage became more negative, particularly in relation to NATO, Ukraine, and impeachment. Similarly, Chinese media oscillated between positive and negative tones depending on the state of bilateral relations.

The present findings lend support to the notion of temporal variation but suggest a different trajectory: rather than an intensification of negativity, there is evidence of softening over time. Across the three electoral cycles, neutrality increased while negativity declined, indicating a broader recalibration in media framing. This does not mean that criticism disappeared – Trump remains more negatively covered than his opponents – but it does suggest a normalization of his political presence and a corresponding shift toward less moralized, more descriptive coverage. Whether this reflects changing journalistic strategies, audience expectations, or the evolving nature of Trump’s political role is beyond the scope of this study, but the trend is clear: the press’s evaluative stance toward Trump is neither static nor monolithic.

---

<sup>65</sup> Trump, the new ‘Teflon man’

### 7.7.6. Synthesis: Confirmation, Complication, and Contribution

Taken together, the present study's findings broadly confirm the central claims of prior scholarship: Trump receives disproportionate attention, his coverage is more negative than that of his rivals, and moralized framing is a recurring feature of journalistic discourse. Yet these findings also complicate and refine those narratives in several ways. Negativity remains pervasive but is now accompanied by – and in many cases outweighed by – neutrality. Patterns identified in earlier studies persist but are less stark, suggesting a gradual rebalancing of media tone. Furthermore, significant cross-national and temporal variations challenge the idea of a uniform global response to Trump and point to the importance of contextual factors in shaping evaluative discourse.

Most importantly, these results demonstrate the value of large-scale, cross-linguistic headline analysis in enriching our understanding of media bias. By capturing subtle shifts in tone, volume, and framing across more than 10,900 headlines, this study not only corroborates key insights from the literature but also reveals complexities that smaller-scale or single-country studies could not detect. The result is a more nuanced, empirically grounded picture of how Trump is represented in the transnational press.

## 7.8. Methodological Reflections and AI-Assisted Analysis

A distinctive contribution of this study lies not only in its empirical findings but also in its methodological approach. The use of artificial intelligence (AI) as a tool for large-scale discourse analysis represents an important innovation in political communication research as well as linguistics and carries significant implications for how future studies may approach multilingual, cross-national corpora. This section reflects on the role AI played in the research process, evaluates its strengths and limitations, and considers its broader methodological consequences for discourse-analytic scholarship.

### 7.8.1. Reliability and Performance of AI Coding

The AI-based classification of headlines proved to be a robust and reliable method for large-scale analysis. Across more than 10,900 headlines from seventeen countries and six continents, the AI consistently produced accurate results in the identification of bias categories (neutral, positive, negative) and in the application of van Leeuwen's (2008) naming framework. No major discrepancies were observed between AI-generated classifications and subsequent human verification in the key categories of evaluative stance, demonstrating that large-scale automated coding can yield valid and replicable results when carefully designed and supervised.

In particular, the AI performed exceptionally well in the bias classification task. Its evaluations closely aligned with those of human coders and with established findings in the literature, correctly identifying both the prevalence of neutrality and the asymmetry in positive and negative coverage. This performance confirms the potential of AI to handle large datasets that would otherwise be prohibitively time-consuming for human researchers. Coding more than 10,900 headlines manually would have required weeks of continuous work and likely would have restricted the scope of the study to a much smaller and less representative sample. The efficiency of AI processing thus enabled the analysis to include multiple languages, diverse national contexts, and three election cycles – a level of breadth and depth that would have been difficult to achieve with human coding alone.

### 7.8.2. Naming, Value Judgement, and the Limits of Prompt Design

While the AI's overall performance was strong, the results also highlight the importance of prompt specificity and the limitations of automated coding when tasks require nuanced linguistic interpretation. In the "name as given" category, for example, the AI successfully captured most conventional naming practices – such as *US President-elect Trump*, *VP Harris*, or *Clinton Foundation* – but it failed to consistently identify marked or evaluative naming forms such as *vitaler Trump*<sup>66</sup>, *shameless Trump*, or *Trump, el salvador*<sup>67</sup>. This shortfall was not due to a fundamental incapacity of the model but rather to the prompt's insufficient precision: the instruction did not fully convey the need to record all evaluative modifiers attached to a candidate's name. A more refined prompt, developed through iterative testing, would likely have yielded more complete results. This illustrates an important methodological lesson: successful AI-based discourse analysis depends as much on prompt engineering as on the capabilities of the model itself.

A similar issue arose in the application of value judgment labels. The AI's assignments – such as classifying *Groper in Chief* as pejorative or *Le résilient et compatissant Joe Biden*<sup>68</sup> as laudative – were generally accurate and aligned with human judgements. However, because the "name as given" output was not always exhaustive, it remains unclear in some cases whether the AI based its evaluations on the candidate's naming itself or on other lexical cues within the headline. This ambiguity slightly reduces the interpretive transparency of the results. Nonetheless, the overall validity of the evaluative classifications remains intact: even if the exact textual anchor was not always as envisioned, the AI's judgements accurately reflected the evaluative tone of the headlines.

### 7.8.3. Workflow, Human Oversight, and Practical Considerations

The practical aspects of AI deployment also provide important methodological insights. The coding process proved most reliable when limited to approximately 50 headlines per batch. Larger inputs occasionally resulted in incomplete outputs, with some headlines being skipped – a problem likely attributable to token or memory limitations in the model. Additionally, the AI occasionally failed to code consecutive headlines with similar structures, requiring manual verification to ensure that no data were omitted. These issues underline the necessity of human oversight: while AI can automate and accelerate large-scale coding, it cannot yet fully replace human supervision, particularly when subtle errors could accumulate over a large dataset.

The need for human involvement also extends to the interpretive stage. As Bendel Larcher (2023) reminds us, discourse analysis must remain closely attuned to the specificities of the material under investigation. AI models, even when highly accurate, cannot infer intent, causality, or broader sociopolitical context. Their output must therefore be critically interpreted within a theoretically informed analytical framework. In this study, the collaboration between AI-generated coding and human interpretation proved highly effective: the AI handled large-scale categorization tasks with speed and consistency, while human expertise ensured accuracy, contextual sensitivity, and theoretical depth.

---

<sup>66</sup> Lively Trump

<sup>67</sup> Trump the savior

<sup>68</sup> The resilient and compassionate Joe Biden

#### 7.8.4. Expanding Analytical Scope: Scale, Language, and Representativeness

One of the most significant methodological advantages of AI-assisted analysis lies in its ability to expand the scale, scope, and representativeness of discourse research. Without automation, this study would have been limited to English, German, and Italian – the languages the researcher could code manually – and to a sample of perhaps 2,000–3,000 headlines. By contrast, the use of AI enabled the inclusion of headlines in Spanish, French, Dutch, Portuguese, Bulgarian, and Polish, producing a corpus of over 10,900 headlines. This expansion dramatically enhanced the comparative power of the analysis, enabling the exploration of cross-linguistic and cross-regional patterns that would have remained inaccessible through manual methods alone.

The multilingual capacity of AI also points to a broader methodological shift in discourse analysis. As natural language processing tools improve, scholars are increasingly able to work with primary materials in their original languages rather than relying on translations – a significant advantage given the subtle but often consequential differences in evaluative language across linguistic contexts. At the same time, this potential is not without risks: translation models, while useful, still introduce errors and interpretive biases, and human language expertise remains indispensable for accurate interpretation.

#### 7.8.5. Broader Implications for Discourse Analysis

The findings of this study demonstrate that AI can serve as a powerful companion tool in discourse analysis, particularly when dealing with large, multilingual, and cross-temporal datasets. It is not a replacement for traditional qualitative methods but a complement to them – one that allows researchers to scale up their analyses while retaining theoretical nuance. Crucially, AI's greatest strengths lie in tasks that are repetitive, time-intensive, and rule-based (such as bias classification or the identification of naming strategies), while human researchers remain essential for interpretive work, contextualization, and theory-building.

The use of AI in this study also highlights the need for further methodological research. Questions remain about the optimal balance between automation and human oversight, the most effective ways to design prompts for discourse-analytic tasks, and the ethical considerations of using machine learning models in the humanities and social sciences. As AI tools become more sophisticated, these questions will become increasingly central to the discipline.

## 8. Conclusion

### 8.1. Rethinking Media Discourse in a Changing Democratic Landscape

This thesis set out to examine how major newspapers across Europe and beyond have represented Donald J. Trump and his Democratic rivals during the last three U.S. presidential election cycles (2016, 2020, and 2024), and what these patterns reveal about bias, discourse, and political communication in a transforming media environment. It also sought to test the methodological potential of artificial intelligence (AI) as a tool for large-scale discourse analysis, exploring how automation can support the study of complex, multilingual datasets. At its core, the study was motivated by a simple but far-reaching research question: How do mainstream news media construct political meaning in their portrayal of Trump and his competitors – and what does this tell us about contemporary media discourse, democratic legitimacy, and the evolving nature of political communication? The findings presented here offer a nuanced and at times surprising answer. Across more than 10,900 headlines from seventeen countries and six continents,

neutrality – not negativity or positivity – emerged as the dominant stance. While Trump was indeed framed more negatively and less positively than his rivals, the scale of this bias was considerably smaller than public narratives or Trump’s own rhetoric would suggest. Patterns of ideological slant were detectable but relatively modest, with both left- and right-leaning outlets largely adhering to neutral reporting norms. Naming conventions, too, were overwhelmingly neutral, and even when evaluative labels appeared, they were rare and often highly moralized. Above all, the results show that political discourse – even around polarizing figures like Trump – is shaped as much by structural dynamics, discursive strategies, and institutional norms as by overt partisan bias. At the same time, this study has underscored the complexity of media discourse in a rapidly changing information environment. The dominance of neutrality does not mean that media coverage is free of bias; rather, it suggests that bias manifests in subtler ways – through asymmetries in tone, frequency, moral framing, and agenda-setting. It also suggests that public perceptions of media hostility, often amplified by political actors themselves, may reflect psychological phenomena such as the *hostile media effect* more than the empirical reality of press coverage. In this sense, the findings offer a corrective to both Trump’s accusations of a “fake news” media and broader populist critiques of journalistic bias.

## 8.2. Reflexive Considerations: The Researcher’s Role in Discourse

“Zu einer vollständigen Diskursanalyse gehört, dass man den kritischen Blick auch auf die eigene Tätigkeit wirft und diese reflektiert. Was haben wir gemacht und welche diskursiven Auswirkungen hat unser Schreiben?”<sup>69</sup> (Bendel Larcher 2023: 261). This reminder – that discourse analysis must include reflection on its own conditions of production – is essential for any responsible conclusion. Throughout this thesis, care has been taken to ensure methodological transparency and reflexivity: the selection of newspapers, the design of the coding schema, and the use of AI for large-scale analysis were all documented in detail to enable reproducibility and invite scrutiny. Nevertheless, no analysis is free from the researcher’s positionality. Decisions about which sources to include, which frameworks to prioritize, and which questions to ask are inevitably shaped by one’s academic background, language competencies, and the constraints of data availability. These choices, in turn, shape the contours of the results. Efforts were made to mitigate individual bias – particularly through the use of AI to perform large-scale, rule-based coding – but even AI-assisted analysis is not immune to interpretive assumptions embedded in prompt design or category definitions. Moreover, the scope of this study was limited to mainstream media due to their accessibility, comparability, and relevance. As a result, it does not capture the full complexity of the contemporary information environment, particularly the alternative media ecosystems where much of the polarization surrounding Trump thrives. Acknowledging these limitations is not a weakness but a necessary component of scholarly rigor. It underscores the provisional nature of discourse-analytic knowledge and situates this thesis as one contribution among many to an ongoing conversation about media, politics, and democracy. It also highlights the importance of critical self-reflection: by examining how our own methodological and epistemological choices shape the narratives we construct, we participate in the very discursive processes we seek to analyze.

## 8.3. The Broader Significance of the Findings

The results of this study speak to several fundamental debates in political communication and media studies. First, they challenge simplified narratives about media bias. Despite widespread claims of systemic hostility, particularly from populist actors, the data reveal a media landscape where neutrality

---

<sup>69</sup> A complete discourse analysis requires that we also take a critical look at our own activities and reflect on them. What have we done, and what discursive effects does our writing have?

predominates and overt partisanship is relatively rare. This is not to say that bias is absent – evaluative asymmetries are real, especially in Trump’s case – but they are far less pronounced than public discourse often assumes. This finding has significant implications for our understanding of democratic legitimacy and trust in media institutions. It suggests that declining trust among certain political groups (Meeks 2020; Mourão et al. 2018) may stem less from the actual content of news reporting than from strategic delegitimization campaigns, selective exposure to partisan media, and broader shifts in the political information environment.

Second, the study contributes to the field of Eurolinguistics by demonstrating the value – and the limits – of cross-national comparative approaches. While certain patterns, such as the dominance of neutrality and the relative scarcity of positive coverage, recur across European contexts, significant national variations remain. These differences likely reflect deeper structural factors – such as media-system typologies, political cultures, and journalistic traditions – that shape how global political events are framed. In this sense, the European media landscape is best understood not as a unified voice but as a mosaic of overlapping but distinct discursive logics.

Third, the findings shed light on the evolving role of morality in political discourse. The study shows that even when overtly evaluative language is rare, it often takes a moral rather than ideological form. Trump is frequently cast not just as ineffective or controversial but as corrupt, dangerous, or unfit – and occasionally as a savior or redeemer. Such moralization has far-reaching implications: it personalizes political debate, elevates questions of character over policy, and shapes how legitimacy itself is constructed and contested in democratic societies.

Finally, the study demonstrates the transformative potential of AI-assisted analysis in the humanities and social sciences. By enabling the systematic coding of a large, multilingual corpus, AI made it possible to identify patterns and trends that would have been difficult or impossible to detect through manual methods alone. At the same time, the challenges encountered – from prompt design limitations to the need for human oversight – highlight the importance of combining computational tools with critical, theory-informed interpretation. Together, these insights point toward a future in which human and machine work collaboratively to advance discourse-analytic scholarship.

#### **8.4. Future Research Directions**

The findings and limitations of this study open up multiple avenues for future research. One important direction concerns the role of non-mainstream and partisan media. As Mourão et al. (2018) and Ross and Rivers (2018) suggest, much of the polarization surrounding media trust and perceptions of bias occurs outside the mainstream press. Examining how Trump and his rivals are portrayed in alternative media ecosystems – including far-right, far-left, and conspiratorial outlets – would provide a more comprehensive picture of the discursive environment and its effects on democratic legitimacy. A second avenue concerns the interactive dynamics between media, elites, and publics. Future research might investigate how media framing interacts with political actors’ communication strategies – particularly Trump’s use of social media to delegitimize mainstream journalism – and how audiences interpret and respond to different forms of coverage. This would help clarify whether neutrality, negativity, or moralization in news discourse actually shapes public opinion or merely reflects pre-existing partisan alignments. A third area of exploration is the application of AI and other computational methods to discourse analysis. While this study demonstrates the promise of AI-assisted approaches, further methodological work is needed to refine prompt design, improve cross-linguistic performance, and enhance the interpretive transparency of automated coding. Future projects could also integrate AI-based discourse analysis with network

analysis, sentiment tracking, or multimodal approaches to capture the full complexity of contemporary media ecologies. Finally, the question of global comparative perspective remains far from exhausted. Extending this research to include additional regions – particularly the Global South, where media systems often operate under different political and economic constraints – would deepen our understanding of how global power dynamics shape media discourse about U.S. politics and its actors.

## 8.5. Concluding Reflections

This thesis set out to investigate how mainstream media portray one of the most polarizing figures in contemporary politics and, in doing so, to illuminate the discursive mechanisms through which political meaning is constructed in the public sphere. The results complicate simplistic narratives of partisan hostility and instead reveal a more complex reality: one in which neutrality predominates, bias manifests in subtle but significant ways, and moralization plays a powerful role in shaping public perceptions. They also underscore the methodological promise of AI-assisted analysis and the importance of reflexivity in discourse research. In the end, this study is as much about media discourse as it is about the act of studying it. As Bendel Larcher (2023: 261) reminds us, a complete discourse analysis must also reflect critically on its own practices – on what it has done, how it has done it, and what discursive effects its writing might have. The present thesis has attempted to do precisely that: to interrogate not only how the media constructs political reality, but also how scholarly analysis itself participates in that construction. It is hoped that this research, with all its findings and limitations, will contribute to a more nuanced, empirically grounded, and self-reflexive understanding of media discourse in an era when democracy itself is being contested, reimagined, and redefined.

## 9. Reflection on the Use of AI in the Scholarly Process

The integration of generative artificial intelligence (AI) tools such as ChatGPT into academic research and writing represents one of the most significant shifts in scholarly practice in decades. As the guidelines of the Catholic University of Eichstätt-Ingolstadt (KU)<sup>70</sup> emphasize, generative AI is a “Taschenrechner-Moment” for disciplines that rely heavily on textual production and analysis – a transformative development comparable to the introduction of the calculator in mathematics (2). Such tools are no longer peripheral: they are increasingly embedded into text processing software, search functions, and research workflows. Rather than resisting their presence, the university rightly calls for a “responsible and critically reflected integration” of AI into scholarly work (2). This reflection follows that call. It documents and critically assesses how I used ChatGPT in the research and writing of my MA thesis, explains my rationale for doing so, and evaluates the opportunities and limitations of AI in the academic process. The aim is to provide a transparent, methodologically sound account of my engagement with AI – an engagement that was always guided by the principles of academic integrity, critical thinking, and intellectual ownership.

The genesis of my thesis topic was entirely my own intellectual work. I conceived the research field, the central question, and the conceptual framework independently. Nevertheless, at the earliest stage of the project, I integrated ChatGPT into my brainstorming process. Alongside traditional methods – including but not limited to extensive searches via Google Scholar and JSTOR, on-site library consultations, and manual note-taking – I used ChatGPT as a “study buddy” to explore possible research avenues and hypotheses. Formulating concise hypotheses proved particularly challenging. By asking ChatGPT for suggestions on possible research questions or subtopics, I was able to view my field from alternative

---

<sup>70</sup>[https://www.ku.de/fileadmin/1903/Personalentwicklung/Hochschuldidaktik/Handreichung\\_KI\\_050925.pdf](https://www.ku.de/fileadmin/1903/Personalentwicklung/Hochschuldidaktik/Handreichung_KI_050925.pdf)

perspectives and consider aspects I might not have noticed otherwise. Crucially, this was never a process of passive acceptance. As Bondi (2025: 56) argues, the key to effective use of AI in writing lies in treating its output as suggestions rather than solutions – suggestions that must be critically accepted or rejected. This is precisely how I approached the tool: I critically evaluated each idea, keeping only those that aligned with my vision for the project and discarding the rest. In this way, AI acted as a catalyst for intellectual refinement, not a substitute for it.

The core of the literature review process remained firmly rooted in traditional scholarship. Over several weeks, I identified, read, and annotated a wide range of scholarly sources – far more than those ultimately cited in the final thesis. ChatGPT played a limited, supplementary role here: occasionally, I used it to clarify complex theoretical concepts or to summarize dense academic arguments as part of my note-taking. However, I was fully aware of the limitations of generative AI in this context. As the KU guidelines stress, AI-generated text is based on probabilistic language models rather than actual knowledge, meaning that outputs can be plausible yet factually inaccurate (1). To mitigate this risk, I double-checked all AI-generated explanations against the original scholarly sources before integrating them into my thinking. In addition, I crafted detailed, specific prompts – a process that required me to articulate my own ideas precisely, thereby reinforcing my critical voice. When AI responses diverged from my intentions, I reformulated my prompts or abandoned the output altogether. This constant interaction forced me to rethink and refine my arguments, making the process more dialogical and reflective.

The most sustained and intensive use of AI occurred during the drafting phase. Here, I approached ChatGPT as a writing tool – a means of generating, structuring, and revising text – but never as a substitute for my own scholarly agency. I drafted the thesis piece by piece with the support of ChatGPT, re-writing and reworking each section multiple times until the final version fully reflected my own voice and intellectual intent. This approach resonates with Bondi's (2025: 57) concept of digital literacy, which includes not only technical proficiency but also “control of the agency of the process” and “retention of decision-making rights.” At no point did I relinquish authorship or accountability. Every sentence was critically reviewed, revised, and, where necessary, rewritten. This allowed me to concentrate more deeply on substantive aspects such as argumentation, literature integration, and hypothesis development.

I find the comparison between AI and the calculator particularly apt. Just as a calculator does not replace the conceptual understanding required to solve a mathematical problem, AI does not replace the intellectual work of research and argumentation. It simply accelerates certain mechanical aspects of writing. Using AI was not a “shortcut.” If anything, it intensified my engagement: crafting effective prompts required a deep understanding of my material, and the iterative cycle of questioning and revision pushed me to think more critically about every argument I presented.

The potential risks of using AI in academic writing – including hallucinated information, loss of authorial voice, over-reliance, and ethical concerns – are well-documented in the literature (Lin et al. 2023). I was conscious of these risks throughout the process and took deliberate steps to mitigate them. The most significant safeguard was rigorous human oversight. Every AI-generated sentence was scrutinized for accuracy, relevance, and consistency with my argument. Whenever I encountered dubious information, I traced it back to authoritative sources or replaced it entirely. This aligns with Bondi's (2025: 63) warning that AI-generated texts are only *plausible*, not necessarily *accurate*, and must therefore be critically validated. A second risk was the potential dilution of my personal writing voice. To counteract this, I edited every section until it authentically reflected my analytical perspective and rhetorical style. This process also deepened my understanding of *discursive identity* — a crucial aspect of scholarly communication that Bondi (2025: 63) identifies as essential to authenticity.

Perhaps the most significant outcome of this process was a shift in how I understand academic writing itself. By delegating certain mechanical tasks (e.g., sentence formulation, structural suggestions) to AI, I was able to devote more time and energy to the intellectual core of the project: interpreting sources, refining arguments, and situating my research within the broader scholarly conversation. Moreover, by engaging critically with AI outputs, I sharpened my analytical skills and became more aware of how arguments are constructed, framed, and communicated. From a professional standpoint, this experience was also invaluable. As someone preparing to enter the teaching profession, I believe it is essential to understand both the possibilities and limitations of AI. Students will inevitably use such tools – with or without guidance. My responsibility as an educator is to help them do so ethically, critically, and effectively. Mastering these tools myself is therefore not just a personal advantage but a professional necessity. Refusing to engage with AI would mean ignoring a technology that, when used responsibly, can strengthen academic work. My thesis is testament to that: it is stronger because of the ways in which AI enriched – rather than replaced – my scholarly process.

In conclusion, my use of AI in the thesis was extensive but always critically controlled, transparent, and methodologically grounded. I used ChatGPT at every stage – from refining research questions and clarifying literature to drafting and revising chapters – but I remained fully responsible for the intellectual substance, structure, and final form of the work. The KU guidelines emphasize that the value of a scholarly work lies not in the mere production of text but in its contribution to knowledge (6). By treating AI as a tool for exploration, reflection, and refinement, I believe I fulfilled that principle: my thesis is a product of human creativity, judgment, and critical thinking – supported, but never replaced, by artificial intelligence.

*Kristina Holler*  
 DE-85072 Eichstätt  
 Kristina.Holler@email.de

## References

All URLs were checked 9 April 2026.

- Ad Fontes (2025), *Media Bias Chart*, <https://app.adfontesmedia.com/chart/interactive> (5 September 2025).
- AllSides (2025), *Media Bias*, <https://www.allsides.com/media-bias> (5 September 2025).
- Akdenizli, Banu / Özçetin, Burak / Gündoğdu, Nazli (2021), “Long Time Allies No More? How News Media in Turkey Covered the United States in 2019”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: the Trump Effect*, 273–289, Lanham: Rowman & Littlefield.
- Allam, Rasha (2021), “The Portrayal of the Trump Administration in the Egyptian Media”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 97–104. Lanham: Rowman & Littlefield.
- Arroyave, Jesus (2021), “The Magnate Became President: How the Colombian Media and Citizens Perceive the United States and Trump”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 81–96, Lanham: Rowman & Littlefield.
- Baghestan, Abbas Ghanbari / Akoje, Topic Peremobowei (2021), “Media Coverage of the United States and Trump in Nigeria: A Mixture of ‘Greatness’ and ‘Arrogance’”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 215–226, Lanham: Rowman & Littlefield.
- Baghestan, Abbas Ghanbari / Osman, Mohd Nizam (2021), “The Image of Donald Trump in Major Malaysian Newspapers: A Deeper ‘Mistrust’ Than Obama’s Time”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 189–200, Lanham: Rowman & Littlefield.
- Bendel Larcher, Sylvia (2023), *Linguistische Diskursanalyse: Ein Lehr- und Arbeitsbuch*. 2nd edn., Tübingen: Gunter Narr Verlag.

- Bhatia, Sudeep / Goodwin, Geoffrey P. / Walasek, Lukasz (2018), "Trait Associations for Hillary Clinton and Donald Trump in News Media: A Computational Analysis", *Social Psychological and Personality Science* 9.2: 123–130. <https://doi.org/10.1177/1948550617751584>.
- Blevens, Fred (2022), "From Chaos and Cage Fighting to Quiet and Calm", in: Gutsche, Robert E., *The Future of the Presidency, Journalism, and Democracy*, 238–254, 1st edn., London: Routledge.
- Bondi, Marina (2025), *AI and Writing: Challenges and Opportunities for ESP and Professional Communication*, Zenodo. <https://doi.org/10.5281/ZENODO.15250813>.
- Bucher, Hans-Jürgen (2017), "Massenmedien als Handlungsfeld I: Printmedien", in: Roth, Kersten Sven / Wengeler, Martin / Ziem, Alexander (eds.), *Handbuch Sprache in Politik und Gesellschaft*, 298–333, [Handbücher Sprachwissen 19], Berlin / Boston: De Gruyter.
- Bykova, Elena / Gavra, Dmitrii / Slutskiy, Pavel (2018), "The Evolution of Trump's Image in Russian Media", *American Behavioral Scientist* 0.0 <https://doi.org/10.1177/0002764218793691>.
- Cazzamatta, Regina / Hafez, Kai (2021), "Mediated Populism: The US Image in German Media", in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 127–144, Lanham: Rowman & Littlefield.
- Cohen, Hart / Gurney, Myra / Castillo, Antonio (2021), "What's He Done Now? Seeing the United States and Trump Through the Australian Media", in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 55–68, Lanham: Rowman & Littlefield.
- D'Alessio, Dave / Allen, Mike (2000), "Media Bias in Presidential Elections: A Meta-Analysis", *Journal of Communication* 50.4: 133–156. <https://doi.org/10.1111/j.1460-2466.2000.tb02866.x>.
- Enli, Gunn (2017), "Twitter as Arena for the Authentic Outsider: Exploring the Social Media Campaigns of Trump and Clinton in the 2016 US Presidential Election", *European Journal of Communication* 32.1: 50–61. <https://doi.org/10.1177/0267323116682802>.
- Eurotopics Media (2025), *Eurotopics*, <https://www.eurotopics.net/en/142186/media> (5 September 2025).
- Fan, Yaxin / Jiang, Feng / Li, Peifeng / Li, Haizhou (2024), "Uncovering the Potential of ChatGPT for Discourse Analysis in Dialogue: An Empirical Study", in: Calzolari, Nicoletta / Kan, Min-Yen / Hoste, Veronique / Lenci, Alessandro / Sakti, Sakriani / Xue Nianwen (eds.), *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, 16998–17010, Torino: ELRA and ICCL. <https://aclanthology.org/2024.lrec-main.1477/>.
- Feldman, Lauren (2014), "The Hostile Media Effect", in: Kenski, Kate / Jamieson, Kathleen Hall (eds.), *The Oxford Handbook of Political Communication*, 549–564, 1st edn., Oxford: Oxford University Press.
- Fowler, Roger (1991), *Language in the News: Discourse and Ideology in the Press*, London / New York: Routledge.
- Garg, Ryan / Han, Jaeyoung / Cheng, Yixin / Fang, Zheng / Swiecki, Zachari (2024), "Automated Discourse Analysis via Generative Artificial Intelligence", in: *Proceedings of the 14th Learning Analytics and Knowledge Conference*, 814–820, Kyoto: ACM. <https://doi.org/10.1145/3636555.3636879>.
- Giessen, Hans / Szwed, Iwona (2025), "Discourses in Europe", in: Grzega, Joachim, *The Routledge Handbook of Eurolinguistics*, 247–261, 1st edn., London: Routledge.
- Grzega, Joachim (2013), *Studies in Europragmatics: Some Theoretical Foundations and Practical Implications*, [Eurolinguistische Arbeiten 7], Wiesbaden: Harrassowitz Verlag.
- Grzega, Joachim (2016a), "Alternative European Values in European Headlines: Competitiveness, Privatization, Solidarity, Socialization, Welfare State", *Journal for EuroLinguistiX* 13: 18–27. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/grzega-161.pdf>.
- Grzega, Joachim (2016b), "Limites, liberté, sécurité – Accès académiques aux pensées de différents Européens au bénéfice d'un public général", *Journal for EuroLinguistiX* 13: 90–100. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/grzega-162.pdf>.
- Grzega, Joachim (2017), "Qualitative and Quantitative Comments on Peace and War from a Eurolinguistic Perspective", *Journal for EuroLinguistiX* 14: 48–62. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/grzega-171.pdf>.
- Grzega, Joachim (2019), "Names, Nations and Nature from a Eurolinguistic View: Notes on Peace-Threatening and Peace-Promoting Language", *Journal for EuroLinguistiX* 16: 23–36. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/grzega-191.pdf>.
- Grzega, Joachim (2020), "Fake News als Wort-Waffe in der Corona-Thematik: Lexikalische, semantische und diskursanalytische Kommentare zu ausgewählten Medien Deutschlands und Österreichs sowie anderer Länder", *Journal for EuroLinguistiX* 17: 1–14. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/grzega-201.pdf>.
- Grzega, Joachim (2023), "Conflict or Peace Language? Cover Texts of 2022 from 17 European and Non-European Countries", *Journal for EuroLinguistiX* 20: 12–30. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/grzega-231.pdf>.
- Grzega, Joachim (2025), "Basic Definitions of Eurolinguistics", in: Grzega, Joachim, *The Routledge Handbook of Eurolinguistics*, 3–16, 1st edn., London: Routledge.

- Gutsche, Robert E. (2015), *Media Control: News as an Institution of Power and Social Control*, New York / London / Oxford / New Delhi: Bloomsbury Academic.
- Hanusch, Nora (2014), "From Words to War: Eine Analyse des metaphorischen Sprachgebrauchs internationaler Printmedien vor Ausbruch des Irakkrieges 2003", *Journal for EuroLinguistiX* 11: 44–50. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/hanusch-141.pdf>.
- Happer, Catherine / Hoskins, Andrew / Merrin, William (eds.) (2019), *Trump's Media War*, Cham: Springer International Publishing.
- Hayes, Andrew F. / Krippendorff, Klaus (2007), "Answering the Call for a Standard Reliability Measure for Coding Data", *Communication Methods and Measures* 1.1: 77–89. <https://doi.org/10.1080/19312450709336664>
- Higdon, Nolan / Marmol, Emil / Huff, Mickey (2022), "Returning to Neoliberal Normalcy", in: Gutsche, Robert E., *The Future of the Presidency, Journalism, and Democracy*, 255–273, 1st edn., London: Routledge.
- Hinck, Robert (2024), "From Political Unknown to an Unwanted Incumbent: Comparing Media Coverage of the 2020 and 2016 U.S. Presidential Election Within Nondemocratic Media", *American Behavioral Scientist* 68.1: 112–141. <https://doi.org/10.1177/00027642231171882>.
- Hippler, Nina (2016), "Der Syrien-Krieg in den Medien – Eurolinguistische Analysen", *Journal for EuroLinguistiX* 13: 28–80. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/hippler-161.pdf>.
- Holler, Kristina Beatrice (2023), "Biased Reporting on China in European Headlines? A Corpus Analysis of Four Newspapers from Large and Small Countries", *Journal for EuroLinguistiX* 20: 1–11. <http://www1.ku-eichstaett.de/SLF/EngluVglSW/ELiX/holler-231.pdf>
- Isakhan, Benjamin / Nwokora, Zim / Pan, Chengxin (2019), "Perceptions of Democracy and the Rise of Donald Trump: A Framing Analysis of Saudi Arabian Media", *Global Media and Communication* 15.2: 159–175. <https://doi.org/10.1177/1742766519846630>.
- Jarvis, Jeff (2019), "Trump and the Press: A Murder-Suicide Pact", in: Happer, Catherine / Hoskins, Andrew / Merrin, William (eds.), *Trump's Media War*, 23–30, Cham: Springer International Publishing.
- Kellner, Douglas (2019), "Trump's War Against the Media, Fake News, and (A)Social Media", in: Happer, Catherine / Hoskins, Andrew / Merrin, William (eds.), *Trump's Media War*, 47–68, Cham: Springer International Publishing.
- Kuypers, Jim A. (2014), *Partisan Journalism: A History of Media Bias in the United States*, [Communication, Media, and Politics], Lanham: Rowman & Littlefield.
- Kuypers, Jim A. (2020), *President Trump and the News Media: Moral Foundations, Framing, and the Nature of Press Bias in America*, [Lexington Studies in Political Communication], Lanham: Lexington Books.
- Leeuwen, Theo van (2008), *Discourse and Practice: New Tools for Critical Discourse Analysis*, [Oxford Studies in Sociolinguistics], Oxford: Oxford University Press.
- Lichter, S. Robert (2014), "Theories of Media Bias", in: Kenski, Kate / Jamieson, Kathleen Hall (eds.), *The Oxford Handbook of Political Communication*, 403–416, 1st edn., Oxford: Oxford University Press.
- Lin, Shu-Min / Chung, Hsin-Hsuan / Chung, Fu-Ling / Lan, Yu-Ju (2023), "Concerns About Using ChatGPT in Education", in: Huang, Yueh-Min / Rocha, Tânia (eds.), *Innovative Technologies and Learning*, 37–49, [Lecture Notes in Computer Science 14099], Cham: Springer Nature Switzerland.
- Lozano, José Carlos / Martínez-Garza, Francisco Javier (2021), "Coverage of Donald Trump in Three Mexican National Newspapers: El Universal, Reforma, and La Jornada, 2017–2019", in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 201–214, Lanham: Rowman & Littlefield.
- Marzi, Giacomo / Balzano, Marco / Marchiori, Davide (2024), "K-Alpha Calculator – Krippendorff's Alpha Calculator: A User-Friendly Tool for Computing Krippendorff's Alpha Inter-Rater Reliability Coefficient", *MethodsX* 12: 102545. <https://doi.org/10.1016/j.mex.2023.102545>.
- Media Bias Fact Check (2025), *Media Bias Fact Check*, <https://mediabiasfactcheck.com/> (5 September 2025).
- Meeks, Lindsey (2020), "Defining the Enemy: How Donald Trump Frames the News Media", *Journalism & Mass Communication Quarterly* 97.1: 211–234. <https://doi.org/10.1177/1077699019857676>.
- Montgomery, Martin (2019), *Language, Media and Culture: The Key Concepts*, London / New York: Routledge.
- Morini, Marco (2020), *Lessons from Trump's Political Communication: How to Dominate the Media Environment*, [Political Campaigning and Communication], Cham: Springer International Publishing.
- Mourão, Rachel R. / Thorson, Esther / Chen, Weiyue / Tham, Samuel M. (2018), "Media Repertoires and News Trust During the Early Trump Administration", *Journalism Studies* 19.13: 1945–1956. <https://doi.org/10.1080/1461670X.2018.1500492>.
- Nord, Lars / Mancini, Paolo / Gerli, Matteo (2017), *The Exceptional Election: Press Coverage of Clinton and Trump in Italy, Sweden and the UK*, Sundsvall: Mid Sweden University.
- Ostyakova, Lidiia / Mikhailova, Anna / Konovalov, Vasily (2025), "Redefining Annotation Practices: Leveraging Large Language Models for Discourse Annotation", in: Panchenko, Alexander / Gubanov, Dmitriy / Khachay, Michael / Kutuzov, Andrey / Loukachevitch, Natalia / Kuznetsov, Andrey / Nikishina, Irina et al. (eds.), *Analysis of Images, Social Networks and Texts*, 131–147, [Lecture Notes in Computer Science 15419], Cham: Springer Nature Switzerland.

- Patterson, Thomas E. (2016), “News Coverage of the 2016 General Election: How the Press Failed the Voters”, *HKS Faculty Research Working Paper Series* RWP16-052, December 2016. <https://www.hks.harvard.edu/publications/news-coverage-2016-general-election-how-press-failed-voters>.
- Pludowski, Tomasz (2021), “Poland and the United States Now Serve Each Other as Mirrors: The Trump Effect”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 243–258, Lanham: Rowman & Littlefield.
- Ross, Andrew S. / Rivers, Damian J. (2018), “Discursive Deflection: Accusation of ‘Fake News’ and the Spread of Mis- and Disinformation in the Tweets of President Trump”, *Social Media + Society* 4.2. <https://doi.org/10.1177/2056305118776010>.
- Rovinskaya, Juliana / Greydina, Nadezhda (2021), “The US Image and Donald Trump Through Russian Eyes”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 259–272, Lanham: Rowman & Littlefield.
- Scheufele, Dietram A. / Iyengar, Shanto (2014), “The State of Framing Research: A Call for New Directions”, in: Kenski, Kate / Jamieson, Kathleen Hall (eds.), *The Oxford Handbook of Political Communication*, 619–632, 1st edn., Oxford: Oxford University Press.
- Siiman, Leo A. / Rannastu-Avalos, Meeli / Pöysä-Tarhonen, Johanna / Häkkinen, Päivi / Pedaste, Margus (2023), “Opportunities and Challenges for AI-Assisted Qualitative Data Analysis: An Example from Collaborative Problem-Solving Discourse Data”, in: Huang, Yueh-Min / Rocha, Tânia (eds.), *Innovative Technologies and Learning*, 87–96, [Lecture Notes in Computer Science 14099], Cham: Springer Nature Switzerland.
- Simpson, Paul / Mayr, Andrea / Statham, Simon (2019), *Language and Power: A Resource Book for Students*, 2nd edn., [Routledge English Language Introductions], London / New York: Routledge.
- Stokes, Jane (2021), “Donald Trump: Balloon Baby or Toilet Head? Satire and Protest in the Press and on the Streets of London”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 119–126, Lanham: Rowman & Littlefield.
- Sun, Lizhou / Song, Fei (2021), “Comprehensive and Balanced: Unveiling the Portrayal of the United States and Trump Through the Chinese Media 2017–2019”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 69–80, Lanham: Rowman & Littlefield.
- Symons, Alex (2019), “Trump and Satire: America’s Carnavalesque President and His War on Television Comedians”, in: Happer, Catherine / Hoskins, Andrew / Merrin, William (eds.), *Trump’s Media War*, 183–198, Cham: Springer International Publishing.
- Taylor, Charlotte (2019), “Conflict and Categorisation: A Corpus and Discourse Study of Naming Participants in Forced Migration”, in: Evans, Matthew / Jeffries, Lesley / O’Driscoll, Jim (eds.), *The Routledge Handbook of Language in Conflict*, 128–144, [Routledge Handbooks in Applied Linguistics], London / New York: Routledge.
- Trump, Donald (2016), “The failing @nytimes is truly one of the worst newspapers. They knowingly write lies and never even call to fact check. Really bad people!” X. <https://x.com/realDonaldTrump/status/709089928324489216>. (19 September, 2025).
- Trump, Donald (2017), “The FAKE NEWS media (failing @nytimes, @NBCNews, @ABC, @CBS, @CNN) is not my enemy, it is the enemy of the American People!” X. <https://x.com/realDonaldTrump/status/832708293516632065>. (19 September, 2025).
- Weaver, David H. / Choi, Jihyang (2014), “The Media Agenda: Who (or What) Sets It?”, in: Kenski, Kate / Jamieson, Kathleen Hall (eds.), *The Oxford Handbook of Political Communication*, 359–376, 1st edn., Oxford: Oxford University Press.
- Wikipedia (2025), *Folha de S. Paulo*, [https://en.wikipedia.org/wiki/Folha\\_de\\_S.Paulo](https://en.wikipedia.org/wiki/Folha_de_S.Paulo) (5 September 2025).
- Wikipedia (2025), *O Globo*, [https://en.wikipedia.org/wiki/O\\_Globo](https://en.wikipedia.org/wiki/O_Globo) (5 September 2025).
- Würth, Anne (2016), “Manipulation bei Bezeichnungen von Regierungschefs in europäischen Zeitungsüberschriften?”, *Journal for EuroLinguistics* 13: 11–17. <http://www1.ku-eichstaett.de/SLF/EnglVglSW/ELiX/wurth-161.pdf>
- Yu, Danni (2025), “Towards LLM-Assisted Move Annotation: Leveraging ChatGPT-4 to Analyse the Genre Structure of CEO Statements in Corporate Social Responsibility Reports”, *English for Specific Purposes* 78: 33–49. <https://doi.org/10.1016/j.esp.2024.11.003>
- Zaharopoulos, Thimios (2021), “Does Trump Represent America? The Image of the United States in Greek Online News Sites During the Trump Presidency”, in: Kamalipour, Yahya R. / Hamelink, Cees J. (eds.), *Global Media Perceptions of the United States: The Trump Effect*, 145–154, Lanham: Rowman & Littlefield.

first version received 13 October 2025

revised version received 2 April 2026

## Appendixes

### 1. H1. Distribution of bias.

#### VAR00002

|         | Observed N | Expected N | Residual |
|---------|------------|------------|----------|
| 4921,00 | 4921       | 5469,5     | -548,5   |
| 6018,00 | 6018       | 5469,5     | 548,5    |
| Total   | 10939      |            |          |

#### Test Statistics

##### VAR00002

|             |                      |
|-------------|----------------------|
| Chi-Square  | 110,011 <sup>a</sup> |
| df          | 1                    |
| Asymp. Sig. | <,001                |

a. 0 cells (0,0%)  
have expected  
frequencies less  
than 5. The  
minimum  
expected cell  
frequency is  
5469,5.

## 2. H2. Candidate-specific bias.

**Candidate (1=Trump, 2=Clinton, 3=Biden, 4=Harris) \* Bias (1=Neutral, 2=Positive, 3=Negative)**  
**Crosstabulation**

|                                                   |                                                            |                                                            | Bias (1=Neutral, 2=Positive, 3=Negative) |          |          | Total   |
|---------------------------------------------------|------------------------------------------------------------|------------------------------------------------------------|------------------------------------------|----------|----------|---------|
|                                                   |                                                            |                                                            | Neutral                                  | Positive | Negative |         |
| Candidate (1=Trump, 2=Clinton, 3=Biden, 4=Harris) | Trump                                                      | Count                                                      | 4503                                     | 566      | 3091     | 8160    |
|                                                   |                                                            | Expected Count                                             | 4489,2                                   | 964,5    | 2706,3   | 8160,0  |
|                                                   |                                                            | % within Candidate (1=Trump, 2=Clinton, 3=Biden, 4=Harris) | 55,2%                                    | 6,9%     | 37,9%    | 100,0%  |
|                                                   | Clinton                                                    | Count                                                      | 310                                      | 235      | 162      | 707     |
|                                                   |                                                            | Expected Count                                             | 389,0                                    | 83,6     | 234,5    | 707,0   |
|                                                   |                                                            | % within Candidate (1=Trump, 2=Clinton, 3=Biden, 4=Harris) | 43,8%                                    | 33,2%    | 22,9%    | 100,0%  |
|                                                   | Biden                                                      | Count                                                      | 896                                      | 354      | 281      | 1531    |
|                                                   |                                                            | Expected Count                                             | 842,3                                    | 181,0    | 507,8    | 1531,0  |
|                                                   |                                                            | % within Candidate (1=Trump, 2=Clinton, 3=Biden, 4=Harris) | 58,5%                                    | 23,1%    | 18,4%    | 100,0%  |
|                                                   | Harris                                                     | Count                                                      | 309                                      | 138      | 94       | 541     |
|                                                   |                                                            | Expected Count                                             | 297,6                                    | 63,9     | 179,4    | 541,0   |
|                                                   |                                                            | % within Candidate (1=Trump, 2=Clinton, 3=Biden, 4=Harris) | 57,1%                                    | 25,5%    | 17,4%    | 100,0%  |
| Total                                             | Count                                                      |                                                            | 6018                                     | 1293     | 3628     | 10939   |
|                                                   | Expected Count                                             |                                                            | 6018,0                                   | 1293,0   | 3628,0   | 10939,0 |
|                                                   | % within Candidate (1=Trump, 2=Clinton, 3=Biden, 4=Harris) |                                                            | 55,0%                                    | 11,8%    | 33,2%    | 100,0%  |

### Chi-Square Tests

|                              | Value                | df | Asymptotic Significance (2-sided) |
|------------------------------|----------------------|----|-----------------------------------|
| Pearson Chi-Square           | 929,234 <sup>a</sup> | 6  | <,001                             |
| Likelihood Ratio             | 838,718              | 6  | <,001                             |
| Linear-by-Linear Association | 94,652               | 1  | <,001                             |
| N of Valid Cases             | 10939                |    |                                   |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 63,95.

### 3. H3. Bias and ideological stance of newspapers.

#### Democrat/Republican \* Bias \* Left/Right Crosstabulation

| Left/Right |                     |                      |                              | neutral | Bias<br>positive | negative | Total  |
|------------|---------------------|----------------------|------------------------------|---------|------------------|----------|--------|
| Left       | Democrat/Republican | Democratic candidate | Count                        | 887     | 446              | 283      | 1616   |
|            |                     |                      | % within Democrat/Republican | 54,9%   | 27,6%            | 17,5%    | 100,0% |
|            |                     | Trump                | Count                        | 2663    | 264              | 1923     | 4850   |
|            |                     |                      | % within Democrat/Republican | 54,9%   | 5,4%             | 39,6%    | 100,0% |
|            | Total               |                      | Count                        | 3550    | 710              | 2206     | 6466   |
|            |                     |                      | % within Democrat/Republican | 54,9%   | 11,0%            | 34,1%    | 100,0% |
| Liberal    | Democrat/Republican | Democratic candidate | Count                        | 105     | 29               | 49       | 183    |
|            |                     |                      | % within Democrat/Republican | 57,4%   | 15,8%            | 26,8%    | 100,0% |
|            |                     | Trump                | Count                        | 337     | 44               | 174      | 555    |
|            |                     |                      | % within Democrat/Republican | 60,7%   | 7,9%             | 31,4%    | 100,0% |
|            | Total               |                      | Count                        | 442     | 73               | 223      | 738    |
|            |                     |                      | % within Democrat/Republican | 59,9%   | 9,9%             | 30,2%    | 100,0% |
| Middle     | Democrat/Republican | Democratic candidate | Count                        | 94      | 33               | 26       | 153    |
|            |                     |                      | % within Democrat/Republican | 61,4%   | 21,6%            | 17,0%    | 100,0% |
|            |                     | Trump                | Count                        | 216     | 40               | 125      | 381    |
|            |                     |                      | % within Democrat/Republican | 56,7%   | 10,5%            | 32,8%    | 100,0% |
|            | Total               |                      | Count                        | 310     | 73               | 151      | 534    |
|            |                     |                      | % within Democrat/Republican | 58,1%   | 13,7%            | 28,3%    | 100,0% |
| Right      | Democrat/Republican | Democratic candidate | Count                        | 429     | 219              | 179      | 827    |
|            |                     |                      | % within Democrat/Republican | 51,9%   | 26,5%            | 21,6%    | 100,0% |
|            |                     | Trump                | Count                        | 1287    | 218              | 869      | 2374   |
|            |                     |                      | % within Democrat/Republican | 54,2%   | 9,2%             | 36,6%    | 100,0% |
|            | Total               |                      | Count                        | 1716    | 437              | 1048     | 3201   |
|            |                     |                      | % within Democrat/Republican | 53,6%   | 13,7%            | 32,7%    | 100,0% |
| Total      | Democrat/Republican | Democratic candidate | Count                        | 1515    | 727              | 537      | 2779   |
|            |                     |                      | % within Democrat/Republican | 54,5%   | 26,2%            | 19,3%    | 100,0% |
|            |                     | Trump                | Count                        | 4503    | 566              | 3091     | 8160   |
|            |                     |                      | % within Democrat/Republican | 55,2%   | 6,9%             | 37,9%    | 100,0% |
|            | Total               |                      | Count                        | 6018    | 1293             | 3628     | 10939  |
|            |                     |                      | % within Democrat/Republican | 55,0%   | 11,8%            | 33,2%    | 100,0% |

**Chi-Square Tests**

| Left/Right |                                 | Value                | df | Asymptotic<br>Significance<br>(2-sided) |
|------------|---------------------------------|----------------------|----|-----------------------------------------|
| Left       | Pearson Chi-Square              | 715,980 <sup>b</sup> | 2  | <,001                                   |
|            | Likelihood Ratio                | 652,139              | 2  | <,001                                   |
|            | Linear-by-Linear<br>Association | 70,002               | 1  | <,001                                   |
|            | N of Valid Cases                | 6466                 |    |                                         |
| Liberal    | Pearson Chi-Square              | 9,935 <sup>c</sup>   | 2  | ,007                                    |
|            | Likelihood Ratio                | 9,100                | 2  | ,011                                    |
|            | Linear-by-Linear<br>Association | ,026                 | 1  | ,873                                    |
|            | N of Valid Cases                | 738                  |    |                                         |
| Middle     | Pearson Chi-Square              | 19,864 <sup>d</sup>  | 2  | <,001                                   |
|            | Likelihood Ratio                | 20,074               | 2  | <,001                                   |
|            | Linear-by-Linear<br>Association | 5,946                | 1  | ,015                                    |
|            | N of Valid Cases                | 534                  |    |                                         |
| Right      | Pearson Chi-Square              | 176,991 <sup>e</sup> | 2  | <,001                                   |
|            | Likelihood Ratio                | 163,723              | 2  | <,001                                   |
|            | Linear-by-Linear<br>Association | 11,914               | 1  | <,001                                   |
|            | N of Valid Cases                | 3201                 |    |                                         |
| Total      | Pearson Chi-Square              | 863,549 <sup>a</sup> | 2  | <,001                                   |
|            | Likelihood Ratio                | 793,339              | 2  | <,001                                   |
|            | Linear-by-Linear<br>Association | 79,527               | 1  | <,001                                   |
|            | N of Valid Cases                | 10939                |    |                                         |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 328,48.

b. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 177,45.

c. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 18,10.

d. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 20,92.

e. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 112,90.

#### 4. H4. Stability across electoral cycles.

**Year \* Bias Crosstabulation**

|       |                |                | Bias    |          |          | Total   |
|-------|----------------|----------------|---------|----------|----------|---------|
|       |                |                | neutral | positive | negative |         |
| Year  | 2016           | Count          | 1884    | 382      | 1365     | 3631    |
|       |                | Expected Count | 1997,6  | 429,2    | 1204,2   | 3631,0  |
|       |                | % within Year  | 51,9%   | 10,5%    | 37,6%    | 100,0%  |
|       | 2020           | Count          | 1810    | 417      | 1059     | 3286    |
|       |                | Expected Count | 1807,8  | 388,4    | 1089,8   | 3286,0  |
|       |                | % within Year  | 55,1%   | 12,7%    | 32,2%    | 100,0%  |
|       | 2024           | Count          | 2324    | 494      | 1204     | 4022    |
|       |                | Expected Count | 2212,7  | 475,4    | 1333,9   | 4022,0  |
|       |                | % within Year  | 57,8%   | 12,3%    | 29,9%    | 100,0%  |
| Total | Count          |                | 6018    | 1293     | 3628     | 10939   |
|       | Expected Count |                | 6018,0  | 1293,0   | 3628,0   | 10939,0 |
|       | % within Year  |                | 55,0%   | 11,8%    | 33,2%    | 100,0%  |

**Chi-Square Tests**

|                                 | Value               | df | Asymptotic<br>Significance<br>(2-sided) |
|---------------------------------|---------------------|----|-----------------------------------------|
| Pearson Chi-Square              | 55,066 <sup>a</sup> | 4  | <,001                                   |
| Likelihood Ratio                | 54,787              | 4  | <,001                                   |
| Linear-by-Linear<br>Association | 41,716              | 1  | <,001                                   |
| N of Valid Cases                | 10939               |    |                                         |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 388,41.

## 5. H5. Naming practices.

### Neutral or marked forms

|         | Observed N | Expected N | Residual |
|---------|------------|------------|----------|
| Neutral | 10735      | 5469,5     | 5265,5   |
| Marked  | 204        | 5469,5     | -5265,5  |
| Total   | 10939      |            |          |

### Test Statistics

|             | Neutral or<br>marked forms |
|-------------|----------------------------|
| Chi-Square  | 10138,217 <sup>a</sup>     |
| df          | 1                          |
| Asymp. Sig. | <,001                      |

a. 0 cells (0,0%) have expected frequencies less than 5. The minimum expected cell frequency is 5469,5.

## 6. H6. Value-judgment labels.

### Democrat/Republican \* Value Judgement Crosstabulation

|                     |                      |                              | Value Judgement |           |            | Total   |
|---------------------|----------------------|------------------------------|-----------------|-----------|------------|---------|
|                     |                      |                              | None            | Laudative | Pejorative |         |
| Democrat/Republican | Democratic candidate | Count                        | 2667            | 43        | 69         | 2779    |
|                     |                      | Expected Count               | 2619,2          | 33,5      | 126,3      | 2779,0  |
|                     |                      | % within Democrat/Republican | 96,0%           | 1,5%      | 2,5%       | 100,0%  |
|                     | Trump                | Count                        | 7643            | 89        | 428        | 8160    |
|                     |                      | Expected Count               | 7690,8          | 98,5      | 370,7      | 8160,0  |
|                     |                      | % within Democrat/Republican | 93,7%           | 1,1%      | 5,2%       | 100,0%  |
|                     | Total                | Count                        | 10310           | 132       | 497        | 10939   |
|                     |                      | Expected Count               | 10310,0         | 132,0     | 497,0      | 10939,0 |
|                     |                      | % within Democrat/Republican | 94,2%           | 1,2%      | 4,5%       | 100,0%  |

### Chi-Square Tests

|                              | Value               | df | Asymptotic Significance (2-sided) |
|------------------------------|---------------------|----|-----------------------------------|
| Pearson Chi-Square           | 39,563 <sup>a</sup> | 2  | <,001                             |
| Likelihood Ratio             | 44,113              | 2  | <,001                             |
| Linear-by-Linear Association | 29,057              | 1  | <,001                             |
| N of Valid Cases             | 10939               |    |                                   |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 33,53.

## 7. H7. European unity of bias.

### Detailed Region \* Bias Crosstabulation

|                 |                          |                          | Bias         |          | Total  |
|-----------------|--------------------------|--------------------------|--------------|----------|--------|
|                 |                          |                          | not negative | negative |        |
| Detailed Region | Europe Center            | Count                    | 400          | 238      | 638    |
|                 |                          | Expected Count           | 383,3        | 254,7    | 638,0  |
|                 |                          | % within Detailed Region | 62,7%        | 37,3%    | 100,0% |
|                 | Europe East              | Count                    | 492          | 221      | 713    |
|                 |                          | Expected Count           | 428,4        | 284,6    | 713,0  |
|                 |                          | % within Detailed Region | 69,0%        | 31,0%    | 100,0% |
|                 | Europe North             | Count                    | 1264         | 1069     | 2333   |
|                 |                          | Expected Count           | 1401,7       | 931,3    | 2333,0 |
|                 |                          | % within Detailed Region | 54,2%        | 45,8%    | 100,0% |
|                 | Europe South             | Count                    | 393          | 243      | 636    |
|                 |                          | Expected Count           | 382,1        | 253,9    | 636,0  |
|                 |                          | % within Detailed Region | 61,8%        | 38,2%    | 100,0% |
|                 | Europe West              | Count                    | 383          | 177      | 560    |
|                 |                          | Expected Count           | 336,5        | 223,5    | 560,0  |
|                 |                          | % within Detailed Region | 68,4%        | 31,6%    | 100,0% |
| Total           | Count                    |                          | 2932         | 1948     | 4880   |
|                 | Expected Count           |                          | 2932,0       | 1948,0   | 4880,0 |
|                 | % within Detailed Region |                          | 60,1%        | 39,9%    | 100,0% |

### Chi-Square Tests

|                                 | Value               | df | Asymptotic<br>Significance<br>(2-sided) |
|---------------------------------|---------------------|----|-----------------------------------------|
| Pearson Chi-Square              | 76,281 <sup>a</sup> | 4  | <,001                                   |
| Likelihood Ratio                | 77,065              | 4  | <,001                                   |
| Linear-by-Linear<br>Association | ,033                | 1  | ,855                                    |
| N of Valid Cases                | 4880                |    |                                         |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 223,54.

### Symmetric Measures

|                    |            | Value | Approximate<br>Significance |
|--------------------|------------|-------|-----------------------------|
| Nominal by Nominal | Phi        | ,125  | <,001                       |
|                    | Cramer's V | ,125  | <,001                       |
| N of Valid Cases   |            | 4880  |                             |

## 8. H8. Europe versus the rest of the world.

### Region \* Bias Crosstabulation

|        |                 | Bias            |          | Total  |
|--------|-----------------|-----------------|----------|--------|
|        |                 | Not negative    | Negative |        |
| Region | Europe          | Count           | 2932     | 1948   |
|        |                 | Expected Count  | 3031,5   | 1848,5 |
|        |                 | % within Region | 60,1%    | 39,9%  |
|        | Not Europe      | Count           | 2137     | 1143   |
|        |                 | Expected Count  | 2037,5   | 1242,5 |
|        |                 | % within Region | 65,2%    | 34,8%  |
| Total  | Count           |                 | 5069     | 3091   |
|        | Expected Count  |                 | 5069,0   | 3091,0 |
|        | % within Region |                 | 62,1%    | 37,9%  |

### Chi-Square Tests

|                                    | Value               | df | Asymptotic<br>Significance<br>(2-sided) | Exact Sig. (2-<br>sided) | Exact Sig. (1-<br>sided) |
|------------------------------------|---------------------|----|-----------------------------------------|--------------------------|--------------------------|
| Pearson Chi-Square                 | 21,432 <sup>a</sup> | 1  | <,001                                   |                          |                          |
| Continuity Correction <sup>b</sup> | 21,217              | 1  | <,001                                   |                          |                          |
| Likelihood Ratio                   | 21,519              | 1  | <,001                                   |                          |                          |
| Fisher's Exact Test                |                     |    |                                         | <,001                    | <,001                    |
| Linear-by-Linear<br>Association    | 21,429              | 1  | <,001                                   |                          |                          |
| N of Valid Cases                   | 8160                |    |                                         |                          |                          |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 1242,46.

b. Computed only for a 2x2 table

### Symmetric Measures

|                    |            | Value | Approximate<br>Significance |
|--------------------|------------|-------|-----------------------------|
| Nominal by Nominal | Phi        | -,051 | <,001                       |
|                    | Cramer's V | ,051  | <,001                       |
| N of Valid Cases   |            | 8160  |                             |

### Risk Estimate

|                                                | Value | 95% Confidence Interval |       |
|------------------------------------------------|-------|-------------------------|-------|
|                                                |       | Lower                   | Upper |
| Odds Ratio for Region<br>(Europe / Not Europe) | ,805  | ,734                    | ,883  |
| For cohort Bias = Not<br>negative              | ,922  | ,891                    | ,954  |
| For cohort Bias = Negative                     | 1,146 | 1,081                   | 1,214 |
| N of Valid Cases                               | 8160  |                         |       |

## 9. H9. Cross-linguistic convergence.

**Language Family \* Bias Crosstabulation**

|                 |          |                          | neutral | Bias<br>positive | negative | Total  |
|-----------------|----------|--------------------------|---------|------------------|----------|--------|
| Language Family | Slavic   | Count                    | 574     | 106              | 275      | 955    |
|                 |          | Expected Count           | 508,2   | 112,5            | 334,3    | 955,0  |
|                 |          | % within Language Family | 60,1%   | 11,1%            | 28,8%    | 100,0% |
|                 |          | Adjusted Residual        | 4,6     | -,7              | -4,3     |        |
|                 | Germanic | Count                    | 2269    | 536              | 1653     | 4458   |
|                 |          | Expected Count           | 2372,2  | 525,1            | 1560,6   | 4458,0 |
|                 |          | % within Language Family | 50,9%   | 12,0%            | 37,1%    | 100,0% |
|                 |          | Adjusted Residual        | -5,4    | ,9               | 5,1      |        |
|                 | Romance  | Count                    | 685     | 139              | 393      | 1217   |
|                 |          | Expected Count           | 647,6   | 143,4            | 426,0    | 1217,0 |
|                 |          | % within Language Family | 56,3%   | 11,4%            | 32,3%    | 100,0% |
|                 |          | Adjusted Residual        | 2,4     | -,4              | -2,2     |        |
| Total           |          | Count                    | 3528    | 781              | 2321     | 6630   |
|                 |          | Expected Count           | 3528,0  | 781,0            | 2321,0   | 6630,0 |
|                 |          | % within Language Family | 53,2%   | 11,8%            | 35,0%    | 100,0% |

### Chi-Square Tests

|                                 | Value               | df | Asymptotic<br>Significance<br>(2-sided) |
|---------------------------------|---------------------|----|-----------------------------------------|
| Pearson Chi-Square              | 34,464 <sup>a</sup> | 4  | <,001                                   |
| Likelihood Ratio                | 34,800              | 4  | <,001                                   |
| Linear-by-Linear<br>Association | 1,630               | 1  | ,202                                    |
| N of Valid Cases                | 6630                |    |                                         |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 112,50.

### Symmetric Measures

|                    |            | Value | Approximate<br>Significance |
|--------------------|------------|-------|-----------------------------|
| Nominal by Nominal | Phi        | ,072  | <,001                       |
|                    | Cramer's V | ,051  | <,001                       |
| N of Valid Cases   |            | 6630  |                             |

## 10. H10. Regional consistency.

### Detailed Region \* Bias \* Candidate Crosstabulation

| Candidate |                 |               |                          | neutral | Bias<br>positive | negative | Total  |
|-----------|-----------------|---------------|--------------------------|---------|------------------|----------|--------|
| Biden     | Detailed Region | Europe Center | Count                    | 42      | 21               | 18       | 81     |
|           |                 |               | Expected Count           | 47,3    | 17,8             | 15,8     | 81,0   |
|           |                 |               | % within Detailed Region | 51,9%   | 25,9%            | 22,2%    | 100,0% |
|           |                 |               | Adjusted Residual        | -1,3    | ,9               | ,6       |        |
|           |                 | Europe East   | Count                    | 114     | 28               | 35       | 177    |
|           |                 |               | Expected Count           | 103,4   | 39,0             | 34,6     | 177,0  |
|           |                 |               | % within Detailed Region | 64,4%   | 15,8%            | 19,8%    | 100,0% |
|           |                 |               | Adjusted Residual        | 1,8     | -2,2             | ,1       |        |
|           |                 | Europe North  | Count                    | 283     | 114              | 97       | 494    |
|           |                 |               | Expected Count           | 288,7   | 108,7            | 96,6     | 494,0  |
|           |                 |               | % within Detailed Region | 57,3%   | 23,1%            | 19,6%    | 100,0% |
|           |                 |               | Adjusted Residual        | -,7     | ,8               | ,1       |        |
|           |                 | Europe South  | Count                    | 62      | 28               | 23       | 113    |
|           |                 |               | Expected Count           | 66,0    | 24,9             | 22,1     | 113,0  |
|           |                 |               | % within Detailed Region | 54,9%   | 24,8%            | 20,4%    | 100,0% |
|           |                 |               | Adjusted Residual        | -,8     | ,8               | ,2       |        |
|           |                 | Europe West   | Count                    | 70      | 24               | 18       | 112    |
|           |                 |               | Expected Count           | 65,5    | 24,6             | 21,9     | 112,0  |
|           |                 |               | % within Detailed Region | 62,5%   | 21,4%            | 16,1%    | 100,0% |
|           |                 |               | Adjusted Residual        | ,9      | -,2              | -1,0     |        |
|           |                 | Total         | Count                    | 571     | 215              | 191      | 977    |
|           |                 |               | Expected Count           | 571,0   | 215,0            | 191,0    | 977,0  |
|           |                 |               | % within Detailed Region | 58,4%   | 22,0%            | 19,5%    | 100,0% |
| Clinton   | Detailed Region | Europe Center | Count                    | 25      | 22               | 19       | 66     |
|           |                 |               | Expected Count           | 27,4    | 21,5             | 17,1     | 66,0   |
|           |                 |               | % within Detailed Region | 37,9%   | 33,3%            | 28,8%    | 100,0% |
|           |                 |               | Adjusted Residual        | -,7     | ,1               | ,6       |        |
|           |                 | Europe East   | Count                    | 1       | 0                | 0        | 1      |
|           |                 |               | Expected Count           | ,4      | ,3               | ,3       | 1,0    |
|           |                 |               | % within Detailed Region | 100,0%  | 0,0%             | 0,0%     | 100,0% |
|           |                 |               | Adjusted Residual        | 1,2     | -,7              | -,6      |        |
|           |                 | Europe North  | Count                    | 89      | 82               | 64       | 235    |
|           |                 |               | Expected Count           | 97,6    | 76,6             | 60,7     | 235,0  |
|           |                 |               | % within Detailed Region | 37,9%   | 34,9%            | 27,2%    | 100,0% |
|           |                 |               | Adjusted Residual        | -1,7    | 1,1              | ,7       |        |
|           |                 | Europe South  | Count                    | 34      | 17               | 12       | 63     |
|           |                 |               | Expected Count           | 26,2    | 20,5             | 16,3     | 63,0   |
|           |                 |               | % within Detailed Region | 54,0%   | 27,0%            | 19,0%    | 100,0% |
|           |                 |               | Adjusted Residual        | 2,2     | -1,0             | -1,3     |        |
|           |                 | Europe West   | Count                    | 23      | 14               | 12       | 49     |
|           |                 |               | Expected Count           | 20,4    | 16,0             | 12,7     | 49,0   |
|           |                 |               | % within Detailed Region | 46,9%   | 28,6%            | 24,5%    | 100,0% |
|           |                 |               | Adjusted Residual        | ,8      | -,6              | -,2      |        |
|           |                 | Total         | Count                    | 172     | 135              | 107      | 414    |
|           |                 |               | Expected Count           | 172,0   | 135,0            | 107,0    | 414,0  |
|           |                 |               | % within Detailed Region | 41,5%   | 32,6%            | 25,8%    | 100,0% |

|        |                 |               |                          |        |       |        |        |
|--------|-----------------|---------------|--------------------------|--------|-------|--------|--------|
| Harris | Detailed Region | Europe Center | Count                    | 9      | 12    | 5      | 26     |
|        |                 |               | Expected Count           | 13,8   | 6,7   | 5,4    | 26,0   |
|        |                 |               | % within Detailed Region | 34,6%  | 46,2% | 19,2%  | 100,0% |
|        |                 |               | Adjusted Residual        | -2,0   | 2,4   | -,2    |        |
|        |                 | Europe East   | Count                    | 27     | 18    | 19     | 64     |
|        |                 |               | Expected Count           | 34,1   | 16,6  | 13,4   | 64,0   |
|        |                 |               | % within Detailed Region | 42,2%  | 28,1% | 29,7%  | 100,0% |
|        |                 |               | Adjusted Residual        | -1,9   | ,4    | 1,9    |        |
|        |                 | Europe North  | Count                    | 102    | 42    | 37     | 181    |
|        |                 |               | Expected Count           | 96,3   | 46,9  | 37,8   | 181,0  |
|        |                 |               | % within Detailed Region | 56,4%  | 23,2% | 20,4%  | 100,0% |
|        |                 |               | Adjusted Residual        | 1,2    | -1,2  | -,2    |        |
|        |                 | Europe South  | Count                    | 20     | 12    | 8      | 40     |
|        |                 |               | Expected Count           | 21,3   | 10,4  | 8,4    | 40,0   |
|        |                 |               | % within Detailed Region | 50,0%  | 30,0% | 20,0%  | 100,0% |
|        |                 |               | Adjusted Residual        | -,4    | ,6    | -,1    |        |
|        |                 | Europe West   | Count                    | 33     | 9     | 6      | 48     |
|        |                 |               | Expected Count           | 25,5   | 12,4  | 10,0   | 48,0   |
|        |                 |               | % within Detailed Region | 68,8%  | 18,8% | 12,5%  | 100,0% |
|        |                 |               | Adjusted Residual        | 2,3    | -1,2  | -1,5   |        |
|        |                 | Total         | Count                    | 191    | 93    | 75     | 359    |
|        |                 |               | Expected Count           | 191,0  | 93,0  | 75,0   | 359,0  |
|        |                 |               | % within Detailed Region | 53,2%  | 25,9% | 20,9%  | 100,0% |
| Trump  | Detailed Region | Europe Center | Count                    | 353    | 47    | 238    | 638    |
|        |                 |               | Expected Count           | 339,1  | 44,2  | 254,7  | 638,0  |
|        |                 |               | % within Detailed Region | 55,3%  | 7,4%  | 37,3%  | 100,0% |
|        |                 |               | Adjusted Residual        | 1,2    | ,5    | -1,4   |        |
|        |                 | Europe East   | Count                    | 432    | 60    | 221    | 713    |
|        |                 |               | Expected Count           | 379,0  | 49,4  | 284,6  | 713,0  |
|        |                 |               | % within Detailed Region | 60,6%  | 8,4%  | 31,0%  | 100,0% |
|        |                 |               | Adjusted Residual        | 4,3    | 1,7   | -5,3   |        |
|        |                 | Europe North  | Count                    | 1137   | 127   | 1069   | 2333   |
|        |                 |               | Expected Count           | 1240,1 | 161,6 | 931,3  | 2333,0 |
|        |                 |               | % within Detailed Region | 48,7%  | 5,4%  | 45,8%  | 100,0% |
|        |                 |               | Adjusted Residual        | -5,9   | -3,9  | 8,1    |        |
|        |                 | Europe South  | Count                    | 345    | 48    | 243    | 636    |
|        |                 |               | Expected Count           | 338,1  | 44,1  | 253,9  | 636,0  |
|        |                 |               | % within Detailed Region | 54,2%  | 7,5%  | 38,2%  | 100,0% |
|        |                 |               | Adjusted Residual        | ,6     | ,7    | -,9    |        |
|        |                 | Europe West   | Count                    | 327    | 56    | 177    | 560    |
|        |                 |               | Expected Count           | 297,7  | 38,8  | 223,5  | 560,0  |
|        |                 |               | % within Detailed Region | 58,4%  | 10,0% | 31,6%  | 100,0% |
|        |                 |               | Adjusted Residual        | 2,6    | 3,0   | -4,3   |        |
|        |                 | Total         | Count                    | 2594   | 338   | 1948   | 4880   |
|        |                 |               | Expected Count           | 2594,0 | 338,0 | 1948,0 | 4880,0 |
|        |                 |               | % within Detailed Region | 53,2%  | 6,9%  | 39,9%  | 100,0% |

|       |                 |               |                          |        |       |        |        |
|-------|-----------------|---------------|--------------------------|--------|-------|--------|--------|
| Total | Detailed Region | Europe Center | Count                    | 429    | 102   | 280    | 811    |
|       |                 |               | Expected Count           | 431,6  | 95,5  | 283,9  | 811,0  |
|       |                 |               | % within Detailed Region | 52,9%  | 12,6% | 34,5%  | 100,0% |
|       |                 |               | Adjusted Residual        | -,2    | ,8    | -,3    |        |
|       |                 | Europe East   | Count                    | 574    | 106   | 275    | 955    |
|       |                 |               | Expected Count           | 508,2  | 112,5 | 334,3  | 955,0  |
|       |                 |               | % within Detailed Region | 60,1%  | 11,1% | 28,8%  | 100,0% |
|       |                 |               | Adjusted Residual        | 4,6    | -,7   | -4,3   |        |
|       |                 | Europe North  | Count                    | 1611   | 365   | 1267   | 3243   |
|       |                 |               | Expected Count           | 1725,7 | 382,0 | 1135,3 | 3243,0 |
|       |                 |               | % within Detailed Region | 49,7%  | 11,3% | 39,1%  | 100,0% |
|       |                 |               | Adjusted Residual        | -5,6   | -1,3  | 6,8    |        |
|       |                 | Europe South  | Count                    | 461    | 105   | 286    | 852    |
|       |                 |               | Expected Count           | 453,4  | 100,4 | 298,3  | 852,0  |
|       |                 |               | % within Detailed Region | 54,1%  | 12,3% | 33,6%  | 100,0% |
|       |                 |               | Adjusted Residual        | ,6     | ,5    | -,9    |        |
|       |                 | Europe West   | Count                    | 453    | 103   | 213    | 769    |
|       |                 |               | Expected Count           | 409,2  | 90,6  | 269,2  | 769,0  |
|       |                 |               | % within Detailed Region | 58,9%  | 13,4% | 27,7%  | 100,0% |
|       |                 |               | Adjusted Residual        | 3,4    | 1,5   | -4,5   |        |
|       | Total           |               | Count                    | 3528   | 781   | 2321   | 6630   |
|       |                 |               | Expected Count           | 3528,0 | 781,0 | 2321,0 | 6630,0 |
|       |                 |               | % within Detailed Region | 53,2%  | 11,8% | 35,0%  | 100,0% |

### Symmetric Measures

| Candidate |                    |                  | Value | Approximate Significance |
|-----------|--------------------|------------------|-------|--------------------------|
| Biden     | Nominal by Nominal | Phi              | ,089  | ,463                     |
|           |                    | Cramer's V       | ,063  | ,463                     |
|           |                    | N of Valid Cases | 977   |                          |
| Clinton   | Nominal by Nominal | Phi              | ,138  | ,447                     |
|           |                    | Cramer's V       | ,097  | ,447                     |
|           |                    | N of Valid Cases | 414   |                          |
| Harris    | Nominal by Nominal | Phi              | ,209  | ,046                     |
|           |                    | Cramer's V       | ,148  | ,046                     |
|           |                    | N of Valid Cases | 359   |                          |
| Trump     | Nominal by Nominal | Phi              | ,131  | <,001                    |
|           |                    | Cramer's V       | ,092  | <,001                    |
|           |                    | N of Valid Cases | 4880  |                          |
| Total     | Nominal by Nominal | Phi              | ,097  | <,001                    |
|           |                    | Cramer's V       | ,069  | <,001                    |
|           |                    | N of Valid Cases | 6630  |                          |

**Chi-Square Tests**

| Candidate |                              | Value               | df | Asymptotic Significance (2-sided) |
|-----------|------------------------------|---------------------|----|-----------------------------------|
| Biden     | Pearson Chi-Square           | 7,701 <sup>b</sup>  | 8  | ,463                              |
|           | Likelihood Ratio             | 8,008               | 8  | ,433                              |
|           | Linear-by-Linear Association | ,421                | 1  | ,516                              |
|           | N of Valid Cases             | 977                 |    |                                   |
| Clinton   | Pearson Chi-Square           | 7,868 <sup>c</sup>  | 8  | ,447                              |
|           | Likelihood Ratio             | 8,175               | 8  | ,417                              |
|           | Linear-by-Linear Association | 2,083               | 1  | ,149                              |
|           | N of Valid Cases             | 414                 |    |                                   |
| Harris    | Pearson Chi-Square           | 15,753 <sup>d</sup> | 8  | ,046                              |
|           | Likelihood Ratio             | 15,236              | 8  | ,055                              |
|           | Linear-by-Linear Association | 7,445               | 1  | ,006                              |
|           | N of Valid Cases             | 359                 |    |                                   |
| Trump     | Pearson Chi-Square           | 83,274 <sup>e</sup> | 8  | <,001                             |
|           | Likelihood Ratio             | 83,646              | 8  | <,001                             |
|           | Linear-by-Linear Association | ,012                | 1  | ,913                              |
|           | N of Valid Cases             | 4880                |    |                                   |
| Total     | Pearson Chi-Square           | 62,563 <sup>a</sup> | 8  | <,001                             |
|           | Likelihood Ratio             | 63,213              | 8  | <,001                             |
|           | Linear-by-Linear Association | 1,228               | 1  | ,268                              |
|           | N of Valid Cases             | 6630                |    |                                   |

- a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 90,59.
- b. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 15,84.
- c. 3 cells (20,0%) have expected count less than 5. The minimum expected count is ,26.
- d. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 5,43.
- e. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 38,79.

# 11. H11. European ideological distinctiveness.

## Left/Right \* Bias \* Region Crosstabulation

| Region     |            |       |                     | neutral | Bias<br>positive | negative | Total  |
|------------|------------|-------|---------------------|---------|------------------|----------|--------|
| Europe     | Left/Right | Left  | Count               | 1280    | 125              | 1134     | 2539   |
|            |            |       | Expected Count      | 1325,0  | 172,6            | 1041,4   | 2539,0 |
|            |            |       | % within Left/Right | 50,4%   | 4,9%             | 44,7%    | 100,0% |
|            | Right      | Right | Count               | 977     | 169              | 640      | 1786   |
|            |            |       | Expected Count      | 932,0   | 121,4            | 732,6    | 1786,0 |
|            |            |       | % within Left/Right | 54,7%   | 9,5%             | 35,8%    | 100,0% |
|            | Total      | Total | Count               | 2257    | 294              | 1774     | 4325   |
|            |            |       | Expected Count      | 2257,0  | 294,0            | 1774,0   | 4325,0 |
|            |            |       | % within Left/Right | 52,2%   | 6,8%             | 41,0%    | 100,0% |
| Not Europe | Left/Right | Left  | Count               | 790     | 65               | 420      | 1275   |
|            |            |       | Expected Count      | 770,6   | 62,5             | 441,8    | 1275,0 |
|            |            |       | % within Left/Right | 62,0%   | 5,1%             | 32,9%    | 100,0% |
|            | Right      | Right | Count               | 171     | 13               | 131      | 315    |
|            |            |       | Expected Count      | 190,4   | 15,5             | 109,2    | 315,0  |
|            |            |       | % within Left/Right | 54,3%   | 4,1%             | 41,6%    | 100,0% |
|            | Total      | Total | Count               | 961     | 78               | 551      | 1590   |
|            |            |       | Expected Count      | 961,0   | 78,0             | 551,0    | 1590,0 |
|            |            |       | % within Left/Right | 60,4%   | 4,9%             | 34,7%    | 100,0% |
| Total      | Left/Right | Left  | Count               | 2070    | 190              | 1554     | 3814   |
|            |            |       | Expected Count      | 2075,0  | 239,9            | 1499,2   | 3814,0 |
|            |            |       | % within Left/Right | 54,3%   | 5,0%             | 40,7%    | 100,0% |
|            | Right      | Right | Count               | 1148    | 182              | 771      | 2101   |
|            |            |       | Expected Count      | 1143,0  | 132,1            | 825,8    | 2101,0 |
|            |            |       | % within Left/Right | 54,6%   | 8,7%             | 36,7%    | 100,0% |
|            | Total      | Total | Count               | 3218    | 372              | 2325     | 5915   |
|            |            |       | Expected Count      | 3218,0  | 372,0            | 2325,0   | 5915,0 |
|            |            |       | % within Left/Right | 54,4%   | 6,3%             | 39,3%    | 100,0% |

**Chi-Square Tests**

| Region     |                              | Value               | df | Asymptotic Significance (2-sided) |
|------------|------------------------------|---------------------|----|-----------------------------------|
| Europe     | Pearson Chi-Square           | 55,404 <sup>b</sup> | 2  | <,001                             |
|            | Likelihood Ratio             | 55,038              | 2  | <,001                             |
|            | Linear-by-Linear Association | 19,618              | 1  | <,001                             |
|            | N of Valid Cases             | 4325                |    |                                   |
| Not Europe | Pearson Chi-Square           | 8,396 <sup>c</sup>  | 2  | ,015                              |
|            | Likelihood Ratio             | 8,237               | 2  | ,016                              |
|            | Linear-by-Linear Association | 7,603               | 1  | ,006                              |
|            | N of Valid Cases             | 1590                |    |                                   |
| Total      | Pearson Chi-Square           | 34,866 <sup>a</sup> | 2  | <,001                             |
|            | Likelihood Ratio             | 33,728              | 2  | <,001                             |
|            | Linear-by-Linear Association | 2,887               | 1  | ,089                              |
|            | N of Valid Cases             | 5915                |    |                                   |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 132,13.

b. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 121,41.

c. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 15,45.

**Symmetric Measures**

| Region     |                    |            | Value | Approximate Significance |
|------------|--------------------|------------|-------|--------------------------|
| Europe     | Nominal by Nominal | Phi        | ,113  | <,001                    |
|            |                    | Cramer's V | ,113  | <,001                    |
|            | N of Valid Cases   |            | 4325  |                          |
| Not Europe | Nominal by Nominal | Phi        | ,073  | ,015                     |
|            |                    | Cramer's V | ,073  | ,015                     |
|            | N of Valid Cases   |            | 1590  |                          |
| Total      | Nominal by Nominal | Phi        | ,077  | <,001                    |
|            |                    | Cramer's V | ,077  | <,001                    |
|            | N of Valid Cases   |            | 5915  |                          |

## 12. H12. Unified naming conventions.

### Region \* Neutral or marked forms Crosstabulation

|        |                 |                 | Neutral or marked forms |        |        |
|--------|-----------------|-----------------|-------------------------|--------|--------|
|        |                 |                 | Neutral                 | Marked | Total  |
| Region | Europe          | Count           | 6489                    | 141    | 6630   |
|        |                 | Expected Count  | 6498,6                  | 131,4  | 6630,0 |
|        |                 | % within Region | 97,9%                   | 2,1%   | 100,0% |
|        | Not Europe      | Count           | 1475                    | 20     | 1495   |
|        |                 | Expected Count  | 1465,4                  | 29,6   | 1495,0 |
|        |                 | % within Region | 98,7%                   | 1,3%   | 100,0% |
| Total  | Count           | 7964            | 161                     | 8125   |        |
|        | Expected Count  | 7964,0          | 161,0                   | 8125,0 |        |
|        | % within Region | 98,0%           | 2,0%                    | 100,0% |        |

### Chi-Square Tests

|                                    | Value              | df | Asymptotic<br>Significance<br>(2-sided) | Exact Sig. (2-<br>sided) | Exact Sig. (1-<br>sided) |
|------------------------------------|--------------------|----|-----------------------------------------|--------------------------|--------------------------|
| Pearson Chi-Square                 | 3,909 <sup>a</sup> | 1  | ,048                                    |                          |                          |
| Continuity Correction <sup>b</sup> | 3,513              | 1  | ,061                                    |                          |                          |
| Likelihood Ratio                   | 4,300              | 1  | ,038                                    |                          |                          |
| Fisher's Exact Test                |                    |    |                                         | ,051                     | ,026                     |
| Linear-by-Linear<br>Association    | 3,909              | 1  | ,048                                    |                          |                          |
| N of Valid Cases                   | 8125               |    |                                         |                          |                          |

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 29,62.

b. Computed only for a 2x2 table

### Risk Estimate

|                                                 | Value | 95% Confidence Interval |       |
|-------------------------------------------------|-------|-------------------------|-------|
|                                                 |       | Lower                   | Upper |
| Odds Ratio for Region<br>(Europe / Not Europe)  | ,624  | ,389                    | 1,000 |
| For cohort Neutral or<br>marked forms = Neutral | ,992  | ,985                    | ,999  |
| For cohort Neutral or<br>marked forms = Marked  | 1,590 | ,999                    | 2,531 |
| N of Valid Cases                                | 8125  |                         |       |

**13. Counts**

| <b>Year/Candidate</b> | <b>Headlines</b> |
|-----------------------|------------------|
| <b>2016</b>           | 3,695            |
| <b>Clinton</b>        | 734              |
| <b>Trump</b>          | 2,961            |
| <b>2020</b>           | 3,335            |
| <b>Biden</b>          | 1,269            |
| <b>Trump</b>          | 2,066            |
| <b>2024</b>           | 4,066            |
| <b>Biden</b>          | 265              |
| <b>Harris</b>         | 543              |
| <b>Trump</b>          | 3,258            |
| <b>Total</b>          | <b>11,096</b>    |

**Democrat/Republican**

|                      | Observed N | Expected N | Residual |
|----------------------|------------|------------|----------|
| Democratic candidate | 2779       | 5469,5     | -2690,5  |
| Trump                | 8160       | 5469,5     | 2690,5   |
| Total                | 10939      |            |          |

**Test Statistics**

|             | Democrat/Rep<br>ublican |
|-------------|-------------------------|
| Chi-Square  | 2646,966 <sup>a</sup>   |
| df          | 1                       |
| Asymp. Sig. | <,001                   |

a. 0 cells (0,0%) have expected frequencies less than 5. The minimum expected cell frequency is 5469,5.

| <b>Year/Event/Candidate</b>             | <b>Headlines</b> |
|-----------------------------------------|------------------|
| <b>2016 Debate 1 Democrat</b>           | 176              |
| <b>2016 Debate 1 Republican</b>         | 224              |
| <b>2016 Debate 2 Democrat</b>           | 126              |
| <b>2016 Debate 2 Republican</b>         | 384              |
| <b>2016 Debate 3 Democrat</b>           | 121              |
| <b>2016 Debate 3 Republican</b>         | 328              |
| <b>2016 Election Day Democrat</b>       | 295              |
| <b>2016 Election Day Republican</b>     | 1,097            |
| <b>2016 Inauguration Day Democrat</b>   | 16               |
| <b>2016 Inauguration Day Republican</b> | 928              |
| <b>2020 Debate 1 Democrat</b>           | 135              |
| <b>2020 Debate 1 Republican</b>         | 417              |
| <b>2020 Debate 2 Democrat</b>           | 102              |
| <b>2020 Debate 2 Republican</b>         | 284              |
| <b>2020 Debate 3 Democrat</b>           | 145              |
| <b>2020 Debate 3 Republican</b>         | 310              |
| <b>2020 Election Day Democrat</b>       | 345              |
| <b>2020 Election Day Republican</b>     | 569              |
| <b>2020 Inauguration Day Democrat</b>   | 542              |
| <b>2020 Inauguration Day Republican</b> | 486              |
| <b>2024 Debate 1 Democrat</b>           | 265              |
| <b>2024 Debate 1 Republican</b>         | 213              |
| <b>2024 Debate 2 Democrat</b>           | 267              |
| <b>2024 Debate 2 Republican</b>         | 375              |
| <b>2024 Election Day Democrat</b>       | 268              |
| <b>2024 Election Day Republican</b>     | 1,202            |
| <b>2024 Inauguration Day Democrat</b>   | 8                |
| <b>2024 Inauguration Day Republican</b> | 1,468            |
| <b>Total</b>                            | <b>11,096</b>    |

| <b>Newspaper</b>                | <b>2016</b>  | <b>2020</b>  | <b>2024</b>  | <b>Total</b>  |
|---------------------------------|--------------|--------------|--------------|---------------|
| <i>China Daily</i>              | 125          | 33           | 18           | <b>176</b>    |
| <i>Corriere della Sera</i>      | 64           | 54           | 63           | <b>181</b>    |
| <i>De Telegraaf</i>             | 59           | 62           | 53           | <b>174</b>    |
| <i>De Volkskrant</i>            | 84           | 81           | 71           | <b>236</b>    |
| <i>Die Welt</i>                 | 125          | 70           | 56           | <b>251</b>    |
| <i>Dnevnik</i>                  |              | 75           | 142          | <b>217</b>    |
| <i>El Mundo</i>                 | 81           | 71           | 90           | <b>242</b>    |
| <i>El Pais</i>                  | 153          | 107          | 119          | <b>379</b>    |
| <i>El Universal</i>             | 236          | 100          | 546          | <b>882</b>    |
| <i>Folha de Sao Paulo</i>       |              | 170          | 462          | <b>632</b>    |
| <i>Gazeta Wyborcza</i>          | 5            | 313          | 430          | <b>748</b>    |
| <i>Irish Independent</i>        | 141          | 141          | 127          | <b>409</b>    |
| <i>La Stampa</i>                | 61           |              |              | <b>61</b>     |
| <i>Le Figaro</i>                | 99           | 88           | 86           | <b>273</b>    |
| <i>Libération</i>               |              | 53           | 49           | <b>102</b>    |
| <i>Neue Zürcher Zeitung</i>     | 145          | 74           | 53           | <b>272</b>    |
| <i>O Globo</i>                  | 31           | 54           | 11           | <b>96</b>     |
| <i>Reforma</i>                  | 237          | 67           | 89           | <b>393</b>    |
| <i>South China Morning Post</i> | 56           | 38           | 33           | <b>127</b>    |
| <i>Tages-Anzeiger</i>           | 37           | 34           | 32           | <b>103</b>    |
| <i>taz, die tageszeitung</i>    | 80           | 48           | 72           | <b>200</b>    |
| <i>The Age</i>                  | 88           | 64           | 68           | <b>220</b>    |
| <i>The Australian</i>           | 100          | 91           | 91           | <b>282</b>    |
| <i>The Guardian</i>             | 682          | 664          | 688          | <b>2,034</b>  |
| <i>The Irish Times</i>          | 163          | 118          | 108          | <b>389</b>    |
| <i>The New York Times</i>       | 474          | 341          | 283          | <b>1,098</b>  |
| <i>The Times</i>                | 166          | 173          | 125          | <b>464</b>    |
| <i>USA Today</i>                | 188          | 142          | 88           | <b>418</b>    |
| <i>Vanguard</i>                 | 15           | 9            | 13           | <b>37</b>     |
| <b>Total</b>                    | <b>3,695</b> | <b>3,335</b> | <b>4,066</b> | <b>11,096</b> |

**14. Bias**

| <b>Year/Candidate</b> | <b>exclude</b> | <b>negative</b> | <b>neutral</b> | <b>positive</b> | <b>Total</b>  |
|-----------------------|----------------|-----------------|----------------|-----------------|---------------|
| <b>2016</b>           | 64             | 1,365           | 1,884          | 382             | <b>3,695</b>  |
| <b>Clinton</b>        | 27             | 162             | 310            | 235             | <b>734</b>    |
| <b>Trump</b>          | 37             | 1,203           | 1574           | 147             | <b>2,961</b>  |
| <b>2020</b>           | 49             | 1,059           | 1,810          | 417             | <b>3,335</b>  |
| <b>Biden</b>          | 2              | 154             | 776            | 337             | <b>1,269</b>  |
| <b>Trump</b>          | 47             | 905             | 1034           | 80              | <b>2,066</b>  |
| <b>2024</b>           | 44             | 1,204           | 2,324          | 494             | <b>4,066</b>  |
| <b>Biden</b>          | 1              | 127             | 120            | 17              | <b>265</b>    |
| <b>Harris</b>         | 2              | 94              | 309            | 138             | <b>543</b>    |
| <b>Trump</b>          | 41             | 983             | 1,895          | 339             | <b>3,258</b>  |
| <b>Total</b>          | <b>157</b>     | <b>3,628</b>    | <b>6,018</b>   | <b>1,293</b>    | <b>11,096</b> |

| <b>Year/Candidate</b> | <b>negative</b> | <b>neutral</b> | <b>positive</b> |
|-----------------------|-----------------|----------------|-----------------|
| <b>2016</b>           | 37,59%          | 51,89%         | 10,52%          |
| <b>Clinton</b>        | 22,91%          | 43,85%         | 33,24%          |
| <b>Trump</b>          | 41,14%          | 53,83%         | 5,03%           |
| <b>2020</b>           | 32,23%          | 55,08%         | 12,69%          |
| <b>Biden</b>          | 12,15%          | 61,25%         | 26,60%          |
| <b>Trump</b>          | 44,82%          | 51,21%         | 3,96%           |
| <b>2024</b>           | 29,94%          | 57,78%         | 12,28%          |
| <b>Biden</b>          | 48,11%          | 45,45%         | 6,44%           |
| <b>Harris</b>         | 17,38%          | 57,12%         | 25,51%          |
| <b>Trump</b>          | 30,56%          | 58,91%         | 10,54%          |
| <b>Total</b>          | <b>33,17%</b>   | <b>55,01%</b>  | <b>11,82%</b>   |

**15. Value Judgement**

| <b>Year/Candidate</b> | <b>laudative</b> | <b>none</b>   | <b>pejorative</b> | <b>Total</b>  |
|-----------------------|------------------|---------------|-------------------|---------------|
| <b>2016</b>           | 33               | 3,492         | 170               | 3,695         |
| <b>Clinton</b>        | 21               | 676           | 37                | 734           |
| <b>Trump</b>          | 12               | 2,816         | 133               | 2,961         |
| <b>2020</b>           | 33               | 3,124         | 178               | 3,335         |
| <b>Biden</b>          | 15               | 1,236         | 18                | 1,269         |
| <b>Trump</b>          | 18               | 1,888         | 160               | 2,066         |
| <b>2024</b>           | 66               | 3,851         | 149               | 4,066         |
| <b>Biden</b>          | 2                | 259           | 4                 | 265           |
| <b>Harris</b>         | 5                | 528           | 10                | 543           |
| <b>Trump</b>          | 59               | 3,064         | 135               | 3,258         |
| <b>Total</b>          | <b>132</b>       | <b>10,467</b> | <b>497</b>        | <b>11,096</b> |